

Aliasing in Fractional Designs

Background

Factorial designs, for example 2^k and 3^k , are discussed elsewhere in this series (*see* **Factorial Experiments**). Fractional factorial designs such as the 2^{k-p} and 3^{k-p} designs are discussed in **Fractional Factorial Designs**, while **Plackett–Burman Designs** discusses Plackett–Burman designs. In this chapter, we discuss the aliasing (confounding or confusion) among effect estimates that occurs with the regular two-level fractional factorial (2^{k-p}) designs, nonregular two-level designs, and the regular three-level fractional factorial (3^{k-p}) designs.

If k factors, each with two levels, are studied by obtaining a single observation at each of the $n = 2^k$ factor-level combinations (design points), then the design is called an *unreplicated 2^k factorial design*. (The two levels of each factor are commonly labeled +1 and –1 with the assignment being arbitrary.) An example of one of the smallest full factorial designs is the 2^3 , that is, $k = 3$ with the factor levels labeled –1 and +1. Table 1 shows the effect (**X**) matrix. The A , B , and C columns represent the design matrix. In other words, the first observation is obtained by setting factors A , B , and C at their –1 or “low” level and observing the response. The other observations are obtained similarly. In practice, the order of experimentation should be randomized in an attempt to balance out any noise, such as a time effect.

In this example there are eight observations so eight quantities can be estimated. The effect matrix

is simply the **X** matrix in a regression setting. The I column represents the intercept or overall mean. The AB , AC , BC , and ABC columns are used to calculate the two- and three-factor interactions which we abbreviate as $2fi$ and $3fi$. (These columns are created by multiplying the columns associated with the component factors element by element.) Note that if all quantities are to be estimated, the model is saturated, and without replication, there is no error term available for **hypothesis testing**.

Aliasing in 2^{k-p} Experiments

In this section we introduce the concept of aliasing in the simplest case, the half fraction or 2^{k-1} , followed by aliasing in the general 2^{k-p} .

Aliasing in 2^{k-1} Experiments

Suppose, instead of the full factorial, a half fraction or 2^{3-1} is used by only running the four rows where $ABC = I$. These rows are marked with an * in Table 1. (Note that a 2^{3-1} design is not commonly used in practice. We include it here for strictly pedagogical reasons.) Inspection of these rows reveals the following equalities: $I = ABC$, $A = BC$, $B = AC$, $C = AB$. Here, $I = ABC$ is denoted the *defining relation* and the pairs are known as *aliases*. We can actually determine the alias pairs by multiplying the defining relation by each effect. When doing so, we note that each column consists of half –1s and half +1s, so if a column is multiplied by itself we get a column full of 1s, or I .

For example, $A^2 = A \cdot A = I$ and $(ABC)^2 = A^2B^2C^2 = I \cdot I \cdot I = I$. Multiplying the defining relation by A gives $A \cdot I = A \cdot ABC$, which simplifies to $A = BC$. The alias pairs $B = AC$ and $C = BD$ are derived similarly. We use this method of determining the alias structure as opposed to writing out the effect matrix and inspecting for equalities. Note that the effect of ABC cannot be separated from the overall mean. This pair of effects, as well as the pairs A and BC , B and AC , and C and AB , are completely *confounded* or confused. In other words, when the experiment is run and the data are analyzed, we will only have four estimates.

Suppose, for example, that the estimate of the alias pair $A = BC$ of effects is quite large and is therefore of interest. Without additional data, we cannot

Table 1 2^3 Full factorial. Rows with * form the 2^{3-1} fraction with $I = ABC$

Row	I	A	B	C	$\cdot$	AB	AC	BC	ABC
1	1	–1	–1	–1	$\cdot$	1	1	1	–1
2*	1	1	–1	–1	$\cdot$	–1	–1	1	1
3*	1	–1	1	–1	$\cdot$	–1	1	–1	1
4	1	1	1	–1	$\cdot$	1	–1	–1	–1
5*	1	–1	–1	1	$\cdot$	1	–1	–1	1
6	1	1	–1	1	$\cdot$	–1	1	–1	–1
7	1	–1	1	1	$\cdot$	–1	–1	1	–1
8*	1	1	1	1	$\cdot$	1	1	1	1

2 Aliasing in Fractional Designs

determine if the effect is due to A , or BC , or both! The only way to structurally remove the confounding is to collect more data. However, if we assume that two-factor and higher-order interactions are negligible, then there is essentially no confounding and we may assume we are estimating the overall mean, A , B , and C . The actual effect matrix used would be just the columns to the left of the dots in Table 1. Note that the column labeled C is also the AB column. This leads to another method of creating this same design, that is, write out the 4×4 effect matrix for a 2^2 factorial in factors A and B and assign factor C to the AB column. Thus we may also refer to $C = AB$ as the *design generator*.

The assumption that $2fis$ are less likely to be important than main effects, and that $3fis$ are less likely to be important than $2fis$ (and so on) is known as the *hierarchical ordering principle*. While many experimenters apply this principle by assuming $3fis$ and higher interactions are negligible, it is indeed a strong assumption that $2fis$ are negligible.

Aliasing in More Highly Fractionated 2^{k-p} Experiments

Generally, screening designs involve more factors and tend to be a bit larger than the example in the previous discussion. For example, if seven factors are studied in $n = 16$ runs, the design is an unreplicated 2^{7-3} or a $1/2^3 = 1/8$ fraction of the 2^7 factorial. Recall that one design generator is required for a 2^{k-1} design. In general, a 2^{k-p} design requires p design generators. Chen and Lin [1] provide a succinct description of design generators and defining relations, which we discuss next.

For any 2^{k-p} design, let the factors be represented by the letters A , B , and so on. Then any product of letters is called a word and we can choose a (nonunique) set of p words, which generate the design. The p selected *generators* determine the design. Let I be the identity defined so that, for all words W , $IW = WI = W$ and $W^2 = I$. This enables us to write the product UW of two words U and W in a minimally reduced form. The set of distinct words formed by all possible products involving the p generators gives the defining relation, which contains 2^p terms, including the identity term I . The words in the defining relation correspond to those products of the design columns that are the same as the identity

I . A design is uniquely determined by the defining relation and vice versa.

Suppose the experimenter wishes to study seven factors labeled A , B , C , D , E , F , and G . One approach for creating this design is to write out the 16×16 effect matrix for a 2^4 factorial in factors A , B , C , and D (the so-called basic factors). The factors E , F , and G are then assigned to columns representing interactions between the first four factors (the design generators). This assignment should not be made arbitrarily as different designs will have different properties. This is elaborated upon subsequently. Table 2 shows such a design where $E = ABC$, $F = BCD$, and $G = ACD$ are the design generators. For the design in Table 3, we have $E^2 = ABCE$, $F^2 = BCDF$, and $G^2 = ACDG$. Combining these we have $I = ABCE = BCDF = ACDG$. Note however, that pairwise and three-way products of these terms are also equal to I , for example $(ABCE)(BCDF) = AB^2C^2DEF = ADEF$. Multiplying each of the other two pairs and the triple gives the defining relation

$$\begin{aligned} I &= ABCE = BCDF = ACDG = ADEF \\ &= BDEG = ABFG = CEF G \end{aligned}$$

Note that the defining relation has $2^3 = 8$ terms, including I , reflecting that it is a one-eighth fraction. Indeed if a 2^{7-4} design were used, the defining relation would have 16 terms.

As we did with the half fraction, we can find the full alias structure by multiplying the defining relation by individual effects. For example, multiplying by A gives $A = BCE = ABCDF = CDG = DEF = ABDEG = BFG = ACEFG$. This implies that when we estimate the effect of A , we are really estimating the combined effects of all terms in this string. If the three-factor and higher-order interactions are null (or negligible), then we can ignore such terms and assume the estimated effect is due to A . An abbreviated form of the alias structure is given in Table 3. Here we use another convention to represent the alias string where plus signs are used to show additivity among the effects in a string.

Suppose instead of design 1 in Table 2, we assigned E , F , and G as follows: $E = ACD$, $F = BCD$, and $G = ABCD$. Then the defining relation and alias pattern (abbreviated) are given in Table 4 for design 2.

Table 2 Example 2^{7-3} experimental design 1

												E =		G =	F =	A =
		A	B	C	D	A B	A C	A D	B C	B D	C D	A B C	A B D	A C D	B C D	B C D
	0	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
1	1	-1	-1	-1	-1	1	1	1	1	1	1	-1	-1	-1	-1	1
2	1	1	-1	-1	-1	-1	-1	-1	1	1	1	1	1	1	-1	-1
3	1	-1	1	-1	-1	-1	1	1	-1	-1	1	1	1	-1	1	-1
4	1	1	1	-1	-1	1	-1	-1	-1	-1	1	-1	-1	1	1	1
5	1	-1	-1	1	-1	1	-1	1	-1	1	-1	1	-1	1	1	-1
6	1	1	-1	1	-1	-1	1	-1	-1	1	-1	-1	1	-1	1	1
7	1	-1	1	1	-1	-1	-1	1	1	-1	-1	-1	1	1	-1	1
8	1	1	1	1	-1	1	1	-1	1	-1	-1	1	-1	-1	-1	-1
9	1	-1	-1	-1	1	1	1	-1	1	-1	-1	-1	1	1	1	-1
10	1	1	-1	-1	1	-1	-1	1	1	-1	-1	1	-1	-1	1	1
11	1	-1	1	-1	1	-1	1	-1	-1	1	-1	1	-1	1	-1	1
12	1	1	1	-1	1	1	-1	1	-1	1	-1	-1	1	-1	-1	-1
13	1	-1	-1	1	1	1	-1	-1	-1	-1	1	1	1	-1	-1	1
14	1	1	-1	1	1	-1	1	1	-1	-1	1	-1	-1	1	-1	-1
15	1	-1	1	1	1	-1	-1	-1	1	1	1	-1	-1	-1	1	-1
16	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1	1

Table 3 Defining relation and alias structure for design 1

I	= $ABCE$	= $BCDF$	= $ACDG$	= $ADEF$	= $BDEG$	= $ABFG$	= $CEFG$
A	+ BCE	+ BFG	+ CDG	+ DEF	+ $ABCDF$	+ $ABDEG$	+ $ACEFG$
B	+ ACE	+ AFG	+ CDF	+ DEG	+ $ABCDG$	+ $ABDEF$	+ $BCEFG$
C	+ ABE	+ ADG	+ BDF	+ EFG	+ $ABCFG$	+ $ACDEF$	+ $BCDEG$
D	+ ACG	+ AEF	+ BCF	+ BEG	+ $ABCDE$	+ $ABDFG$	+ $CDEFG$
E	+ ABC	+ ADF	+ BDG	+ CFG	+ $ABEFG$	+ $ACDEG$	+ $BCDEF$
F	+ ABG	+ ADE	+ BCD	+ CEG	+ $ABCEF$	+ $ACDFG$	+ $BDEFG$
G	+ ABF	+ ACD	+ BDE	+ CEF	+ $ABCEG$	+ $ADEFG$	+ $BCDFG$
AB	+ CE	+ FG	+ $ACDF$	+ $ADEG$	+ $BCDG$	+ $BDEF$	+ $ABCEFG$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
ABD	+ ACF	+ AEG	+ BCG	+ BEF	+ CDE	+ DFG	+ $ABCDEF$

Table 4 Defining relation and alias structure for design 2

I	= $ACDE$	= $BCDF$	= $ABCDG$	= $ABEF$	= BEG	= AFG	= $CDEFG$
A	+ FG	+ BEF	+ CDE	+ $ABEG$	+ $BCDG$	+ $ABCDF$	+ $ACDEFG$
B	+ EG	+ AEF	+ CDF	+ $ABFG$	+ $ACDG$	+ $ABCDE$	+ $BCDEFG$
C	+ ADE	+ BDF	+ $ABDG$	+ $ACFG$	+ $BCEG$	+ $DEFG$	+ $ABCEF$
D	+ ACE	+ BCF	+ $ABCG$	+ $ADFG$	+ $BDEG$	+ $CEFG$	+ $ABDEF$
E	+ BG	+ ABF	+ ACD	+ $AEFG$	+ $CDFG$	+ $BCDEF$	+ $ABCDEG$
F	+ AG	+ ABE	+ BCD	+ $BEFG$	+ $CDEG$	+ $ACDEF$	+ $ABCDFG$
G	+ AF	+ BE	+ $ABCD$	+ $CDEF$	+ $ABEFG$	+ $ACDEG$	+ $BCDFG$
AB	+ EF	+ AEG	+ BFG	+ CDG	+ $ACDF$	+ $BCDE$	+ $ABCDEF$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
DG	+ ABC	+ ADF	+ BDE	+ CEF	+ $ACEG$	+ $BCFG$	+ $ABDEFG$

Resolution and Aberration

Not all designs of the same resolution have the same *WLP*. When comparing two designs of the same resolution, say A and B , design A has less aberration than B if its first nonzero c_i is less than that of B . A also has less aberration than B if it is of higher resolution than B . For our two designs, then, D_1 has less aberration than D_2 . A design with *minimum aberration* within a class of designs is such that no other design exists with less aberration. (See Fries and Hunter [3] or **Minimum Aberration** for more information about aberration.)

All 2^{k-p} fractional factorials have a structure in which effects are either totally aliased (confounded) or not aliased at all and can be represented by a defining relation. In other designs, we may have effects that are partially aliased. Two effects are totally aliased if the **correlation** of their columns in the $\mathbf{X}$ matrix is -1 or $+1$, while if two columns have nonzero, but not perfect correlation, then their effects are partially aliased. Designs that have partially aliased effects are called *nonregular designs*. A popular series of nonregular designs was developed by Plackett and Burman [4] (see **Plackett–Burman Designs**). These designs have run sizes that are multiples of 4 and are a special class of Hadamard Matrices. Table 5 shows the 12-run Plackett and Burman design.

Table 5 Plackett–Burman 12-run design[illegible]

matrix $\mathbf{A} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{X}_1 = [(a_{ij})]$. Here, element a_{ij} measures the degree of aliasing between **main effect** i and interaction j . For example, if the first four columns of Table 5 are used to study four factors, the matrix $\mathbf{A}$ is as given in Table 6, ignoring higher-order interactions. From this we see that each main effect is partially confounded with three *2fis*, those that do not involve the main effect itself: $A = -1/3BC = -1/3BD = 1/3CD$, $B = -1/3AC = -1/3AD = -1/3CD$, $C = -1/3AB = 1/3AD = -1/3BD$, and $D = -1/3AB = 1/3AC = -1/3BC$. If 11 factors are studied in this 12-run design, then the main effect of each factor is partially confounded with 45 *2fis*, those not involving the factor. Plackett and Burman designs have this general property and are therefore resolution III designs. A complete alias table is given in Lin and Draper [5]. They also provide alias tables for Plackett and Burman designs of higher orders.

Aliasing in 3^{k-p} Designs

For three-level designs, we change our notation to 0, 1, 2 to denote the three levels of a factor. Because each factor has three levels, each main effect has 2 **degrees of freedom** and each *2fi* has 4 degrees of freedom. Since each column in the effect matrix has $df = 2$, we need two columns to designate a *2fi*. For example, the columns AB and AB^2 together represent the interaction effect

between factors A and B . The column AB is calculated as $A + B$ and AB^2 as $A + 2B$. Both calculations are done modulo 3 meaning the solution is the remainder left after dividing the quantity by 3. For example, $3(\text{mod}3) = 0$, $4(\text{mod}3) = 1$, and so on. It can be shown that A^2B results in the same sum-of-squares calculation as AB^2 so only one is needed. We use the convention that the first letter in an interaction is always first order and if it is not, we square the entire term to achieve this. Specifically, $A^2B(\text{mod}3) = (A^2B)^2(\text{mod}3) = A^4B^2(\text{mod}3) = AB^2(\text{mod}3)$.

So the effect matrix for the 3^3 design will have three columns for the three main effects, six columns for the three *2fis*, and four columns for the *3fi*: A , B , C , AB , AB^2 , AC , AC^2 , BC , BC^2 , ABC , ABC^2 , AB^2C , and AB^2C^2 . We can create a 3^{4-1} design by assigning the fourth factor D to one of the interaction columns, as we did when creating a 2^{k-1} design. If we let $D = ABC^2$ for example, then we have the defining relation $I = ABC^2D$. We determine the aliases of an effect by multiplying by both I and I^2 , or ABC^2D and $A^2B^2CD^2$. For our example, this leads to the alias structure shown in Table 7. More highly fractionated designs can be created by adding terms to the defining relation. See Montgomery [6] for more information on aliasing in 3^{k-p} designs.

Summary

In regular designs such as the 2^{k-p} and 3^{k-p} series, the defining relation can be used to determine what effects are confounded with each other, that is, effects that are aliases. Multiplying an effect by each term in the relation results in the effects that are confounded. In nonregular designs, the regression matrix $\mathbf{X}$ may be created by listing all effects. Then the alias matrix $\mathbf{A} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{X}_1$ shows the alias pattern, including effects that are partially confounded. This method may also be used for regular designs.

Table 6 Alias matrix of four factors in the Plackett–Burman 12-run design

	AB	AC	AD	BC	BD	CD
A	0	0	0	−0.33	−0.33	0.33
B	0	−0.33	−0.33	0	0	−0.33
C	−0.33	0	0.33	0	−0.33	0
D	−0.33	0.33	0	−0.33	0	0

Table 7 Alias structure of 3^{4-1} design with defining relation $I = ABC^2D$

$A = AB^2CD^2 = BC^2D$	$BC = AB^2D = ACD$
$B = AB^2C^2D = AC^2D$	$BD = AB^2C^2D^2 = AC^2$
$C = ABD = ABCD$	$CD = ABD^2 = ABC$
$D = ABC^2D^2 = ABC$	$AB^2 = ACD^2 = BCD^2$
$AB = ABCD^2 = CD^2$	$AD^2 = AB^2C = BC^2D^2$
$AC = AB^2D^2 = BCD$	$BD^2 = AB^2C^2 = AC^2D^2$
$AD = AB^2CD = BC^2$	

References

- [1] Chen, J. & Lin, D.K.J. (1991). On the identity relationships of 2^{k-p} designs, *Journal of Statistical Planning and Analysis* **28**, 95–98.
- [2] Box, G.E.P., Hunter, J.S. & Hunter, W.G. (2005). *Statistics for Experimenters: Design, Innovation, and Discovery*, John Wiley & Sons, New York, p. 242.
- [3] Fries, A. & Hunter, W.G. (1980). Minimum aberration 2^{k-p} designs, *Technometrics* **22**, 601–608.
- [4] Plackett, R.L. & Burman, J.P. (1946). The design of optimum multivariable experiments, *Biometrika* **33**, 305–325.
- [5] Lin, D.K.J. & Draper, N.R. (1993). Generating alias relationships for two-level Plackett and Burman designs, *Computational Statistics and Data Analysis* **15**, 147–157.
- [6] Montgomery, D.C. (2001). *Design and Analysis of Experiments*, 6th Edition, John Wiley & Sons, New York, pp. 356–365.

RICHARD N. MCGRATH AND DENNIS K.J. LIN

Analysis of Variance

Introduction

It is hard to imagine a more commonly used data analysis technique than the analysis of variance, or simply, **ANOVA**. First introduced by Fisher [1, 2], (with Mackenzie [3]) for the analysis of agriculture experiments, which could be arranged as Latin squares (see **Latin Squares and Related Experimental Designs**) or randomized block designs (see **Blocking**), ANOVA has become ubiquitous in statistical practice. It is the standard approach for the analysis of most industrial experiments (e.g., see Box *et al.* [4], Montgomery [5], or Wu and Hamada [6]).

The aim of ANOVA, in short, is to determine if there exists a relationship between a dependent variable (or response variable) and one or more experimental factors. A relationship is identified by observing significantly different mean responses at the various settings of the experiment factors (factor levels). The ANOVA approach partitions the variability, and associated **degrees of freedom**, of the response into variability due to error and factor effects. Variability that is not explained by chance variation is taken as evidence of a significant effect.

Methodology

In many investigations, experimenters are interested in exploring the impact of two or more factors on the mean response of a process. The analysis of variance is a natural way to address this goal. Many studies have more than one factor, but for ease of presentation the single-factor (one-way) ANOVA is first presented, followed by generalizations to multifactor studies.

Suppose an experiment has been conducted with a single factor (e.g., temperature or pressure) at I levels. Further, suppose that there are n_i observations taken at each level ($i = 1, 2, \dots, I$), and thus, the experiment has $N = \sum_{i=1}^I n_i$ total observations. Let y_{ij} denote the j th observation ($j = 1, 2, \dots, n_i$) taken at level i . The one-way ANOVA model can be written as

$$y_{ti} = \mu_i + \epsilon_{ij} \quad \text{for } i = 1, 2, \dots, I \quad \text{and} \\ j = 1, 2, \dots, n_i \quad (1)$$

where μ_i is the mean response for the i th treatment, and the errors, ϵ_{ij} , are independent and identically distributed (i.i.d.) $N(0, \sigma^2)$ for all i and j . An equivalent alternative form for the ANOVA model is

$$y_{ti} = \mu + \alpha_i + \epsilon_{ij} \quad \text{for } i = 1, 2, \dots, I \quad \text{and} \\ j = 1, 2, \dots, n_i \quad (2)$$

where μ is the grand system mean over the I levels of the factor and α_i is the effect of the factor at level i . For the most part, this presentation focuses on model (2) because it is easily understandable in most industrial settings. That is, by setting the factor at level i , the system mean is changed by α_i beyond the grand mean μ .

The above model (2), as written, is not identifiable. If, for instance, one rewrites equation (2) as a regression model (e.g., Kutner *et al.* [7]), then the model matrix is not invertible. To address this issue, additional constraints are placed on the model parameters. The most common such constraint is to restrict the factor effects (e.g., the α_i 's) so that they sum to zero (more on this later). Although there are other constraints that can be used to make the model parameters estimable, this formulation lends itself to simple interpretation.

An important question that ANOVA addresses is whether there is evidence of differences among the mean response at the different factor levels. The factor effects, in conjunction with the constraint $\sum_i \alpha_i = 0$, can be interpreted as the impact on the response of setting the factor to level i above the process mean of the I levels, μ . So, if the mean response is the same at each level of the factor, then $\alpha_i = 0$ for all i . Consequently, an important objective of ANOVA is to test the null hypothesis that all the factor effects (i.e., the α_i 's) are equal to zero *versus* the alternate hypothesis that they are not.

The operational idea behind one-way ANOVA is to use two estimators of σ^2 , one which is robust to the hypotheses and one which is not, to assess the evidence against the null hypothesis. Begin by constructing the total **sum of squares**,

$$SS_T = \sum_{i=1}^I \sum_{j=1}^{n_i} (y_{ij} - \bar{y})^2 \quad (3)$$

where $\bar{y}$ is the sample mean of all of the observations. The total sum of squares measures the total variability

2 Analysis of Variance

in the data. It can be written as,

$$SS_T = \sum_{i=1}^I \sum_{j=1}^{n_i} (y_{ij} - \bar{y}_i)^2 + \sum_{i=1}^I n_i (\bar{y}_i - \bar{y})^2 \quad (4)$$

where $\bar{y}_i$ is the sample mean of observations at level i . This decomposition is a key feature of ANOVA.

The first term in the above decomposition, called the *error* (or *residual*) *sum of squares* and denoted SS_E , measures the squared deviations of observations from their respective treatment mean. Since the error variance is assumed to be constant across levels of the factor, the SS_E should not be impacted by departures from the null hypothesis and when properly normalized can be used to estimate σ^2 . The second term (the between treatment sum of squares, denoted SS_B) in the decomposition of the total sum of squares measures the squared deviations of the treatment means from the grand mean. Under the null hypothesis, this sum, properly normalized, can also be used to estimate σ^2 . On the other hand, when the null hypothesis is false this sum will be inflated.

Under the above model, the variance of the response at each level of the factor is σ^2 . Therefore, σ^2 can be estimated by pooling the sample variances from each level. Specifically, the pooled estimate of the population variance is

$$MS_E = \frac{\sum_{i=1}^I \sum_{j=1}^{n_i} (y_{ij} - \bar{y}_i)^2}{N - I} \quad (5)$$

This estimator of σ^2 is called the *error mean square*, and is just the error sum of squares divided by its degrees of freedom. Similarly, σ^2 can be estimated using the between treatments sum of squares. Specifically, the between treatments mean square is

$$MS_B = \frac{\sum_{i=1}^I n_i (\bar{y}_i - \bar{y})^2}{I - 1} \quad (6)$$

This is simply the between treatment sum of squares divided by its degrees of freedom.

The above information is summarized in Table 1. The idea behind ANOVA is to use the two estimators of σ^2 to test the null hypothesis of the equality of treatments. This is done *via* their ratio $F = MS_B/MS_E$. Under the null hypothesis, both the mean squares are unbiased estimators of σ^2 , and the ratio has an F distribution with $(I - 1, N - I)$ degrees of freedom. When the null hypothesis is false, only the MS_E is an unbiased estimator, and MS_B will get large. Thus large values of the F statistic provide evidence against the null hypothesis. The p value for assessing the evidence against the null of the equality of treatment effects is found from the F distribution with $(I - 1, N - I)$ degrees of freedom.

Parameter Estimation and Contrasts

The parameters of the ANOVA model (2) are unknown and have to be estimated from the data. The lack of identifiability of the model in equation (2) is addressed by placing restrictions on the model parameters (i.e., $\sum_i \alpha_i = 0$). Specifically, one can rewrite equation (2) as a linear model with $\alpha_I = -\alpha_1 - \alpha_2 \cdots - \alpha_{I-1}$, without loss of generality. Next, the process mean, μ , and the α_i 's are estimated by **least squares**. When each group has the same number of observations (the data are balanced), $\hat{\mu} = \bar{y}$ and $\hat{\alpha}_i = \bar{y}_i - \bar{y}$. Here, the least squares estimators reinforce the intuitive interpretation of the factor effects. That is, the estimate of the factor effect is simply the difference between the treatment sample mean and the overall mean. So, the treatment effects measure the impact on the process mean of setting the factor to its i th level, after adjusting for the grand mean.

When the null hypothesis is true, one expects that the treatment sample means will be close to the grand mean of the data. Of course, there are

Table 1 One-way ANOVA table

Source of variation	Sum of squares	Degrees of freedom	Mean square	F statistic
Between treatments	SS_B	$I - 1$	$MS_B = \frac{SS_B}{I - 1}$	$F = \frac{MS_B}{MS_E}$
Error	SS_E	$N - I$	$MS_E = \frac{SS_E}{N - I}$	
Total	SS_T	$N - 1$		

other restrictions one could have chosen (e.g., see Searle [8, Chapter 5]) that would, in turn, give different meaning to the model parameters. It is important to note that software packages will report the same ANOVA table, but can use different default model restrictions. Therefore, in practice, one should be careful to ascertain the meaning of the model parameter estimates.

When the null hypothesis is rejected and one concludes that there is evidence that the treatment effects (or, equivalently, the treatment means) are different, the next obvious step is to investigate how these differences occurred. One way to approach this is to perform hypothesis tests for the equality for each of the pairwise differences of the treatment means, or factor effects (see Kutner *et al.* [7] for a discussion of multiple comparisons tests). More generally, a contrast is a linear combination of the treatment means in model (1), $L = \sum_i c_i \mu_i$ where $\sum_i c_i = 0$. The contrast is a comparison of two or more of the treatment means, and the pairwise differences used in the multiple comparisons test as a special case. It turns out that the contrast can be estimated by $\hat{L} = \sum_i c_i \bar{y}$. Hypotheses about the contrast can be evaluated and **confidence intervals** can be constructed using a t -distribution (see Kutner *et al.* [7, Chapter 17]).

ANOVA with Two Factors

The one-way ANOVA model is used to investigate the impact of one factor with two or more levels. Of course, many studies have more than one factor. We first consider the case when the sample sizes at each setting of the factors are identical (i.e., the data are balanced), in which case the analysis is a simple generalization of the one-factor model. When there are different numbers of observations at each setting, the data are said to be unbalanced.

To understand the multifactor ANOVA model, it is easiest to present the two-factor model and then mention how to generalize to any number of factors. In the same spirit as one-way ANOVA, multifactor ANOVA attempts to partition the total system variability, and degrees of freedom, of the data into variability due to error (within treatments) and variability due to the effects of factors (between treatments).

The two-factor ANOVA model, with factors labeled A and B , can be written as

$$y_{ijk} = \mu + \alpha_i + \beta_j + \eta_{ij} + \epsilon_{ijk}; \quad \text{for } i = 1, 2, \dots, I, \\ j = 1, 2, \dots, J, \text{ and } k = 1, 2, \dots, n_{ij} \quad (7)$$

where α_i is the effect of factor A at level i , and β_j is the effect of factor B at level j , η_{ij} is the interaction effect at levels i and j of factors A and B (see **Interactions** for details on interactions). As before, the errors, ϵ_{ijk} , are i.i.d. $N(0, \sigma^2)$ random variables for all i, j , and k . Let $N = \sum_{i=1}^I \sum_{j=1}^J n_{ij}$ be the total number of observations.

The two-factor ANOVA model, similar to the one-factor model, is not identifiable. Again, constraints on the parameters are required to make the model identifiable. Specifically, it is common to impose the restrictions $\sum_i \alpha_i = 0$, $\sum_j \beta_j = 0$, $\sum_i \eta_{ij} = 0$ and $\sum_j \eta_{ij} = 0$. For this model and these restrictions, the factor effects can be thought of as the impact on the system mean response when changing the levels of one of the factors. The interaction effect, η_{ij} is the joint impact on the system mean when factor A is at level i and factor B is at level j , beyond the individual effects of each factor.

Let $\bar{y}$ be the sample mean of all the observations, $\bar{y}_{i..}$ be the mean response at level i of factor A , $\bar{y}_{.j.}$ be the mean response at level j of factor B , and $\bar{y}_{ij.}$ be the mean response at levels i and j of factors A and B . For balanced data, the total sum of squares is the same as defined above, but the variability is partitioned to represent the model components; that is,

$$SS_T = SS_A + SS_B + SS_{AB} + SS_E \quad (8)$$

where the sums of squares for A , B and the interaction effect are, respectively,

$$SS_A = \sum_{i=1}^I \sum_{j=1}^J \sum_{k=1}^{n_{ij}} (\bar{y}_{i..} - \bar{y})^2 \quad (9)$$

$$SS_B = \sum_{i=1}^I \sum_{j=1}^J \sum_{k=1}^{n_{ij}} (\bar{y}_{.j.} - \bar{y})^2 \quad (10)$$

$$SS_{AB} = \sum_{i=1}^I \sum_{j=1}^J \sum_{k=1}^{n_{ij}} (\bar{y}_{ij.} - \bar{y}_{i..} - \bar{y}_{.j.} + \bar{y})^2 \quad (11)$$

Note that we are now using SS_B to denote the sum of squares for factors B .

4 Analysis of Variance

Table 2 Two-way ANOVA table

Source of variation	Sum of squares	Degrees of freedom	Mean square	F statistic
Factor A	SS_A	$I - 1$	$MS_A = \frac{SS_A}{I - 1}$	$F = \frac{MS_A}{MS_E}$
Factor B	SS_B	$J - 1$	$MS_B = \frac{SS_B}{J - 1}$	$F = \frac{MS_B}{MS_E}$
AB interaction	SS_{AB}	$(I - 1)(J - 1)$	$MS_{AB} = \frac{SS_{AB}}{(I - 1)(J - 1)}$	$F = \frac{MS_{AB}}{MS_E}$
Error	SS_E	$IJ(N - 1)$	$MS_E = \frac{SS_E}{IJ(N - 1)}$	
Total	SS_T	$N - 1$		

As before, the significance of each effect is assessed using the appropriate F statistic, with degrees of freedom associated with the numerator and denominator components, respectively. The corresponding ANOVA table for this model is shown in Table 2.

Note, the above sums of squares decomposition is appropriate when the treatment sample sizes are all the same. When they are not all equal, the decomposition breaks down. That is, when the n_{ij} are not all the same, the total sum of squares is not equal to the sum of SS_A , SS_B , SS_{AB} , and SS_E . This, in turn, impacts the distribution of the parameter estimates, the hypotheses being tested and the corresponding test statistics. Considerable attention has been drawn to AVOVA with unbalanced data, and several analysis strategies proposed (e.g., see Searle [9] and Milliken and Johnson [10]). Hypothesis tests in the unbalanced setting are often conducted by comparing the total sum of squares for a “full” model (the model with all effects present) to the total sum of squares for the “reduced” model (a model without the effect of interest). This is frequently called the *Type III analysis* (e.g., Milliken and Johnson [10, p. 150]). In this case, one is investigating the marginal explanatory power of an effect after including all other effects in the model. When all of the n_{ij} are equal, the Type III analysis hypothesis tests coincide with those for balanced datasets.

Generalizations

To generalize the ANOVA model to more than two factors, one simply augments the two-factor model with the appropriate number of main effects and

interactions. Indeed, one can also add all interactions of any order, if the degrees of freedom permit. When the degrees of freedom are not sufficient to estimate all interactions and experimental error, the higher-order interactions are usually not included in the model, under the assumption that they are negligible.

Checking Assumptions

Application of ANOVA is not simply the construction of a table and the computation of F statistics. Indeed a key step is the verification of the underlying model assumptions and this is a routine part of any ANOVA endeavor.

The ANOVA model assumes that the data are random samples from normal populations with the same variance and potentially different means. So the ϵ 's should be a random sample from an $N(0, \sigma^2)$ distribution. Though the exact errors are unavailable, these errors can be estimated by the **residuals** (the observed minus the expected response) (see **Residuals**). The validity of the model assumptions are usually assessed using several simple plots, some of which are now discussed.

The normality assumption is usually validated using a normal probability plot. This is done by plotting the ordered residuals *versus* the **percentiles** of the standard **normal distribution**. If the plot reveals a straight-line pattern, this is taken as evidence that the residuals are normally distributed, otherwise one concludes that the normality assumption is violated.

To verify that the residuals appear to be i.i.d., several separate plots are constructed. The plots attempt to see if there are any systematic departures from the i.i.d. assumption. Common plots to visualize

Table 3 Example data for the plutonium experiment

A	B	Replicate 1	Replicate 2	Replicate 3
Low	Low	4.08	4.54	4.07
Medium	Low	4.48	5.23	4.19
High	Low	5.66	5.04	5.59
Low	Medium	4.05	4.65	4.16
Medium	Medium	5.15	4.07	4.51
High	Medium	4.86	4.14	4.98
Low	High	4.64	4.95	4.36
Medium	High	4.11	4.66	4.14
High	High	5.36	5.11	4.93

such violations are time plots (plot the residuals against the order in which they were collected), plotting the residuals *versus* their treatment and also plotting the residuals *versus* the expected response. In each case, any systematic pattern that is visualized (short of random scatter) is taken as evidence that the i.i.d. assumption has been violated.

If the residual plots reveal departures from the model assumptions, then it is common to take remedial measures. These measures usually take the form of transforming the responses, the predictors, or both. For instance, when the distribution of the residuals appears to exhibit **skewness** or nonconstant variance, a Box–Cox transformation (*see* **Box and Cox Transformation**) of the response variable is often used to make the distribution of the residuals appear approximately normal. In more complicated settings, one may elect to model both the mean and the variance (*see* **Generalized Linear Models**).

Example

Researchers at the Los Alamos National Laboratory were interested in studying a treatment process that creates a type of plutonium alloy with a high decay rate that can eventually be used for accelerated testing (for details on the physics, see Freibert *et al.* [11]). The experiment is modified for this illustration. Consider a designed experiment to determine which of two factors, at three levels (low, medium, and high) impact the decay energy (in MeV) and should be subsequently studied to optimize the process. The factors *A* and *B* relate to different chemical compositions and processing temperatures, respectively. The data for three replicates of the nine possible levels settings are shown in Table 3.

The ANOVA table for these data is shown in Table 4. Notice that the *p*-value for the *F*-statistic for factor *A* is quite small (*p* value = 0.0027) whereas the *p* values for the remaining two effects are relatively large. Consequently, only factor *A* has a significant effect on the mean response, which increased on average from 4.39 MeV when *A* was at its low level to 5.07 MeV at the high level of factor *A*. There is no clear evidence that factor *B* affects the mean response or that there is any interaction between the two factors. One should not conclude from the analysis that the processing temperature (factor *B*) is not important. Instead, the data indicated that there was not a significant difference in the process mean over the range of temperatures explored.

In the interests of space, the residual analysis and multiple comparisons tests are not presented. However, a complete investigation of the data will always include these steps. Such analysis must be routine practice.

Comments

There are some final things that should be mentioned. Firstly, this brief presentation has considered only data that arise from completely randomized designs. Randomized block (*see* **Blocking**), **Latin square** (*see* **Latin Squares and Related Experimental Designs**), and split-plot (*see* **Split-Plot Designs**) designs are all examples of designs with **randomization** restrictions (e.g., see Montgomery [5]). When conducting an ANOVA in these settings, it is important to integrate the additional random effects into the analysis.

One must also be cognizant that the data are assumed to be continuous. Frequently, the responses are counts or discrete values. Direct application of

6 Analysis of Variance

Table 4 ANOVA table for the plutonium experiment

Source of variation	Sum of squares	Degrees of freedom	Mean square	<i>F</i> -statistic	<i>P</i> value
Factor <i>A</i>	2.42	2	1.21	8.34	0.0027
Factor <i>B</i>	0.32	2	0.16	1.09	0.3565
<i>AB</i> interaction	1.09	4	0.27	1.87	0.1601
Error	2.62	18	0.15		
Total	6.45	26			

ANOVA in these settings is inefficient. One possible solution is to transform the data (e.g., square root or log). Such transformations can also be a good strategy for dealing with continuous data that do not follow a normal distribution. An alternate solution for discrete data is to use a **generalized linear model** [12] instead of the ANOVA model.

References

- [1] Fisher, R.A. (1921). Studies in crop variation. I. An examination of the yield of dressed grain from Broadbalk, *Journal of Agricultural Science* **11**, 107–135.
- [2] Fisher, R.A. (1925). Applications of student's distribution, *Metron* **5**, 90–104.
- [3] Fisher, R.A. & Mackenzie, W.A. (1923). Studies in crop variation. II. The manurial response of different potato varieties, *Journal of Agricultural Science* **13**, 311–320.
- [4] Box, G.E.P., Hunter, J.S. & Hunter, W.G. (2005). *Statistics for Experimenters: Design, Innovation, and Discovery*, 2nd Edition, John Wiley & Sons, New York.
- [5] Montgomery, D.C. (2005). *Design and Analysis of Experiments*, 6th Edition, John Wiley & Sons, New York.
- [6] Wu, C.F.J. & Hamada, M. (2000). *Experiments: Planning, Analysis and Parameter Design Optimization*, John Wiley & Sons, New York.
- [7] Kutner, M.H., Nachtsheim, C.J., Neter, J. & Li, W. (2004). *Applied Linear Statistical Models*, 5th Edition, McGraw-Hill/Irwin.
- [8] Searle, S.R. (1971). *Linear Models*, John Wiley & Sons, New York.
- [9] Searle, S.R. (1987). *Linear Models for Unbalanced Data*, John Wiley & Sons, New York.
- [10] Milliken, G.A. & Johnson, D.E. (1984). *Analysis of Messy Data*, Wadsworth, California.
- [11] Freibert, F., Olivas, J.D. & Coonley, M.S. (2002). Researchers cast first “spiked” plutonium alloy, *Actinide Research Quarterly* **2**, 1–6.
- [12] McCullagh, P. & Nelder, J.A. (1989). *Generalized Linear Models*, Chapman & Hall, London.

Related Articles

Blocking; Interactions; Latin Squares and Related Experimental Designs; Split-Plot Designs; Signal-to-Noise Ratios for Robust Design; Box and Cox Transformation.

DEREK BINGHAM

Blocking

Introduction and Definition

When planning an experiment, experimental units can be grouped into *blocks* in such a way that units in the same block are expected to give similar responses. This allows more precise comparisons of treatments to be made because, when performing a data analysis, the variation between blocks can be separated from the variation within blocks.

The *experimental units* are the objects to which the *treatments* (typically combinations of levels of several input factors) are applied, so that each unit receives only one treatment, while different units can receive different treatments. They are typically *runs* of the process which is the subject of experimentation, or people (*subjects*) who are under study. In an experiment on the mixing of pastry dough that we use throughout as an illustration, the treatments were the eight combinations of low and high levels of the flow rate of materials into the mixer, the initial moisture content in the mix, and the screw speed at which the mixer was turned. The experimental units were eight runs of a food extrusion cooker used to mix the dough. In other contexts, the experimental units could be different products, animals, pieces of land, and so on. Irrespective of the treatments applied, the units will give different responses. If there are no known patterns of variation among the units, the treatments should be applied to the units completely at random. This is the only objective way to allocate treatments and it ensures that differences between treatments can be estimated free from **bias** (*see Randomization in Experimental Designs*).

Often there are predictable patterns of variation among the experimental units, for example, runs made on the same day can be expected to give more similar responses than runs made on different days or subjects of similar ages can be expected to give more similar responses than subjects of greatly different ages. By identifying these predictable patterns of variation, we can aim to ensure that the treatments can be assessed by comparing the effect of different treatments run on the same day, rather than on different days. This allows more precise **estimation** of the differences between treatments. In the pastry dough experiment, only four runs could be made per day, so that two blocks, each of four runs, were used.

In order to be useful, blocks have to be chosen so that units within each block are expected to give similar responses; but formally, blocking is defined as a form of restricted **randomization**. In a *completely randomized design*, given a treatment design (i.e., the treatments and the number of replicates of each) and a set of experimental units, the treatments are assigned to the experimental units completely at random (*see Randomization in Experimental Designs*). In a *randomized block design* the randomization is restricted, so that the treatments in a specific block are allocated to the units in that block completely at random, with the randomization being done independently for each block. Thus we need to decide not just the replication of each treatment, but also the allocation of treatments to blocks. In the example, after defining which four treatments would go in each block, the order of the blocks was selected at random and then the run order within each block was selected at random.

Note that the formal definition of blocking does not imply a good way to choose blocks for a specific experiment. In order to be successful, the units within the same block must give more similar responses than the units in different blocks. Of course, it is not possible before running the experiment to guarantee that this will be true. The statistical theory tells the experimenters that they should try to identify suitable blocks composed of homogeneous units. All the experimenters' prior knowledge, insight, and experience are needed to ensure that a good set of blocks is actually chosen.

In some applications the block size will be very clear, for example if testing is being carried out in a number of different centers on six lines (units) in each center, it is clear that we have blocks of six units each. In many cases, however, there will be some choice as to how many units are in each block, for example if blocks correspond to the runs made on the same day there might be some flexibility over how many runs can be made each day. The usual advice in such situations would be to make the blocks as large as is reasonably practicable, but this must be balanced against the concern that large blocks will compromise the similarity of the units in the blocks, against the danger of being unable to complete the full number of units in a day for example, and against the possible deterioration in experimental technique if experiments are performed too hurriedly. In the example, the experimenter was confident that four

2 Blocking

runs could be made each day, so this was the block size used.

See Mead [1, Chapters 2, 7, and 9], Hinkelmann and Kempthorne [2, Chapters 5 and 9], or other general books on the design of experiments for further discussion of randomization and blocking. Sometimes the experimental units are divided into “blocks”, but the treatments are not randomized to the units within the blocks. In such cases, it is probably advisable to analyze the data as if the experiment had been properly randomized. However, it must be realized that the interpretation of such an analysis is very strongly dependent on the assumption that units within a block are homogeneous and so is much less robust than it would have been had the design been carefully randomized.

Complete Blocks

If we plan to have equal replication of each treatment and if it is sensible to choose the block size to be equal to the number of treatments, then it is possible and advisable to use a *complete block design*, in which each treatment appears exactly once in each block. For example, if our pastry dough experiment had only four treatments, a complete block design is shown in Table 1. Before being used in an experiment, this *combinatorial design* should be randomized, by randomly allocating the treatments to the units in each block independently. The design is now known as a *randomized block design* or more correctly as a *randomized complete block design (RCBD)*. One particular randomization of the combinatorial design of Table 1 is shown in Table 2. Note that for the complete block design, there is no need to randomize the blocks, since each block contains exactly the same set of treatments.

Table 1 Complete block design for four treatments, numbered 1–4, in two blocks

Block	
I	II
1	1
2	2
3	3
4	4

Table 2 Randomized complete block design for four treatments, numbered 1–4, in two blocks

Unit	Block	
	I	II
I	1	1
II	4	3
III	2	2
IV	3	4

The RCBD is relatively easy to set up and to use, and it is very simple to analyze the resulting data. For continuous responses, letting $y_{ij(r)}$ denote the j th unit, $j = 1, \dots, t$, in block i , $i = 1, \dots, b$, with treatment r applied, $r = 1, \dots, t$, the data can be analyzed using the linear model

$$Y_{ij(r)} = \mu + b_i + \tau_r + \epsilon_{ij} \quad (1)$$

where $\sum_{i=1}^b b_i = 0$, $\sum_{r=1}^t \tau_r = 0$ and the ϵ_{ij} are assumed to be independent and to have constant variance σ^2 . The analysis resulting from this model can also be justified by the randomization carried out, assuming only that treatment and unit effects are additive. This model leads to an analysis of variance having the form shown in Table 3, where $n = bt$ is the total number of units (*see Analysis of Variance*).

The variance for estimating any treatment contrast (i.e., any comparison among treatments) from this design is the corresponding variance from a completely randomized design multiplied by σ^2/σ_u^2 , where the subscript u indicates that the variance is from the unblocked design. Thus, as long as the variance between units within blocks is less than the variance between units across blocks, blocking leads to improved estimation. If the blocks do, indeed, form

Table 3 Skeleton analysis of variance for randomized complete block design

Stratum	Source	df
Blocks	Error	$b - 1$
Units	Treatments	$t - 1$
	Error	$n - b - t + 1 = (b - 1)(t - 1)$
	Total	$n - b = b(t - 1)$
Total		$n - 1 = bt - 1$

homogeneous units, the reduction in variance can be very large.

Incomplete Blocks

General Ideas

It is often necessary to use *incomplete blocks*, that is, those that do not contain all treatments. In particular, when using **factorial** treatment structures as in our example, there is usually a large number of treatments and for many practical situations it is advisable to use blocks that are smaller, sometimes considerably smaller, than the number of treatments. Incomplete blocks raise the new problem of how to arrange the treatments into blocks, that is, how to choose which subset of the treatments should be used in each block. Consider a trivial example of three treatments, each replicated twice, in three blocks, each containing two units. The possible arrangements in blocks (up to symmetry) are given in Table 4. In this simple case it is obvious that design 1 is best, since the other designs do not allow a comparison (within blocks) of the treatment labeled 1 with either of the other treatments. Note that the table gives the *combinatorial design*, which simply shows which treatments appear together in blocks. Before being used for a particular experiment, it must be randomized. Randomly reordering the blocks, we get the order (II, I, III). Then we randomly allocate the treatments in each block, obtaining the orders (II, I), (I, II), and (II, I), respectively, giving the plan shown in Table 5. Of course, this randomization must be carried out afresh for each new experiment. However, the properties of the design depend only on the combinatorial design and the method of randomization and not on the outcome of the randomization.

In general, it will be much less obvious what the best design (i.e., arrangement of treatments to blocks) will be, but we try to choose designs which allow for good comparisons between treatments. The

Table 4 Designs for three treatments in blocks of size 2

Design 1 Block			Design 2 Block			Design 3 Block		
I	II	III	I	II	III	I	II	III
1	1	2	1	2	2	1	2	3
2	3	3	1	3	3	1	2	3

Table 5 Design 1 from Table 4 after randomization

Unit	Block		
	I	II	III
I	3	1	3
II	1	2	2

precise meaning of this depends on the objectives of the experiment, which in turn define the treatment structure, the relevant treatment contrasts, and the model for treatments. For example, if the objective is to find the best of several products, we will be interested in direct comparisons of pairs of treatments and should choose a design which minimizes the variances of these comparisons. If, on the other hand, the objective is to optimize the running of some industrial process, we might fit a multifactor polynomial model to the responses and should choose a design that allows the parameters of such a model to be well fitted.

If the design is randomized by randomly permuting blocks and by randomizing treatments to units within each block (independently), then the implied linear model is similar to that for a complete block design,

$$Y_{ij(r)} = \mu + b_i + \tau_r + \epsilon_{ij} \quad (2)$$

Formally, the block effects should be considered as random, giving the linear mixed model

$$Y_{ij(r)} = \mu + \delta_i + \tau_r + \epsilon_{ij} \quad (3)$$

but it usually makes little difference and often the data are analyzed as if the blocks were fixed. The analysis of variance (*see Analysis of Variance*) also takes the same form as for the RCBD, but the variances of particular treatment contrasts depend on how the treatments are allocated to blocks. For some estimated treatment contrasts the variances will be greater than the corresponding variances from a completely randomized design multiplied by σ^2/σ_u^2 . In the analysis with fixed blocks, define the *efficiency factor* E_k for a contrast C_k as

$$E_k = \frac{V_u(\hat{C}_k)/\sigma_u^2}{V(\hat{C}_k)/\sigma^2} \quad (4)$$

so that contrast C_k is better estimated from the blocked design if $E_k > \sigma^2/\sigma_u^2$. Note that all efficiency factors are 1 in a complete block design.

4 Blocking

Efficiency factors of less than 1 represent a loss of information on treatment comparisons. This information is “lost” in the sense that it cannot be obtained from comparisons between units in the same block. If block effects are considered to be random, then this information can be recovered by comparing block totals (or means). The variance relevant for assessing this information is the variance between blocks, which will typically be high, so this *interblock information* is often of little use. However, there are some circumstances, discussed later, under which the interblock information is both useful and necessary for a proper analysis to be carried out. Residual maximum likelihood (REML) analysis is sometimes used to combine the interblock analysis with the usual, intrablock analysis, but this is valid only for large numbers of blocks and can have poor properties in smaller experiments. John and Williams [3, Chapter 7] give a thorough description of interblock information and the combined analysis.

Unstructured Treatments

When the treatments are categories of a single qualitative factor, it is usually the case that the experimenters are equally interested in all pairwise comparisons of treatments, for example, to find the best treatment or simply to check whether there are differences between the treatments. In this case, we should aim to make the design *balanced*, that is, $V(\hat{\tau}_r - \hat{\tau}_s)$ should be equal for all $r, s; r \neq s$. Balanced designs exist only for some specific numbers of treatments, replications, and block sizes.

For equally replicated treatments and blocks of equal sizes, it is easily shown that a design is balanced if all pairs of treatments appear together in blocks equally often. This makes it very simple to find small *balanced incomplete block designs* (BIBDs) when they exist (see **Block Designs, Incomplete**). For example, a BIBD for 6 treatments in 10 blocks each of 3 units is given in Table 6. For larger problems several methods of construction are available and several lists of BIBDs are available, for example, Cochran and Cox [4, Chapter 9] or <http://www.designtheory.org/>.

When no BIBD exists, a nearly balanced design is chosen. However, the definition of nearly balanced is not clear and it is usual to choose a design to optimize a criterion such as A-, D- or E-efficiency (see **Optimal Design**). A design is said to be

Table 6 Balanced incomplete block design for 6 treatments in 10 blocks of 3

I	Block								
	II	III	IV	V	VI	VII	VIII	IX	X
1	1	1	1	1	2	2	2	3	3
2	2	3	4	5	3	4	5	4	4
3	4	5	6	6	6	5	6	5	6

- *A-optimally* blocked if it minimizes the average variance of estimated pairwise comparisons among treatments.
- *D-optimally* blocked if it minimizes the volume of a joint confidence region for the parameters.
- *E-optimally* blocked if it maximizes the minimum efficiency factor for any estimated contrast.
- *M-optimally* blocked if it maximizes the mean of the efficiency factors for pairwise comparisons among treatments.
- *(M,S)-optimally* blocked if it minimizes, among M-optimally blocked designs, the variance of the efficiency factors for pairwise comparisons among treatments.

If it exists, a BIBD optimizes all of these criteria. See John and Williams [3, Chapter 2] for a discussion of these criteria for optimal blocking.

For relatively small designs, it is possible to do a complete search of all possible block designs (see <http://www.designtheory.org/>). For larger designs, a search procedure must be used which will not be guaranteed to find the **optimal design**, but which has a high probability of finding an optimal, or very nearly optimal, design. Such searches are usually based on *interchange algorithms*, which arrange a given treatment design in blocks to optimize the relevant criterion. The general structure of an interchange algorithm is as follows:

1. Randomly allocate the treatments to blocks.
2. If the resulting design allows the model with fixed blocks to be fitted, evaluate the criterion and call this the *current design*. Otherwise, go to 1.
3. Systematically interchange treatments between units in different blocks, calculating the criterion for each design obtained. If the new design is better than the current design, it becomes the current design. Continue until no possible interchange improves the design.

4. The current design is the near-optimally blocked design.

Usually the whole procedure is restarted several times, called *tries*, and the best design from any try is kept. Interchange algorithms for block designs are discussed by Atkinson and Donev [5, Chapters 14 and 15].

Note that the criteria that are commonly optimized are based entirely on the intrablock analysis. This is most appropriate, since the interblock analysis is rarely useful (except as a last resort in trying to obtain information from badly designed experiments).

Factorial Treatments

When the treatments are combinations of levels of several factors as in our example, it is usually not the case that all treatment comparisons are of equal interest (see **Factorial Experiments**). Instead main effects and low-order interactions are usually more important and so the design should be arranged in blocks in such a way that these low-order effects can be well estimated even at the cost of very poor estimation of the higher-order effects. With some regular treatment structures and specific, convenient block sizes, standard designs exist. For example, our treatments are a single replicate of a 2^3 factorial design and there are two blocks of four units each. Then the design in Table 7 allows all main effects and two-factor interactions to be estimated with efficiency factors of 100%, while the three-factor interaction cannot be estimated in the intrablock stratum. The three-factor interaction is said to be *confounded* with block effects.

Note that the design in block I in Table 7 is a fractional factorial design (see **Fractional Factorial Designs**), specifically a half replicate, and the design in block II is the complementary half replicate. If two replicates of the 2^3 treatment set were to be

Table 7 Single replicate of a 2^3 treatment set in two blocks of four units

Block I			Block II		
X_1	X_2	X_3	X_1	X_2	X_3
-1	-1	-1	-1	-1	1
-1	1	1	-1	1	-1
1	-1	1	1	-1	-1
1	1	-1	1	1	1

used, we can confound the three-factor interaction in one replicate and a two-factor interaction in the other replicate, as in Table 8, where the two-factor interaction X_2X_3 is confounded in blocks III and IV. In the overall design, the effects X_2X_3 and $X_1X_2X_3$ can each be estimated with an efficiency factor of 50% and are said to be *partially confounded* with block effects.

Note that the information that is lost can be recovered in the interblock analysis, but that these effects will usually be estimated with high variance. For example, the analysis of variance for the design in Table 8 has the form shown in Table 9. A design with one or more main effects confounded with blocks is known as a *split-plot* design (see **Split-Plot Designs**). Mead [1, Chapter 15], Hinkelmann and Kempthorne [2, Chapter 11], and many other books on design of experiments describe confounding in more detail. Tables of standard confounded designs can be found in Cochran and Cox [4, Chapter 6] and Wu and Hamada [6, Chapter 3].

The theory of confounding extends to fractional factorial treatment sets (see **Fractional Factorial Designs**). For example, the resolution V (see **Factorial Designs, Resolution of**) half replicate of the 2^5 design with defining contrast $-X_1X_2X_3X_4X_5$ can be arranged in two blocks by confounding the two-factor interaction X_4X_5 . This automatically confounds this effect's alias, the three-factor interaction $X_1X_2X_3$ (see **Aliasing in Fractional Designs**). The design is given in Table 10. Again, published tables of these regularly confounded designs are available, for example, Wu and Hamada [6, Chapters 4 and 5], and

Table 8 Two replicates of a 2^3 treatment set in four blocks of four units

Block I			Block II		
X_1	X_2	X_3	X_1	X_2	X_3
-1	-1	-1	-1	-1	1
-1	1	1	-1	1	-1
1	-1	1	1	-1	-1
1	1	-1	1	1	1
Block III			Block IV		
X_1	X_2	X_3	X_1	X_2	X_3
-1	-1	-1	-1	-1	1
-1	1	1	-1	1	-1
1	-1	-1	1	-1	1
1	1	1	1	1	-1

6 Blocking

Table 9 Skeleton analysis of variance for the design of Table 8

Stratum	Source	<i>df</i>	<i>EF</i>
Blocks	X_2X_3	1	0.5
	$X_1X_2X_3$	1	0.5
	Error	1	
	Total	3	
Units	X_1	1	1
	X_2	1	1
	X_3	1	1
	X_1X_2	1	1
	X_1X_3	1	1
	X_2X_3	1	0.5
	$X_1X_2X_3$	1	0.5
	Error	5	
	Total	12	
Total		15	

several comprehensive statistical programs include general routines for constructing regular confounded designs. Recent work on confounded fractional factorial designs includes Zhu [7], Cheng and Tang [8], and Butler [9].

If the treatment design is not equally replicated, is an irregular fraction, or is chosen to be able to fit a polynomial model, or if the block sizes are such that no regular design exists, it is necessary to use a search procedure. This situation also arises if the “blocks” are individual stages in a sequential design (*see Sequential Experimentation*). Again, for small problems a complete search might be possible, but otherwise an interchange algorithm will be used. This works exactly as for unstructured treatments, except that slightly different optimality criteria will

Table 10 Half replicate of a 2^5 treatment set in two blocks of eight units

Block I					Block II				
X_1	X_2	X_3	X_4	X_5	X_1	X_2	X_3	X_4	X_5
−1	−1	−1	−1	−1	−1	−1	1	−1	1
−1	−1	−1	1	1	−1	−1	1	1	−1
−1	1	1	−1	−1	−1	1	−1	−1	1
−1	1	1	1	1	−1	1	−1	1	−1
1	−1	1	−1	−1	1	−1	−1	−1	1
1	−1	1	1	1	1	−1	−1	1	−1
1	1	−1	−1	−1	1	1	1	−1	1
1	1	−1	1	1	1	1	1	1	−1

be appropriate (*see Optimal Design*). It is usually necessary to specify a model and then arrange the design in blocks to optimize some function of the variance matrix of the parameter estimates. Useful optimality criteria include the following:

- *weighted-A-efficiency*, which minimizes a weighted average of the variances of the parameters of interest;
- *D-efficiency*, which minimizes the volume of a joint confidence region of the parameters of interest;
- *E-efficiency*, which maximizes the minimum efficiency factor among the parameters of interest; and
- *weighted-M-efficiency*, which maximizes a weighted mean of the efficiency factors of the parameters of interest.

A-efficiency is sometimes defined without weights, but this definition is scale dependent (e.g., if the coding of factor levels is changed, the definition of the criterion changes), and so should not be used without making clear what the scale is. Cook and Nachtsheim [10], Atkinson and Donev [11], Trinca and Gilmour [12], and Goos [13, Chapter 2] discuss algorithms for blocking in more detail.

Extensions

The idea of blocking can be extended to situations in which there is more than one blocking factor, that is, there is more than one way of defining sets of units which can be expected to be similar and the randomization is restricted to allow for all of them.

The first possibility is *nested* blocking factors, in which the units are grouped together in blocks and the blocks themselves are grouped into superblocks. For example, if the units are runs of a process and the blocks are days, then it might be that the

Table 11 A 4×4 Latin square

Row	Column			
	I	II	III	IV
I	1	2	3	4
II	2	3	4	1
III	3	4	1	2
IV	4	1	2	3

experiment is run on four days, two consecutive days and then, somewhat later, another two consecutive days. In such a case, it is often sensible to restrict the randomization of blocks to days in such a way that a good design is obtained with respect to the superblocks. For example, the blocks in Table 8 can be grouped into two superblocks, one containing blocks I and II and one containing blocks III and IV. In this case, each superblock contains a single replicate of each treatment and the design is said to be *resolved* (often only *resolvable* designs, i.e., those that can be resolved, are defined, but this is a purely mathematical concept). The idea of nesting can, of course, be generalized to any number of levels, but it is unusual to have more than three strata, except in generalizations of split-plot designs (see **Split-Plot Designs**) in which the main effects of different factors are confounded at different strata. Bailey [14] and Gilmour and Trinca [15] discuss nested block designs further.

Another possibility is *crossed* blocking factors, in which the units are grouped into two sets of blocks, often called *rows* and *columns*, so that each row-column combination contains a unit. For example, a study of process operating conditions might involve several runs made each day with each run made by a different operator, so that days are rows and operators are columns. Then a design should be chosen which is a good block design with respect to rows and a good block design with respect to columns. If the numbers of rows and columns are both equal to the numbers of treatments, then a *Latin square* can be used, for example, the 4×4 Latin square shown in Table 11 is a design for comparing four treatments in four rows and four columns (see **Latin Squares and Related Experimental Designs**). For factorial designs, confounding techniques can be used for some sizes of rows and columns, but in general a search algorithm will be needed. The idea of crossing can also be extended to more than two blocking factors, but this is unusual in practice. Algorithms for crossed blocking are described by Gilmour and Trinca [16].

References

- [1] Mead, R. (1988). *The Design of Experiments*, Cambridge University Press.
- [2] Hinkelmann, K. & Kempthorne, O. (1994). *The Design and Analysis of Experiments*, John Wiley & Sons, Vol. 1.
- [3] John, J.A. & Williams, E.R. (1995). *Cyclic and Computer Generated Designs*, 2nd Edition, Chapman & Hall.
- [4] Cochran, W.G. & Cox, G.M. (1957). *Experimental Designs*, 2nd Edition, John Wiley & Sons.
- [5] Atkinson, A.C. & Donev, A.N. (1992). *Optimum Experimental Designs*, Oxford University Press.
- [6] Wu, C.F.J. & Hamada, M. (2000). *Experiments*, John Wiley & Sons.
- [7] Zhu, Y. (2003). Structure function for aliasing patterns in $2^{(l-n)}$ design with multiple groups of factors, *Annals of Statistics* **31**, 995–1011.
- [8] Cheng, C.-S. & Tang, B.X. (2005). A general theory of minimum aberration and its applications, *Annals of Statistics* **33**, 944–958.
- [9] Butler, N.A. (2006). Optimal blocking of two-level factorial designs, *Biometrika* **93**, 289–302.
- [10] Cook, R.D. & Nachtsheim, C.J. (1989). Computer-aided blocking of factorial and response-surface designs, *Technometrics* **31**, 339–346.
- [11] Atkinson, A.C. & Donev, A.N. (1989). The construction of exact D-optimum experimental designs with applications to blocking response surface designs, *Biometrika* **76**, 515–526.
- [12] Trinca, L.A. & Gilmour, S.G. (2000). An algorithm for arranging response surface designs in small blocks, *Computational Statistics and Data Analysis* **33**, 25–43.
- [13] Goos, P. (2002). *The Optimal Design of Blocked and Split-Plot Experiments*, Springer.
- [14] Bailey, R.A. (1999). Choosing designs for nested blocks, *Listy Biometryczne* **36**, 85–126.
- [15] Gilmour, S.G. & Trinca, L.A. (2006). Response surface experiments on processes with high variation, in *Response Surface Methodology and Related Topics*, A.I. Khuri, ed, World Scientific.
- [16] Gilmour, S.G. & Trinca, L.A. (2003). Row-column response surface designs, *Journal of Quality Technology* **35**, 184–193.

STEVEN G. GILMOUR

Computer Experiments

Introduction

Investigators in virtually all areas of science and engineering application have become dependent on *computer models* as tools in designing, testing, and evaluating physical systems. The physical systems of interest range from manufacturing supply chains to complex chemical reactions to the global climate. Computer models can be based on theoretical knowledge of the physical system of interest, empirically derived relationships, or both. In any case, they are computer programs that evaluate or solve mathematical equations written to describe the physical system of interest. Hence a global climate model is a program written to simulate or predict climate patterns years or centuries in the future, on the basis of specified *initial conditions* representing the state of the climate system at the beginning of simulated time and *forcing functions* representing possibly changing aspects of natural and man-made influences on the system.

Because these models are often quite complex, their behavior can sometimes be surprising, even to their authors. Hence empirical studies based on the model, or *computer experiments*, may be performed for a number of purposes. The phrase “computer experiment” is intended to suggest an activity similar in many ways to traditional experimentation, but where the computer model is (at least one of) the object(s) of interest. Following standard experimental language, *treatments* correspond to the values selected for model inputs or other quantities that must be specified before the model can be evaluated or “run”, and *observations* correspond to the values of the outputs, or summary functions of them, resulting from such runs. *Experimental design* in this context primarily involves the selection of input vectors or specific model runs to be included in the experiment. As with traditional or *physical* experiments, the overall goal of a computer experiment is generally the collection of information that will support an empirical understanding of the relationship between inputs and outputs.

While it is useful to view computer experiments and physical experiments as similar, there are important differences. One of these differences is the treatment of experimental error. Deterministic

computer models produce exactly the same output values if run multiple times under the same conditions. This means that the usual emphasis on replication, **randomization**, and **blocking** in physical experiments is not relevant in computational studies. (An important exception is with computer models, such as discrete-event simulators, that include stochastic simulation as part of the calculation; replication often *does* make sense in studying these models.) Another important difference is the large dimension of the input and output vectors in many computer models. It is not unusual for computer experiments to involve tens, hundreds, or even thousands of input and/or output variables, while most physical experiments (except perhaps in certain screening contexts) involve far fewer factors and responses.

Categories of Computer Experiments

The reasons for performing computer experiments are as many and as varied as the reasons for performing physical experiments. In this section, three general categories of computer experiments are briefly described. The first two correspond to what might be called *factor screening* and *follow-up studies*, respectively, in a physical experiment context. The third category involves joint consideration of the model and the physical system it mimics.

Assessment of the Relative Importance of Inputs

As with experimental programs in physical contexts, one early goal is to identify the factors/inputs that are most important in determining the values (or changes in values) of the outputs. This is perhaps even more critical in computer experiments, due to the very large number of input variables often encountered. One popular experimental paradigm is based on the assignment of a probability distribution to the input vector and the selection of input values through restricted random processes designed to allow **estimation** of characteristics of the resulting distribution of outputs. This exercise can be thought of as an empirical “transformation of variables” problem and is often performed to identify how much of the propagated variance might be eliminated by reducing uncertainty in each individual input or in groups of related

2 Computer Experiments

inputs. These techniques are often referred to as *sensitivity analysis* or *uncertainty analysis*; several of them are described in [1], along with case studies from a number of application areas. Other methods, which might be called *quasi-modeling*, are based on using rank-transformed data in traditional linear regression or **correlation** analyses; ranks are used to linearize monotonic relationships between inputs and outputs, for example [2]. Group screening ideas have also been used in computer experiments, usually in sequential form, to efficiently identify those inputs that require more intensive follow-up, for example [3].

Quantitative Understanding of a Computer Model

After the subset of most influential inputs has been identified, experimental emphasis often progresses to a more detailed investigation of how these inputs affect outputs. Construction of an approximation to the model function, sometimes called a *metamodel* or *surrogate*, is often desirable at this point, to facilitate rapid output approximations for applications that require a large number of output values (e.g., **Monte Carlo** studies in which the computer model is evaluated inside the simulation loop), or to produce a simpler functional form to aid in understanding the model. A series of such model surrogates is used as the basis of an optimization procedure in [4], with an application involving the design of a helicopter rotor blade. Linear regression models have sometimes been used for this purpose; an early contribution [5] describes extended **central composite** designs for third-order polynomial models and demonstrates their use in a computer experiment. However, much of the statistical research in recent years has focused on the use of Gaussian process models, similar to the Kriging models used in geostatistics, rather than traditional linear models, for example [6] and [7]. Gaussian process models have generally preferable properties as metamodels for deterministic computer models because they can be easily constructed to interpolate (rather than smooth) the noiseless output data, and are much more flexible (and so can more faithfully reproduce nonlinear and asymptotic behavior) than traditional regression models. Excellent detailed references for the use of Gaussian stochastic processes in this setting can be found in [8] and [9]. A demonstration of how such models can be used both to screen inputs and to

explore the structure of input–output relationships through functional analysis of variance (**ANOVA effects**), lower-dimensional components of the predicted model function, is given in [10]. Designs for fitting Gaussian process models have been constructed using several criteria including entropy [11] and asymptotically justified distance-based indices [12].

Understanding and Using Relationships between a Model and Its Context

Everything said to this point has emphasized experimental activity directed at understanding the computer model *per se*, rather than the physical system it mimics. Since the latter is the object of ultimate interest, experimental effort eventually moves to studies involving both. For these purposes, suppose the computer model is denoted

$$y = M(\mathbf{x}, \boldsymbol{\theta}) \quad (1)$$

where y is the output of interest, $\mathbf{x}$ is the corresponding vector of inputs, and $\boldsymbol{\theta}$ are model parameters such as physical rate coefficients. In most real applications, many output variables are actually of interest, but we limit this description to a scalar value for simplicity. The model is designed to mimic *reality* where, for example, y_R results from $\mathbf{x}_R$. Computer experiments requiring both physical measurement y_R and model output y may be carried out for the following purposes:

- **Model calibration** – determination of $\boldsymbol{\theta}$ so that y matches y_R as well as possible when $\mathbf{x}$ matches $\mathbf{x}_R$. Joint analysis of physical and model data for calibration is developed in [13] and demonstrated using computer models of a charged particle accelerator and a spot welding process.
- **Model validation** – for a calibrated model, determination of how well y matches y_R when $\mathbf{x}$ matches $\mathbf{x}_R$. A validation study involving CORSIM, a traffic simulation model, using physical traffic data collected in an urban setting is reported in [14].
- **Inverse problems** – for a calibrated and validated model, determination of $\mathbf{x}$ leading to $y = y_R$ when $\mathbf{x}_R$ is unknown. A Gaussian stochastic process is used in [15] to make inference about the initial conditions ($\mathbf{x}$) leading to specified

observations (y) in a one-dimensional fluid flow problem.

These descriptions have been overly simplified for illustration; in each case, single unique solutions generally are not available and uncertainty characterization is critical.

Experiments involving both model outputs and physical data generally must account for the fact that computer models are not perfect representations of reality. In [16], a metamodel for the computer model and an estimate of the “discrepancy function”, or difference between y and y_R as a function of $\mathbf{x}$, are simultaneously constructed. This approach can also be used for the joint analysis of multiple computer models of a common physical system. The broader problem of combining data from computer models, physical experiments, and **expert opinion** in a unified analysis is considered in [17].

Gaussian Process Models

As noted in previous section, much of the most recent work in the design and analysis of computer experiments is based on the use of a Gaussian stochastic process model as a statistical surrogate for the computer model. (See, e.g., [8] or [9] for a detailed treatment of these models.) Continuing with the notation of previous section, the output y is modeled *a priori* as a *random function*, normally distributed at each $\mathbf{x}$ and jointly multivariate normal at any collection of $\mathbf{x}$'s, and characterized by functional representations of the mean and variance:

$$\begin{aligned} E(y(\mathbf{x})) &= \mu(\mathbf{x}), \quad V(y(\mathbf{x})) = \sigma^2(\mathbf{x}), \\ \text{Cov}(y(\mathbf{x}), y(\mathbf{x}')) &= C(\mathbf{x}, \mathbf{x}'), \quad \mathbf{x}, \mathbf{x}' \in \Delta \end{aligned} \quad (2)$$

where Δ is the space of allowable values of the input vector. In practice, the second-order **moments** are usually assumed to be stationary, and the mean function often is as well:

$$\begin{aligned} E(y(\mathbf{x})) &= \mu, \quad V(y(\mathbf{x})) = \sigma^2, \\ \text{Cov}(y(\mathbf{x}), y(\mathbf{x}')) &= \sigma^2 R(|\mathbf{x} - \mathbf{x}'|), \quad \mathbf{x}, \mathbf{x}' \in \Delta \end{aligned} \quad (3)$$

where $|\mathbf{x} - \mathbf{x}'|$ denotes the vector of absolute values of the elements of $\mathbf{x} - \mathbf{x}'$, and $R(-)$ is a *correlation function*. R is often a parameterized function, for example,

$$R(|\mathbf{x} - \mathbf{x}'|) = \exp(-|\mathbf{x} - \mathbf{x}'|^T \Lambda |\mathbf{x} - \mathbf{x}'|) \quad (4)$$

where $\Lambda = \text{diag}(\lambda_1, \lambda_2, \lambda_3, \dots, \lambda_k)$ and k is the dimension of $\mathbf{x}$, is called the *product Gaussian correlation function*, with parameters (λ 's) that control how quickly the positive correlation between output values decreases with distance in each coordinate of the input vector. An experimental design of size N , $D = \{\mathbf{x}_1, \mathbf{x}_2, \mathbf{x}_3, \dots, \mathbf{x}_N\}$, is selected, the corresponding runs of the model are executed, and an N vector of outputs $\mathbf{y}_D$ results. For this model and known parameters, the frequentist best linear unbiased predictor (BLUP), or Bayesian prediction based on squared-error loss, of y at a new input vector $\mathbf{x}_0$ is as follows:

$$\hat{y}(\mathbf{x}_0) = \mu + \mathbf{r}_{\mathbf{x}_0, D}^T \mathbf{R}_{D, D}^{-1} (\mathbf{y}_D - \mu \mathbf{1}) \quad (5)$$

where $\mathbf{r}_{\mathbf{x}_0, D}$ is an N vector of correlations, $R(|\mathbf{x}_0 - \mathbf{x}_1|)$, $R(|\mathbf{x}_0 - \mathbf{x}_2|)$, $R(|\mathbf{x}_0 - \mathbf{x}_3|)$, $\dots$, $R(|\mathbf{x}_0 - \mathbf{x}_N|)$, and $\mathbf{R}_{D, D}$ is the $N \times N$ correlation matrix with (i, j) element $R(|\mathbf{x}_i - \mathbf{x}_j|)$. The corresponding **standard error** of prediction, or posterior standard deviation, is

$$\text{se}(\hat{y}(\mathbf{x}_0)) = \sigma \sqrt{1 - \mathbf{r}_{\mathbf{x}_0, D}^T \mathbf{R}_{D, D}^{-1} \mathbf{r}_{\mathbf{x}_0, D}} \quad (6)$$

As noted, these expressions assume that μ , σ^2 , and any parameters of $R(-)$ are known. Simple extensions apply when μ is unknown; for example, the expression for $\hat{y}(\mathbf{x}_0)$ in which $\hat{\mu} = \mathbf{y}_D^T \mathbf{R}_{D, D}^{-1} \mathbf{1} / \mathbf{1}^T \mathbf{R}_{D, D}^{-1} \mathbf{1}$ is substituted for μ in the frequentist BLUP, but this still requires knowledge of the correlation parameter values. Since it is usually not reasonable to regard these parameters as known quantities in practice, fitting a Gaussian process model usually involves estimation of the stochastic process parameters, often by the method of maximum likelihood, or through a more comprehensive Bayesian treatment employing priors on the unknown parameters. In any case, the aim is to provide simple predictions of what y would be at any $\mathbf{x}_0$ required, along with a useful measure of uncertainty for that prediction. For any form of correlation function for which $R(\mathbf{0}) = 1$, the prediction exactly reproduces the output and yields standard errors of zero at the design points. When R is a smooth function of its arguments, predictions and standard errors vary smoothly over other points in the input space. The Gaussian stochastic process model also provides a structure within which good experimental designs can be constructed for settings that require derivation of a simulator of M , for example [11] and [12]. Extensions of the basic theory have also been developed

4 Computer Experiments

for settings in which the output is vector valued [18], derivatives of the output with respect to the inputs are computed [19], ancillary *truncation errors* computed in the numerical solution of model equations are available [20], and multiple computer models of the same physical system are used [21].

Experimental Design

As noted above, the design component of a computer experiment often is mainly concerned with the selection of input vectors specifying N runs of the computer model. Perhaps the most popular design used is the Latin hypercube (*see* [22], **Latin Hypercube Designs**), which was initially proposed as an input *sampling* plan for probability-weighted output integration, but is also widely used when the primary goal of the experiment is metamodeling (or model approximation). Especially when the planned analysis is based on a Gaussian stochastic process or other “spatial” modeling methodology, experimental designs that are in some sense “space-filling” are natural choices since they offer “local” support to flexible modeling throughout the experimental region. For this reason, various forms of *uniform designs* (*see* **Uniform Experimental Designs**) have also been proposed for computer experiments. By contrast, key properties of traditional experimental designs, such as orthogonality, may not be so important in this context when linear models are not the basis of data analysis. However, **orthogonal array**-based Latin hypercubes described by [23] use orthogonal structure to insure that Latin hypercubes cover low-dimensional subspaces of the experimental region well.

Optimal design arguments specifically formulated for the Gaussian stochastic process model include entropy (or D optimality for prediction throughout Δ) [11], and distance-based criteria such as the *maximin* (*Mm*) *distance criterion* which has been shown in [12] to be asymptotically equivalent to the entropy criterion under certain conditions. *Mm*, which stipulates that the closest pairs of design points should be as far apart as possible and that the number of point pairs in the design separated by this minimum distance should be minimized, is also intuitively appealing for computer experiments due to the spatial characteristics. The *Mm* criterion has also been used to generate designs within attractive classes, for example, Latin hypercubes [24].

Integrated mean square error of prediction has also been used as an optimality criterion for the construction of designs [25, 26]. A more detailed review of experimental designs for computer experiments is offered in [9].

A Brief Example

Mitchell and Morris [27] presented data from a computer experiment carried out to demonstrate Gaussian process modeling methodology, using a model of methane combustion developed by Frenklach and Rabinowitz [28]. For this demonstration, the output response of interest is the logarithm of ignition delay in the modeled system. The experimentally controlled model inputs are seven reaction rates of interest; the ranges of the reaction rates were determined by physical knowledge of the system, and each was transformed to the unit interval for this exercise. A design of 50 input vectors was selected from the points on a full five-level **factorial** point set, $\{0, \frac{1}{4}, \frac{1}{2}, \frac{3}{4}, 1\}^7$, using a *maximum entropy* criterion associated with the Gaussian process model, essentially finding the design for which the determinant of $\mathbf{R}_{D,D}$ is maximized. Note that this requires specification of the form and parameters of the spatial correlation function. For reasons described in the paper, the authors used a product exponential correlation function:

$$R(\|\mathbf{x} - \mathbf{x}'\|) = \exp(-\mathbf{1}^T \mathbf{\Lambda} \|\mathbf{x} - \mathbf{x}'\|), \quad \mathbf{\Lambda} = 4.6 \times \mathbf{I} \quad (7)$$

as the basis of computing an experimental design. That design, along with the 50 response values computed from the corresponding runs of the computer model, are listed in Table 1.

With the data in hand, the authors used a *piecewise cubic correlation* function described in [7] for model fitting and prediction. In this reanalysis, a Gaussian correlation function was assumed, leading to maximum likelihood estimates of the correlation parameters:

$$\begin{aligned} \hat{\lambda}_1 &= 0.085, \hat{\lambda}_2 = 0.024, \hat{\lambda}_3 = 0.498, \hat{\lambda}_4 = 0.363, \\ \hat{\lambda}_5 &= 0.347, \hat{\lambda}_6 = 0.167, \hat{\lambda}_7 = 0.090 \end{aligned} \quad (8)$$

These parameter estimates are themselves useful indices of “importance” for each input, because they

Table 1 Input and output values for the methane combustion example

Scaled input							Output
1	2	3	4	5	6	7	
0.75	0.00	0.75	0.75	0.00	1.00	1.00	5.0056
1.00	0.50	1.00	0.75	0.00	0.25	0.50	5.1090
1.00	0.00	1.00	0.25	0.00	1.00	0.00	5.3045
0.50	0.00	1.00	0.00	0.00	0.50	0.75	5.3423
1.00	1.00	1.00	1.00	0.50	1.00	0.00	5.3495
0.25	0.50	1.00	0.75	0.50	1.00	1.00	5.4478
0.50	1.00	0.75	1.00	0.25	0.75	1.00	5.5674
1.00	0.25	0.75	0.00	0.25	0.00	1.00	5.6656
1.00	0.25	1.00	1.00	1.00	0.25	0.00	5.7137
0.00	1.00	1.00	0.50	0.25	1.00	0.50	5.8572
0.75	0.75	1.00	0.00	0.00	0.00	0.00	5.8651
1.00	0.00	0.75	0.75	0.50	0.75	0.25	5.8721
0.50	0.00	0.25	1.00	0.00	0.25	0.75	6.0216
1.00	1.00	0.50	0.25	0.00	1.00	1.00	6.0325
0.50	0.00	1.00	0.25	1.00	0.75	1.00	6.2131
0.25	0.50	0.50	0.75	0.00	1.00	0.00	6.2171
1.00	0.25	0.50	1.00	1.00	1.00	0.50	6.2543
0.00	0.75	0.75	1.00	1.00	1.00	0.50	6.3491
1.00	0.50	0.00	1.00	0.00	0.75	1.00	6.3746
1.00	0.00	1.00	0.00	0.00	0.75	0.00	6.4065
0.25	0.00	1.00	0.50	0.75	0.50	0.00	6.4489
0.50	1.00	0.00	1.00	0.00	1.00	0.25	6.4493
0.25	0.00	0.50	0.25	0.25	0.00	1.00	6.5214
0.00	0.00	0.00	1.00	0.25	1.00	1.00	6.5603
1.00	0.50	1.00	0.00	1.00	1.00	0.50	6.5656
0.25	1.00	0.75	0.75	0.50	0.00	0.25	6.5930
0.75	1.00	0.75	0.25	0.25	0.75	0.25	6.7111
0.00	0.75	0.50	0.75	0.25	0.25	0.50	6.7319
0.75	0.75	0.25	0.50	0.00	0.25	1.00	6.7549
1.00	0.25	0.00	0.50	0.00	0.50	0.00	6.8957
0.25	1.00	0.25	1.00	0.75	1.00	0.75	6.9749
1.00	1.00	0.00	1.00	0.25	0.00	0.50	7.2242
0.25	0.00	0.00	0.25	0.00	0.75	0.50	7.3542
0.00	0.25	0.25	1.00	0.75	0.25	0.75	7.4006
0.50	0.50	0.50	0.50	0.50	0.00	0.00	7.4111
0.50	0.00	0.50	0.00	0.75	1.00	0.25	7.5044
0.75	0.50	0.25	0.50	1.00	1.00	1.00	7.5563
0.50	0.50	0.75	0.00	1.00	0.25	0.00	7.5708
0.75	1.00	0.25	0.75	1.00	0.75	0.25	7.6225
1.00	0.00	0.00	0.50	0.50	0.50	1.00	7.6311
1.00	0.75	0.00	0.00	0.50	1.00	0.75	7.6953
0.00	0.50	0.50	0.00	0.75	0.50	0.75	7.7907
0.00	1.00	0.00	0.25	0.00	0.00	1.00	7.8535
0.75	0.25	0.00	1.00	1.00	0.00	1.00	7.8710
0.50	0.75	0.25	0.00	0.25	0.50	0.50	7.8725
1.00	0.50	0.25	0.25	0.75	0.75	0.00	7.9182
0.00	0.00	0.00	0.50	1.00	1.00	0.25	7.9315
1.00	0.00	0.00	0.00	1.00	0.25	0.00	8.2060
0.00	1.00	0.25	0.25	0.75	1.00	0.00	8.4206
0.00	1.00	0.00	0.75	1.00	0.50	1.00	8.5372

represent the rate at which correlation appears to decrease with distance in each dimension of the input space. For example, correlation parameter estimates associated with inputs 3, 4, and 5 are relatively large, suggesting that the output may vary more quickly, or in a more “interesting” fashion, as these inputs are varied.

The nature of the input–output relationship can be investigated more directly by use of plots of $\hat{y}$ as functions of one or more elements of $\mathbf{x}$. Panels 1–7 of Figure 1 each display 20 “traces” of $\hat{y}$, the predicted response plotted as a function of one input for fixed, randomly selected values of the other six inputs. Plots such as these display the general shape of the relationship between each input and the response, and also provide an indication of how much that input “interacts” with the others based on how much the shape of the relationship changes from trace to trace. These plots support our initial impression, based on estimates of the correlation parameters, that reaction rates 3, 4, and 5 are most influential in the determination of ignition delay. A further examination of the joint influence of these three inputs can be made by examining contour plots of the output for pairs of inputs. Panels 8 and 9 of Figure 1 display such joint plots for inputs 3 and 4, and 3 and 5 respectively, with all other inputs fixed at the midpoints of their coded scales (0.5).

The general appearance of the plots in Figure 1 suggests that a quadratic polynomial might also serve as a useful model approximation. A full quadratic polynomial, fitted by **least squares** to these 50 data points, does account for most of the variation in the output values, as indicated by $R^2 = 99.6\%$. **Residuals** from this polynomial fit vary between -0.18 and $+0.11$ (compared to a total data range of 3.53), so within these tolerances, a fitted polynomial may well be a reasonable descriptive approximation of the computer model. On the other hand, the Gaussian process predictor reproduces the 50 data values without error. The standard errors constructed under the Gaussian process model are also reasonable representations of prediction uncertainty, for example, they are exactly 0 at the 50 points where the response is known, and are larger at more “remote” points in the design space. Standard errors of prediction constructed in the linear regression framework have no clear interpretation since they are based on

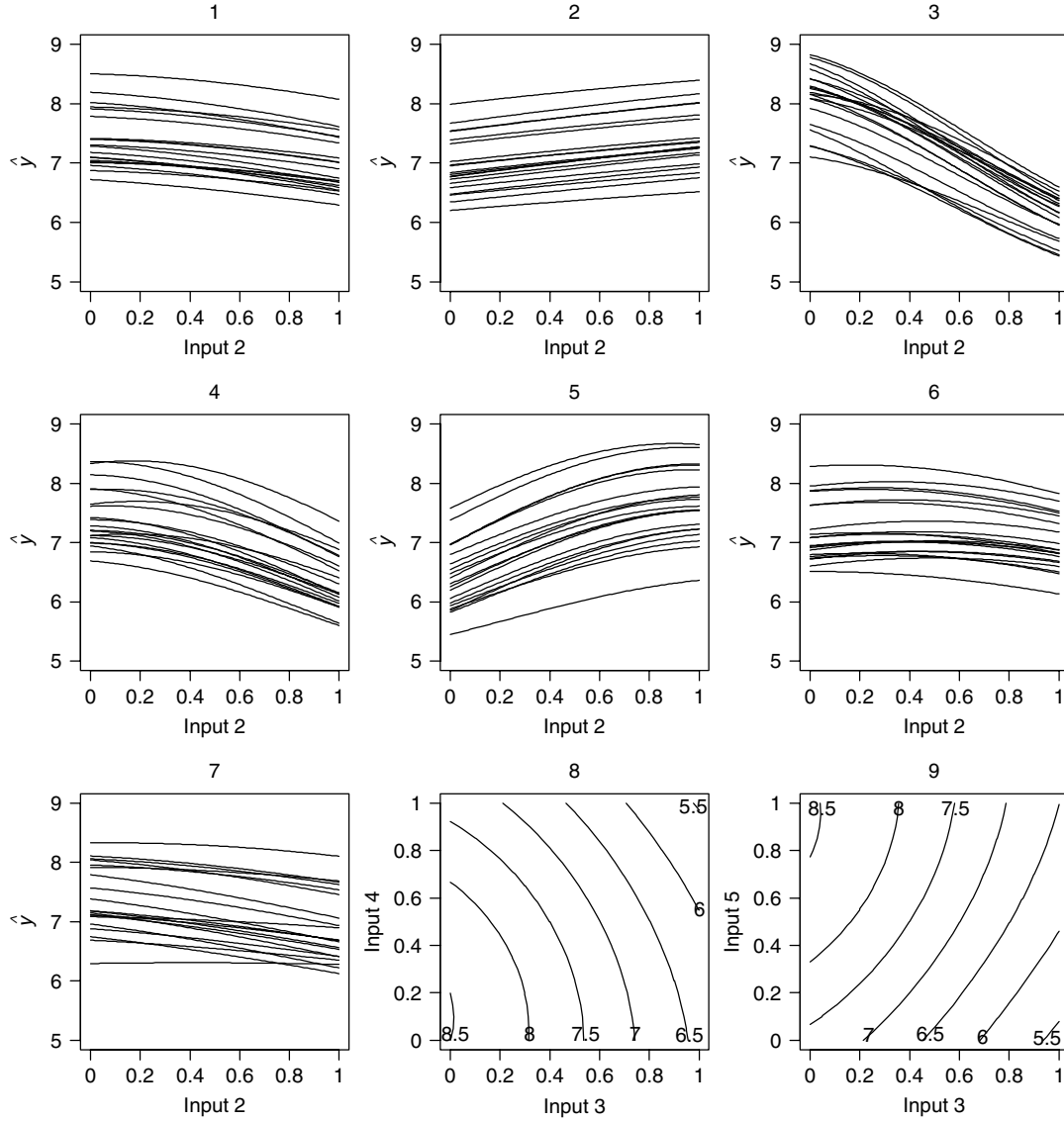


Figure 1 Graphs displaying the form of $\hat{y}$ from the Gaussian process model as a function of scaled inputs. Panels 1–7: 20 traces of $\hat{y}$ as a function of a single input, each with randomly chosen values of the remaining inputs. Panels 8 and 9: contour plots of $\hat{y}$ as a function of inputs 3 and 4, and 3 and 5, with all other inputs fixed at 0.5

assumptions about random “error” that are not realistic in this context.

Summary and Conclusion

Computer experiments are becoming increasingly important in engineering, industry, and the applied sciences. As physical systems of interest grow more

complex and computing power increases, this trend can be expected to continue. Statistical ideas and methods are critical to the design and analysis of these studies, just as they are in traditional physical experimentation. These methods do not necessarily take exactly the same form in both realms, but the differences in effective methods generally reflect the real differences between the properties of the

data produced in the two kinds of experimentation. While computer experiments often do not involve measurement error and do not require traditional strategies such as blocking and randomization, they often involve much larger input spaces and more complex factor (input) versus response (output) relationships than do most physical experiments. An additional difficulty in computer experiments is the fact that the physical system giving rise to creation of the model is usually the object of real interest, rather than the computer model itself. Effective statistical methods, many based on Gaussian stochastic process models, have been developed for use in computer experiments over the last few decades, and further developments in this area will continue to be an important area of research in the future.

References

- [1] Saltelli, A., Chan, K. & Scott E.M. (eds) (2000). *Sensitivity Analysis*, John Wiley & Sons, New York.
- [2] Iman, R.L. & Conover, W.J. (1980). Small-sample sensitivity analysis techniques for computer models, with an application to risk assessment (with discussion), *Communications in Statistics-Theory and Methods* **A9**, 1749–1874.
- [3] Kleijnen, J.P.C., Bettonvil, B. & Persson, F. (2006). Screening for the important factors in large discrete-event simulation models: sequential bifurcation and its applications, in *Screening: Methods for Experimentation in Industry, Drug Discovery, and Genetics*, A. Dean & S. Lewis, eds, Springer, New York.
- [4] Booker, A.J., Dennis, J.E., Frank, P.D., Serafini, D.B., Torczon, V. & Trosset, M.W. (1999). A rigorous framework for optimization of expensive functions by surrogates, *Structural Optimization* **17**, 1–13.
- [5] Baker, F.D. & Bargmann, R.E. (1985). Orthogonal central composite designs of the third order in the evaluation of sensitivity and plant growth simulation models, *Journal of the American Statistical Association* **80**, 574–579.
- [6] Sacks, J., Welch, W.J., Mitchell, T.J. & Wynn, H.P. (1989). Design and analysis of computer experiments (with discussion), *Statistical Science* **4**, 409–423.
- [7] Currin, C., Mitchell, T.J., Morris, M.D. & Ylvisaker, D. (1991). Bayesian prediction of deterministic functions, with applications to the design and analysis of computer experiments, *Journal of the American Statistical Association* **86**, 953–963.
- [8] Koehler, J.R. & Owen, A.B. (1996). Computer experiments, in *Handbook of Statistics*, S. Ghosh & C.R. Rao, eds, Elsevier, Amsterdam, Vol. 13.
- [9] Santner, T.J., Williams, B.J. & Notz, W.I. (2003). *The Design and Analysis of Computer Experiments*, Springer-Verlag, New York.
- [10] Schonlau, M. & Welch, W.J. (2006). Screening the input variables to a computer model via analysis of variance and visualization, in *Screening: Methods for Experimentation in Industry, Drug Discovery, and Genetics*, A. Dean & S. Lewis, eds, Springer, New York.
- [11] Shewry, M.C. & Wynn, H.P. (1987). Maximum entropy sampling, *Journal of Applied Statistics* **14**, 165–170.
- [12] Johnson, M.E., Moore, L.M. & Ylvisaker, D. (1990). Minimax and maximin distance designs, *Journal of Statistical Planning and Inference* **26**, 131–148.
- [13] Higdon, D., Kennedy, M., Cavendish, J.C., Cafoe, J.A. & Ryne, R.D. (2004). Combining field data and computer simulations for calibration and prediction, *SIAM Journal on Scientific Computing* **26**, 448–466.
- [14] Sacks, J., Roupail, N.M., Park, B. & Thakuria, P. (2002). Statistically-based validation of computer simulation models in traffic operations and management (with discussion), *Journal of Transportation and Statistics* **5**, 18–22.
- [15] Morris, M.D. & Solomon, A.D. (1995). Design and analysis for an inverse problem arising from an advection-dispersion process, *Technometrics* **37**, 293–302.
- [16] Kennedy, M.C. & O'Hagan, A. (2000). Predicting the output from a complex computer code when fast approximations are available, *Biometrika* **87**, 1–13.
- [17] Reese, C.S., Ryan, K.J., Wilson, A.G., Hamada, M. & Martz, H.F. (2004). Integrated analysis of computer and physical experiments, *Technometrics* **46**, 153–164.
- [18] Drignei, D. (2006). Empirical Bayesian analysis for high-dimensional computer output, *Technometrics* **48**, 230–240.
- [19] Morris, M.D., Mitchell, T.J. & Ylvisaker, D. (1993). Bayesian design and analysis of computer experiments: use of derivatives in surface prediction, *Technometrics* **35**, 243–255.
- [20] Drignei, D. & Morris, M.D. (2006). Empirical Bayesian analysis for computer experiments involving finite difference codes, *Journal of the American Statistical Association* **101**, 1527–1536.
- [21] Kennedy, M.C. & O'Hagan, A. (2001). Bayesian calibration of computer models, *Journal of the Royal Statistical Society, Series B* **63**, 425–262.
- [22] McKay, M.D., Beckman, R.J. & Conover, W.J. (1979). A comparison of three methods for sampling values of input variables in the analysis of output from a computer code, *Technometrics* **21**, 239–245.
- [23] Tang, B. (1993). Orthogonal array based Latin hypercube sampling, *Journal of the American Statistical Association* **88**, 1392–1397.
- [24] Morris, M.D. & Mitchell, T.J. (1995). Exploratory designs for computational experiments, *Journal of Statistical Planning and Inference* **43**, 381–402.
- [25] Sacks, J., Schiller, S.B. & Welch, W.J. (1989). Designs for computer experiments, *Technometrics* **31**, 41–47.

- [26] Park, J.-S. (1994). Optimal Latin-hypercube designs for computer experiments, *Journal of Statistical Planning and Inference* **39**, 95–111.
- [27] Mitchell, T.J. & Morris, M.D. (1992). Bayesian design and analysis of computer experiments, *Statistica Sinica* **2**, 359–379.
- [28] Frenklach, M. & Rabinowitz, M. (1989). Optimization of large reaction systems, *Proceedings of the 12th IMACS World Congress on Scientific Computation*, Gerfihn, Paris, Vol. 3, pp. 602–604.

Related Article

Latin Hypercube Designs.

MAX D. MORRIS

Factorial Experiments

Introduction

Statistically designed experiments have been used to accelerate learning since their introduction by R.A. Fisher in the first half of the twentieth century. Fisher's book, *The Design of Experiments* [1], revolutionized agricultural research by providing methods, now used all over the world, for evaluating the results of small sample experiments [2]. One of Fisher's major contributions was the development of factorial experiments, which can simultaneously study several factors. Such experiments ran completely counter to the common wisdom of isolating one factor at a time, with all other factors held constant. Quoting R.A. Fisher: "No aphorism is more frequently repeated in connection with field trials, than that we must ask Nature few questions, or, ideally, one question, at a time. The writer is convinced that this view is wholly mistaken. Nature, he suggests, will best respond to a logical and carefully thought out questionnaire. A factorial design allows the effect of several factors and interactions between them, to be determined with the same number of trials as are necessary to determine any one of the effects by itself with the same degree of accuracy." [3].

New challenges and opportunities arose when Fisher's ideas were applied in the industrial environment. The opportunity for rapid feedback led G.E.P. Box to develop **response surface** methodology, with new families of designs and a new strategy of **sequential experimentation**. Box's contributions were stimulated by problems that arose during his 8 years of work as a statistician for Imperial Chemical Industries (ICI). The book by G.E.P. Box, W.G. Hunter and J.S. Hunter, *Statistics for Experimenters* [4], has become a classic in the area of design of experiments. Statistically designed experiments are now recognized as essential for rapid learning and thereby, for reducing time-to-market, while preserving high quality and peak performance. We present here a brief review of factorial experiments or statistically designed experiments. We focus mostly, but not exclusively, on factorial designs and cover several areas of applications of such designs.

Factorial designs have been extensively studied in the literature. A partial list of references includes [1, 4–16].

Basic Elements of Factorial Designs

The first step in planning any experiment is to clearly define its goals. The next step is to determine how performance will be assessed. The third step is to list the controllable factors that may affect performance. It is very useful to carry out these steps in a "brainstorming" session, bringing together individuals who are familiar with different aspects of the process under study. Typically, the list of potential factors will be long and often there will be disagreements about which factors are most important. The experimental team needs to decide which factors are indeed the most important, which of those will be systematically varied in the experiment, and how to treat the ones that will be left out. For each factor that is included in the experiment, the team must decide on its experimental range. Some factors may be difficult or impossible to control during actual production or use, but can be controlled for the purpose of an experiment. Such factors are called *noise factors* (see **Robust Design**) and it can be beneficial to include them in the experiment. In general, for an experiment to be considered statistically designed, the following nine issues need to be ascertained (Table 1). Some of these issues, such as those on determination of the experimental array (6) and the number of replications (7), are purely statistical in nature. Some other points deal with basic problem solving and management such as statement of the problem (point 1) and budget control (point 9).

A typical example is reported on by Milman, Sirota and Steinberg [17]. The goal of this experiment was to discover how required safety modifications would affect performance characteristics of a pyrotechnic device. The test results for each experimental unit were summarized in a graph of pressure against time and four characteristics of the pressure profile were important: initial delay time, maximum pressure, time to maximum pressure, and time from 10 to 90% of pressure increase. After some debate, the team decided to include three factors in the experiment: (a) the type of main charge, (b) the type of booster charge, and (c) the type of initiator. Four different main charges were used, along with two booster charges and three initiators. Thus the experiment is a $2 \times 3 \times 4$ factorial, with 24 possible ways to select the booster charge, initiator, and main charge that would go into a unit. The experiment included all 24 of these combinations, with two units of each type.

2 Factorial Experiments

Table 1 Issues for consideration in statistically designed experiments

	Issue	Description
1	Problem definition	Plain language description of the problem, how it was identified, who is in charge of the process/product under consideration
2	Response variable(s)	What will be measured and with what device; how the data will be collected
3	Control factors	What factors can be controlled by the operator, or by design decisions
4	Factor levels	What are the current factor levels and what are the reasonable alternatives
5	Noise factors	What factors cannot be controlled during regular operations, but could be controlled in an experiment, and what factors simply cannot be controlled, but could be observed and recorded
6	Experimental array	What main effects and interactions are of interest. What is the total practically possible number of runs. What is the detection power required against what alternatives. What are the experimental runs
7	Number of replications and order of experimentation	What is the practically possible number of replications per experimental run. In what order will the experimental runs be conducted
8	Data analysis	How will the data be analyzed. Who will write the final report
9	Budget and project control	Who is in charge of the project. What is the allocated budget. What is the timetable. What resources are needed

In general terms, a factorial experiment will involve k factors and the j th factor will be assigned l_j levels in the experiment. The resulting experiment is then said to be an $l_1 \times l_2 \times \cdots \times l_k$ factorial experiment. If several factors have the same number of levels, it is common to combine the common factors in a shorthand notation. For example, an experiment with three factors, each having two levels, and a fourth factor, having three levels, is a $2^3 \times 3$ factorial experiment.

A *full factorial experiment* includes all the possible combinations of the factor levels. The pyrotechnic experiment is an example of a full factorial experiment. As the number of factors increases, a full factorial experiment may be too large to run. Instead, specially selected subsets of the combinations, known as *fractional factorial designs*, may be used (see **Fractional Factorial Designs; Orthogonal Arrays; Plackett–Burman Designs**; Plackett and Burman [18]).

In the pyrotechnic experiment, all the factors are *nominal factors*; that is, they define options that cannot be placed on any ordered scale. Often experiments include factors like temperature, concentration, speed, and so on, which are *numerical* or *quantitative factors*. These factors have levels on a clear mathematical scale. An intermediate case is an *ordinal factor*, whose levels can be ordered from smallest to largest, but are not associated with a mathematical scale. As an example, an experiment on shock absorbers in a car might include “test track” as one

of the factors, with three different test tracks being used. It is not difficult to set up test tracks so that one is clearly the easiest, one the hardest and one is intermediate, but there is no obvious way to quantify how much harder one is than another.

Special considerations may go into choosing the levels of quantitative factors, as we may typically expect that responses will vary smoothly as their levels are changed. Thus, for example, the number and spacing of levels for these factors will affect whether we can represent their relationships to the responses as linear, or if additional terms are needed to capture nonlinear relationships. Selecting levels close together will usually allow a linear fit to work well, but if they are too close, the effect of the factor may not come through at all. Moving the levels further apart will increase the chances of seeing an effect, but might require a nonlinear relationship, which will require adding some intermediate levels for the factor.

Two-Level Factorial Designs

A common experimental design is one with all input factors set at two levels each. These levels are called *high* and *low* or “+1” and “−1”, respectively. A design with all possible high/low combinations of all the input factors is called a *two levels full factorial design*.

If there are k factors, each at two levels, a full factorial design has 2^k runs.

Consider the two-level, full factorial design for three factors, namely the 2^3 design. This implies eight runs (not counting replications or **center point** runs). Graphically, we can represent the 2^3 design by the cube shown in Figure 1. The arrows show the direction of increase of the factors. The numbers “1” through “8” at the corners of the design box reference the “standard order” of runs (see Table 2).

The left-most column of Table 2, numbers 1 through 8, specifies a (nonrandomized) run order called the *standard order*. These numbers are also shown in Figure 1. For example, run 1 is made at the “low” setting of all three factors.

We can readily generalize the 2^3 standard order matrix to a two-level full factorial with k factors. The first (X_1) column starts with -1 and alternates in sign for all 2^k runs. The second (X_2) column starts with -1 repeated twice, then alternates with two in

a row of the opposite sign until all 2^k places are filled. The third (X_3) column starts with -1 repeated four times, then four repeats of $+1$'s, and so on. In general, the i th column (X_i) starts with 2^{i-1} repeats of -1 followed by 2^{i-1} repeats of $+1$. Note that if we have k factors, each at two levels, there will be 2^k different combinations of the levels. In the case $k = 3$, the number of experimental runs is $2^3 = 8$.

Running trials corresponding to all possible factor combinations means that we can estimate all the main and interaction effects. In a 2^3 full factorial design there are three main effects, three two-factor interactions, and a three-factor interaction, all of which appear in the full model as follows:

$$\begin{aligned} Y = & \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 \\ & + \beta_{12} X_1 * X_2 + \beta_{23} * X_2 * X_3 + \beta_{13} X_1 * X_3 \\ & + \beta_{123} X_1 * X_2 * X_3 \end{aligned} \quad (1)$$

From the factorial experiment we calculate the main effects. In general, the effect of factor i is computed as:

$$E_i = \frac{\sum_{h: x_i = (+)} y_h - \sum_{h: x_i = (-)} y_h}{N/2},$$

where $h = 1, \dots, N, i = 1, \dots, k$ (2)

with N being the total number of trials (number of experimental runs $\times$ number of replicates) and $x_i = (+)$, $x_i = (-)$ corresponding to all values of the response variable at the high and low levels, respectively.

A full factorial design allows us to estimate all eight “ β ”-coefficients.

An effect corresponds to an increase of a factor by two units, from -1 to $+1$. The regression coefficient represents an increase in the factor X by one unit, so we have that for the main effects and the regression intercept:

$$\beta_i = \frac{E_i}{2} \text{ and } \beta_0 = \frac{\sum_{h=1}^N y_h}{N} \quad (3)$$

Formulas for computing interactions under a variety of experimental designs and details on other designs are available in **Foldover Designs; Fractional Factorial Designs; Block Designs, Incomplete; Interactions; Latin Hypercube Designs;**

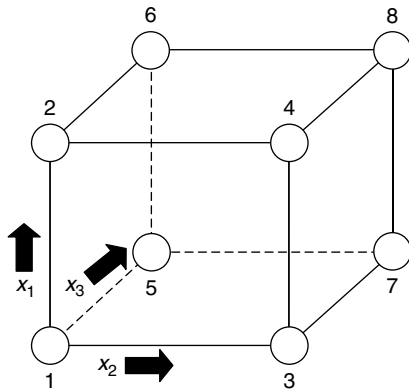


Figure 1 A cube plot of a full factorial 2^3 design

Table 2 A full factorial 2^3 design in standard order

A 2^3 two-level, full factorial design table showing runs in “standard order”

Run	X_1	X_2	X_3
1	-1	-1	-1
2	1	-1	-1
3	-1	1	-1
4	1	1	-1
5	-1	-1	1
6	1	-1	1
7	-1	1	1
8	1	1	1

4 Factorial Experiments

Table 3 The hybrid full factorial experiment

Experimental run	Factors			Responses				
	Paste	Dielectric	Oven	Y_1	Y_2	Y_3	Y_4	Y_5
1	−1	−1	−1	119.3	119.4	113.2	129.4	113.6
2	1	−1	−1	39.4	35.7	29.8	46.3	68.7
3	−1	1	−1	234.8	226	234.7	235.8	228.5
4	1	1	−1	127.9	135.8	129.7	140.0	141.5
5	−1	−1	1	2.0	10.1	11.7	7.5	23.7
6	1	−1	1	35.1	24.6	35.4	29.9	44.8
7	−1	1	1	41.1	36.3	47.9	43.9	55.6
8	1	1	1	85.07	72.8	86.4	83.1	72.4

Latin Squares and Related Experimental Designs; Lenth's Method for the Analysis of Unreplicated Experiments; Main Effects; Main Effect Designs; Optimal Design; Orthogonal Arrays; Plackett–Burman Designs; Screening Designs; Split-Plot Designs; Supersaturated Designs; Uniform Experimental Designs; Box–Behnken Designs; Central Composite Designs; Mixture Experiments; Response Surface Methodology; Robust Design.

Factorial Designs and Industrial Experiments

We illustrate factorial experiments by a simple full factorial experiment with three factors at two levels each, using an example presented by Kenett and Steinberg [19]. An organization developing, manufacturing and deploying telecommunication equipment faced a major crisis with functional defects in smart phones being reported by sales and marketing. Such reports had a negative effect on sales, and created a potential product recall with dramatic economic consequences. It became critical to quickly focus on the cause of the functional defects so that target actions could be launched in order to contain the growing problem. Management nominated an investigation team consisting of engineers from development and production, quality experts and a statistician. The team studied the problem reports and ascertained that the defects were caused by a reduction in interlayer resistance of a specific hybrid component on the smart phones. After intuitive screening of various factors that could affect the hybrid's performance, it was decided to focus on three nominal factors (a) dielectric type, (b) paste type, and (c) oven.

Production had switched to lower cost dielectric and paste types. Two ovens were used to produce the hybrid components and one of them was located in an area where recent construction projects created dust and other environmental disturbances. None of these changes were originally believed to have a negative impact on the hybrid components. The team decided to verify these assumptions.

A 2^3 full factorial experiment was designed to analyze the impact of these three factors on the resistance of the hybrid microcircuit. The factors' levels were labeled “−1” and “+1” and each experimental run was conducted with five replicates consisting of hybrids manufactured in identical conditions. The response variable, Y , is interlayer resistance measured in ohms. The experiment lasted 1 month with measurements being repeated every week. Table 3 presents results derived at the end of the first week.

A graphical display of the data is shown in Figure 2 as a cube layout with the average response per experimental run indicated at the corners of the cube. The effect of the oven clearly shows up with oven “−1” producing consistently larger resistances than oven “+1”. An interaction between oven and paste is also apparent, with the effect of the paste depending on the oven used.

Main effect plots on the left of Figure 3 connect average responses over different factor level combinations. Again, we clearly see that oven “+1” produced lower resistance. Splitting the main effect of an individual factor, over the levels of another factor, produced interaction plots. The top right corner interaction plot in Figure 3 corresponds to the interaction between oven and paste. Using paste “−1” in oven “−1” gave high resistances; all the other combinations of these two factors were problematic.

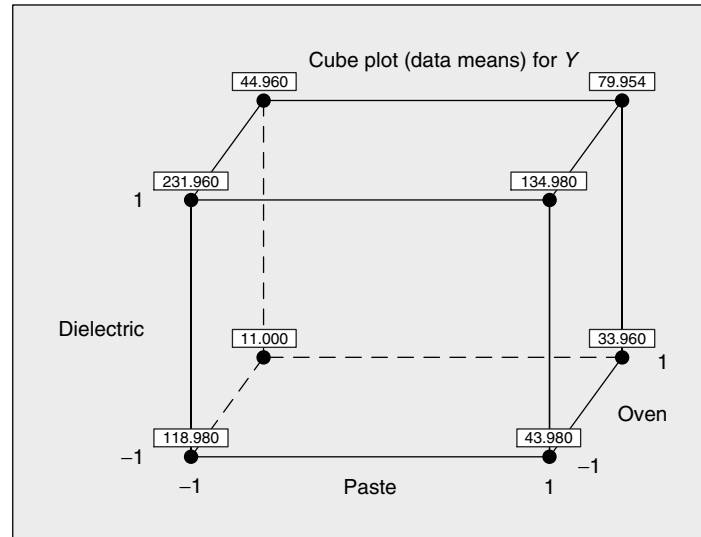


Figure 2 Cube plot of hybrid 2^3 full factorial experiment

The main goal of the experiment was to assess the main effects of the three factors and their interactions. From Table 4 we can see that, except for the paste by dielectric interaction, all effects and interactions are significantly nonzero.

In Figure 4, a normal and half-normal probability plot of the main effects and interactions are presented (see **Half-Normal Plot**, [20]). Again only the paste by dielectric interaction appears as nonsignificant. These conclusions allowed management to focus the recall on batches produced in oven “+1” with

dielectric “-1” and contain the problem. The figures and analysis in this example were generated with the MINITAB™ version 15 software.

Factorial Designs in Marketing Applications

Factorial designs are extensively used in marketing applications. Suppose you wanted to book an airline flight and you had a choice of spending \$400 or \$700 for a ticket. If this were the only consideration

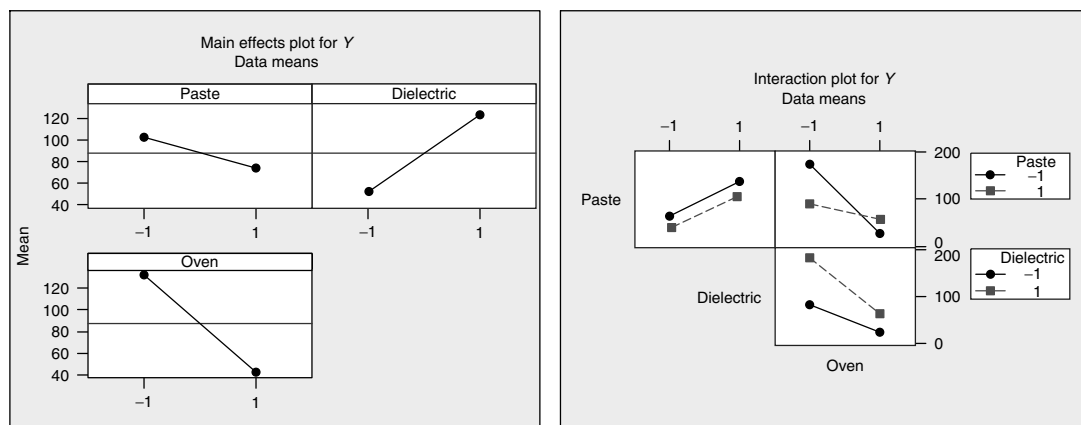


Figure 3 Main effect and interaction plots of hybrid 2^3 full factorial experiment

6 Factorial Experiments

Table 4 Effects and coefficients of hybrid 2^3 full factorial experiment

Estimated effects and coefficients for Y (coded units)					
Term	Effect	Coefficient	SE coefficient	T	P
Constant		87.47	1.306	67.00	0.000
Paste	−28.51	−14.25	1.306	−10.92	0.000
Dielectric	70.98	35.49	1.306	27.18	0.000
Oven	−90.01	−45.00	1.306	−34.47	0.000
Paste * dielectric	−2.49	−1.24	1.306	−0.95	0.348
Paste * oven	57.48	28.74	1.306	22.01	0.000
Dielectric * oven	−31.01	−15.50	1.306	−11.87	0.000
Paste * dielectric * oven	8.50	4.25	1.306	3.26	0.003

then the choice is clear: the lower priced ticket is preferable. What if the only consideration in booking a flight was sitting in a regular or extra-wide seat? If seat size was the only consideration then you would probably prefer an extra-wide seat. Finally, suppose you can take either a direct flight which takes 3 h or a flight that stops once and takes 5 h. Virtually everyone would prefer the direct flight. In a real purchase situation, however, consumers do not make choices based on a single attribute like price, comfort, or convenience. Consumers examine a range of features or attributes and then make judgments or trade-offs to determine their final purchase choice. Marketing specialists developed a specific application of factorial designs to these types of problems and labeled it *conjoint analysis*. Conjoint analysis examines feature

trade-offs to determine the combination of attributes that will be most satisfying to the consumer. In other words, by using conjoint analysis a company can determine the optimal features for their product or service. In addition, conjoint analysis will identify the best advertising message by identifying the features that are most important in product choice. Conjoint analysis presents choice alternatives between products/services defined by sets of attributes. This is illustrated by the following choice: would you prefer a flight with regular seats, that costs \$400 and takes 5 h, or a flight which costs \$700, has extra-wide seats and takes 3 h? Extending this, we see that if seat comfort, price and duration are the only relevant attributes, there are potentially eight flight choices (see Table 5).

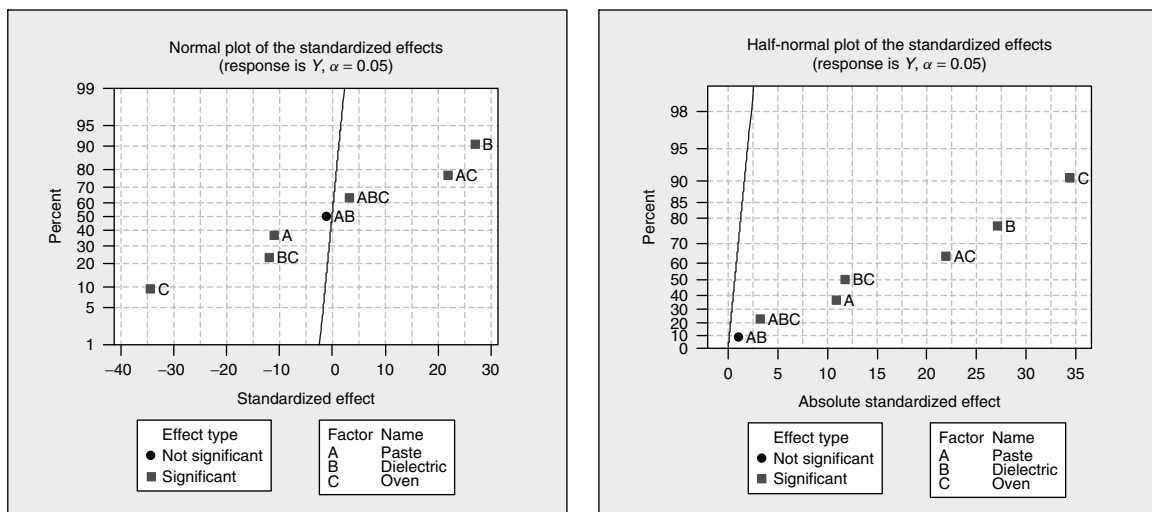


Figure 4 Normal and half-normal probability plot of main effects and interactions

Table 5 Factorial design used in conjoint analysis

Choice	Seat comfort	Price	Duration
1	Extra-wide	\$700	5 h
2	Extra-wide	\$700	3 h
3	Extra-wide	\$400	5 h
4	Extra-wide	\$400	3 h
5	Regular	\$700	5 h
6	Regular	\$700	3 h
7	Regular	\$400	5 h
8	Regular	\$400	3 h

Given the above alternatives, product 4 is very likely the most preferred choice while product 5 is probably the least preferred product. How the other choices are ranked is determined by what is important to that individual. Conjoint analysis can be used to determine the relative importance of each attribute, attribute level, and combinations of attributes. If the most preferable product is not feasible for some reason (perhaps the airline simply cannot provide extra-wide seats and a 3-h arrival time at a price of \$400) then the conjoint analysis will identify the *next* most preferred alternative. If you have other information on travelers, such as background demographics, you might be able to identify market segments for which distinct products may be appealing. For example, the business traveler and the vacation traveler may have very different preferences which could be met by distinct flight offerings.

In evaluating products, consumers will always make trade-offs. A traveler may like the comfort and arrival time of a particular flight, but reject purchase owing to the cost. In this case, cost has a high *utility* value. Utility can be defined as a number which represents the value that consumers place on an attribute. In other words, it represents the relative “worth” of the attribute. A low utility indicates less value; a high utility indicates more value. For more on conjoint analysis see [21–24]. Factorial designs have been widely implemented in the social sciences and many other disciplines (see Winer [7]).

Factorial Designs and Statistical Education

A major challenge in teaching statistics is to convey how statistics is used in solving real problems. Introductory courses often fail to provide students with a sense of how to apply statistical ideas. Students,

on the other hand, often complain that introductory statistics courses are a dry encounter with pages of formulas. Factorial designs have been successfully incorporated in such courses to overcome these difficulties [25–27].

Factorial Designs and Clinical Trials

Factorial designs are also used in **clinical trials**. Quoting the US Food and Drug Administration (FDA) guidance to industry on statistical principles for clinical trials: “In a factorial design, two or more treatments are evaluated simultaneously through the use of varying combinations of the treatments. The simplest example is the 2×2 factorial design in which subjects are randomly allocated to one of the four possible combinations of two treatments, A and B. These are: A alone; B alone; both A and B; neither A nor B. In many cases, this design is used for the specific purpose of examining the interaction of A and B. The statistical test of interaction may lack **power** to detect an interaction if the **sample size** was calculated based on the test for main effects. This consideration is important when this design is used for examining the joint effects of A and B, in particular, if the treatments are likely to be used together. Another important use of the factorial design is to establish the dose-response characteristics of the simultaneous use of treatments C and D, especially when the efficacy of each monotherapy has been established at some dose in prior trials. A number, m , of doses of C is selected, usually including a zero dose (placebo), and a similar number, n , of doses of D. The full design then consists of $m \times n$ treatment groups, each receiving a different combination of doses of C and D. The resulting estimate of the response surface may then be used to help identify an appropriate combination of doses of C and D for clinical use” [28].

In Barnett *et al.* [29], a clinical trial with a factorial design is presented. Five hundred and eighty-five patients with threatened stroke were followed in a randomized clinical trial for an average of 26 months to determine whether aspirin or sulfinpyrazone, singly or in combination, influences the subsequent occurrence of continuing transient ischemic attacks, stroke or death. Eighty-five subjects suffered a stroke, and 42 died. Aspirin reduced the risk of continuing ischemic attacks, stroke or death by 19% and

also reduced risk for the “harder”, more important events of stroke or death by 31%, but this effect was sex-dependent: among men, the risk reduction for stroke or death was 48%, whereas no significant trend was observed among women. For sulfinpyrazone, no risk reduction of ischemic attacks was observed, and the 10% risk reduction of stroke or death was not statistically significant. No interaction (i.e., synergism or antagonism) was observed between the two drugs. The authors conclude the analysis of this full factorial experiment with a statement that aspirin is an efficacious drug for men with threatened stroke [29].

Computer Intensive Methods and Factorial Designs

In analyzing data from factorial experiments an underlying model is fitted to the data. Two-level factorial experiments are used to estimate parameters of a linear model which, in turn, depends on the estimability properties of the experimental design. When a model is misspecified, an error of the third kind is said to have occurred. Bootstrapping can be used to flag errors of the third kind or, alternatively, validate a specific model. Use of an inadequate model will often lead to an overestimate of the residual variance and to inflated standard errors for the model parameters. Comparison of **bootstrap** standard errors (SEs) to those from a regression analysis is therefore a valuable diagnostic. If the bootstrap standard errors are clearly smaller than those from fitting a regression model to the experimental data, this is a likely sign that the model is inadequate [30].

In recent years, computer simulators have become an increasingly popular platform for running experiments in the product and/or process development process [31]. Some of the issues involved in computer experiments (*see Computer Experiments*) are common to physical experiments such as the need to study many factors, to make effective local approximations to input–output relationships and to exploit sequential methods. New issues have also arisen: the ability to use a very large number of factor settings, the lack of “experimental error”, the need to combine computer and lab or field data, to name a few. These new issues have stimulated the development of new statistical methods, which are often given the acronym DACE, for the design and analysis of computer experiments, the

title of a milestone publication in this area [32]. Computer simulations provide a platform to running experiments so as to simultaneously optimize and robustify products and processes. The experiments conducted in this context are extensions of the basic factorial experiments. These extensions include space filling experiments such as Latin hypercube (*see Latin Hypercube Designs*), uniform (*see Uniform Experimental Designs*) and lattice designs. For an implementation approach of these techniques, labeled the *stochastic emulator strategy*, see [33].

Computer simulations are used nowadays to study a wide range of phenomena, from automobile performance in crash tests to the impact of chemical mixtures on process yields and integrated circuit design on product performance. Actual physical experimentation in the laboratory or test bed can today be replaced with a “virtual” experiment in which a computer runs a program that simulates the behavior of the true process. Experiments can be run on a simulator at a much faster rate than in the laboratory, allowing companies to dramatically reduce the time, cost and effort needed to conduct experiments. Moreover, process settings that might have been thought risky or even potentially dangerous can easily be tested. When the physical testing is destructive, computer simulations are obviously much more economical than actual experiments.

References

- [1] Fisher, R.A. (1935). *The Design of Experiments*, 1st Edition, Oliver & Boyd, London.
- [2] Box, J. (1978). *R.A. Fisher; The Life of a Scientist*, John Wiley & Sons, New York.
- [3] Fisher, R. (1926). The arrangement of field experiments, *Journal of the Ministry of Agriculture of Great Britain* **33**, 503–513.
- [4] Box, G.E.P., Hunter, W.G. & Hunter, J.S. (1978). *Statistics for Experimenters. An Introduction to Design, Data Analysis, and Model Building*, John Wiley & Sons, New York.
- [5] Rao, C.R. (1947). Factorial experiments derivable from combinatorial arrangements of arrays, *Journal of the Royal Statistical Society Supplement* **9**, 128–139.
- [6] Cox, D.R. (1958). *Planning of Experiments*, John Wiley & Sons, New York.
- [7] Winer, B.J. (1962). *Statistical Principles in Experimental Design*, McGraw-Hill, New York.
- [8] John, P.W.M. (1971). *Statistical Design and Analysis of Experiments*, Macmillan, New York.

- [9] Steinberg, D.M. & Hunter, W. (1984). Experimental design: review and comment, (with discussion) *Technometrics* **26**, 71–130.
- [10] Hamada, M. & Balakrishnan, N. (1998). Analyzing unreplicated factorial experiments: a review with some new proposals (with discussion), *Statistica Sinica* **8**, 1–41.
- [11] Kenett, R.S., Zacks, S., (1998). *Modern Industrial Statistics: Design and Control of Quality and Reliability*, 2nd edition, Duxbury Press, San Francisco. 2003, Chinese edition 2004.
- [12] Wu, C.F.J. & Hamada, M. (2000). *Experiments: Planning, Analysis, and Parameter Design Optimization*, John Wiley & Sons, New York.
- [13] Brenneman, W.A. & Nair, V.N. (2001). Methods for identifying dispersion effects in unreplicated factorial experiments: a critical analysis and proposed strategies, *Technometrics* **43**(4), 388–405.
- [14] Montgomery, D.C. (2004). *Design and Analysis of Experiments*, 5th edition, John Wiley & Sons, New York.
- [15] Mukerjee, R. & Wu, C.F.J. (2006). *A Modern Theory of Factorial Design*, Springer, New York.
- [16] Steinberg, D.M. (1988). Factorial experiments with time trends, *Technometrics* **30**, 259–269.
- [17] Milman, B., Sirota, I. & Steinberg, D.M. (1995). Improving the safety of a pyrotechnic igniter through a controlled experiment, *Propellants, Explosives, Pyrotechnics* **20**, 294–299.
- [18] Plackett, R.L. & Burman, J.P. (1946). The design of optimum multifactorial experiments, *Biometrika* **33**, 305–325.
- [19] Kenett, R.S. & Steinberg, D.M. (2001). On the application of bootstrapping in analyzing designed experiments, Proceedings of the First Annual Conference of the European Network of Business and Industrial Statistics (ENBIS), Oslo.
- [20] Daniel, C. (1959). Use of half-normal plots in interpreting factorial experiments, *Technometrics* **1**, 311–341.
- [21] Green, P.E. & Rao, V.R. (1971). Conjoint measurement for quantifying judgmental data, *Journal of Marketing Research* **8**, 355–363.
- [22] Green, P.E. & Srinivasan, V. (1978). Conjoint analysis in consumer research: issue and outlook, *Journal of Consumer Research* **5**, 103–123.
- [23] Gustafsson, A., Hermann, A. & Huber, F. (2001). *Conjoint Measurement: Methods and Applications*, Springer, Berlin.
- [24] Rossi, P.E., Allenby, G.M. & McCulloch, R. (2005). *Bayesian Statistics and Marketing*, John Wiley & Sons, West Sussex.
- [25] Hunter, W.G. (1975). “101 Ways to design an experiment, or some ideas about teaching design of experiments, Technical Report number 413, University of Wisconsin - Madison, <http://curiouscat.com/bill/101doe.cfm>.
- [26] Hunter, W.G. (1977). Some ideas about teaching design of experiments, with 2^5 examples of experiments conducted by students, *The American Statistician* **31**, 12–17.
- [27] Kenett, R.S. & Steinberg, D.M. (1987). Some experiences teaching factorial designs in introductory statistics courses, *Journal of Applied Statistics* **14**, 219–228.
- [28] International Conference on Harmonization (ICH), (1998). Guidance for Industry E9 Statistical Principles for Clinical Trials, www.fda.gov/Cder/guidance/ICH/E9-fnl.PDF.
- [29] Barnett, H.J.M., Gent, M., Sackett, D.L., (1978). A randomized trial of aspirin and sulfinpyrazone in threatened stroke: the Canadian cooperative study group, *N Engl J Med* **299**, 53–59.
- [30] Kenett, R.S., Rahav, E. & Steinberg, D.M. (2006). Bootstrap analysis of designed experiments, *Quality and Reliability Engineering International* **22**, 659–667.
- [31] Kenett, R.S. & Steinberg, D.M. (2006). New frontiers in design of experiments, *Quality Progress* **39**(8), 61–65.
- [32] Sacks, J., Welch, W.J., Mitchell, T.J. & Wynn, H.P. (1989). Design and analysis of computer experiments, *Statistical Science* **4**, 409–435.
- [33] Bates, R.A., Kenett, R.S., Steinberg, D.M. & Wynn, H.P. (2006). Achieving robust design from computer simulations, *Quality Technology and Quantitative Management* **3**(2), 161–177.

Related Article

Assessment of Experimental Designs.

RON S. KENETT AND DAVID M. STEINBERG

Foldover Designs

Introduction

Fractional factorial designs are frequently used in the design of experiment area (see **Fractional Factorial Designs**). One consequence of using a fractional factorial design is the **aliasing** of the factorial effects. Consider a 2^{7-4} example given in Montgomery [1], in which a human performance analyst studied eye focus time. Seven two-level factors chosen for study were acuity or sharpness of vision (factor 1), distance from target to eye (factor 2), target shape (factor 3), illumination level (factor 4), target size (factor 5), target density (factor 6), and subject (factor 7). Factors 1–3 constitute a full factorial design with eight runs and are called *basic factors*. The other four factors are defined by *generators* (see **Fractional Factorial Designs**): $4 = 12$, $5 = 13$, $6 = 23$, and $7 = 123$, and are called *generated factors*. The *defining relations* (see **Fractional Factorial Designs**) are given by:

$$\begin{aligned} I &= 124 = 135 = 236 = 1237 = 2345 = 1346 \\ &= 347 = 1256 = 257 = 167 = 456 = 1457 \\ &= 2467 = 3567 = 1234567 \end{aligned} \quad (1)$$

The analysis of the experiment shows that main effects of factors 1, 2, and 4 are significant. However, since each **main effect** is fully aliased with some interactions (e.g., from equation (1), factor $1 = 24 = 35 = 67 = \dots$), it was not clear from the original design whether main effect 1 or two-factor interaction 24 (or 35) was significant. A standard follow-up strategy to break the aliasing of effects is to add a *foldover design* (or simply, *foldover*) by reversing the signs of one or more columns of the initial design. (See, e.g., Box *et al.* [2]; Kutner *et al.* [3]; Montgomery [1]; and Wu and Hamada [4]). Following Li and Lin [5], we denote a foldover plan γ by the collection of columns whose signs are reversed in the foldover design. Then in the eye focus time example, there are $2^7 = 128$ foldover plans to generate a foldover design.

We refer to the *combined design* as the combination of the initial design and its foldover design. When a foldover is added to the initial design, then the resulting combined design will have a better design structure than the initial design. Different

foldover methods are discussed in the next section. Throughout the article we focus on foldover of two-level designs.

Foldover Methods

Foldover on All Columns

A classic foldover approach is to reverse the signs of all columns. This special type of foldover plan is called a *full-foldover* plan in Li and Lin [5]. The resulting full-foldover design for the eye focus example is given in the second part of Table 1. The full-foldover plan has been among the most popular foldover techniques. A main advantage of the full-foldover plan is that any odd-length words in the defining relation of the initial designs do not appear in the combined design. Consequently, if the initial design is of resolution III, the combined design using the full-foldover is of at least resolution IV because all words of length 3 of the initial design disappear in the combined design. (For definitions of *word length* and *resolution* of a fractional factorial design, see **Fractional Factorial Designs** and **Minimum Aberration**.)

To see why this happens and how the use of foldover can untangle aliased effects, consider the eye focus time example. After the initial design was performed, the main effect of factor 2 – the largest among seven main effects – can be calculated as follows:

$$\begin{aligned} l_2 &= \frac{1}{4}(-y_1 - y_2 + y_3 + y_4 - y_5 - y_6 + y_7 + y_8) \\ &= 38.375 \end{aligned} \quad (2)$$

Assuming negligible interactions between three or more factors, we obtain from equation (1):

$$2 = 14 = 36 = 57 \quad (3)$$

Thus,

$$l_2 = 38.375 \rightarrow 2 + 14 + 36 + 57 \quad (4)$$

where the notation means that this contrast is actually estimating the main effect of factor 2 plus the sum of the three two-factor interactions. In the full-foldover design given in the second part of Table 1, factors are denoted by -1 , -2 , -3 , -4 , -5 , -6 , and -7 . It can be shown that it has generators: $4 = -12$, $5 =$

2 Foldover Designs

Table 1 Initial and foldover designs for the eye focus time experiment

Run	Initial design							Response
	1	2	3	4 = 12	5 = 13	6 = 23	7 = 123	
1	–	–	–	+	+	+	–	85.5
2	+	–	–	–	–	+	+	75.1
3	–	+	–	–	+	–	+	93.2
4	+	+	–	+	–	–	–	145.4
5	–	–	+	+	–	–	+	83.7
6	+	–	+	–	+	–	–	77.6
7	–	+	+	–	–	+	–	95
8	+	+	+	+	+	+	+	141.8

Run	Foldover design (on all columns)							Response
	–1	–2	–3	–4	–5	–6	–7	
9	+	+	+	–	–	–	+	91.3
10	–	+	+	+	+	–	–	136.7
11	+	–	+	+	–	+	–	82.4
12	–	–	+	–	+	+	+	73.4
13	+	+	–	–	+	+	–	94.1
14	–	+	–	+	–	+	+	143.8
15	+	–	–	+	+	–	+	87.3
16	–	–	–	–	–	–	–	71.9

–13, 6 = –23, and 7 = 123. We estimate the main effect for factor 2 in the foldover design by

$$l'_2 = \frac{1}{4}(y_9 + y_{10} - y_{11} - y_{12} + y_{13} + y_{14} - y_{15} - y_{16})$$

$$= 37.725 \rightarrow 2 - 14 - 36 - 57 \quad (5)$$

Overall, the main effect for factor 2 can be estimated from the combined 16-run design:

$$\frac{1}{2}(l_2 + l'_2) = \frac{1}{2}(38.375 + 37.725)$$

$$= 38.05 \rightarrow 2 \quad (6)$$

Note that the main effect for factor 2 is “de-aliased” from interactions 14 + 36 + 57 by equations (4) and (5). In fact, all defining words of length 3 in equation (1) do not appear in the defining relations for the combined design. Consequently, all main effects are de-aliased from two-factor interactions.

Foldover on One Column

This type of foldover design was first discussed in [2]. If a foldover on one column is added, then neither this factor nor its associated interactions are aliased with other effects. For example, in the eye focus time

experiment, the largest effect obtained from the first set of eight runs was factor 2. Therefore, it might have been argued that it would be desirable to use a foldover to de-alias factor 2 and all the interactions of other variables with 2.

To see why this could have been achieved by adding a foldover design that switches the sign of column 2 only, we write down the defining relations for the foldover design:

$$I = -124 = 135 = -236 = -1237 = -2345$$

$$= 1346 = 347 = -1256 = -257 = 167 = 456$$

$$= 1457 = -2467 = 3567 = -1234567 \quad (7)$$

From equations (1) and (7), the defining relations for the combined design are as follows:

$$I = 135 = 1346 = 347 = 167$$

$$= 456 = 1457 = 3567 \quad (8)$$

Note that none of the defining words in equation (8) is associated with factor 2.

Optimal Foldover Designs

For many years the traditional foldover approaches have been focused on the foldover on all columns

or just one particular column. Obviously, a foldover can be conducted on *any* set of columns. There has been resurgent research on foldover designs in recent years, after the optimal foldover was considered by Li and Mee [6] and Li and Lin [5]. The former [6] provides a sufficient condition for the situation where a better foldover plan exists compared with the full-foldover plan. The other paper [5] obtains optimal foldover designs for commonly used two-level fractional factorial designs with 16 and 32 runs. Both papers demonstrate that foldover on only a subset of all columns can be a better choice in terms of commonly used criteria, such as the minimum **aberration** criterion. (For the definition of *aberration*, see **Fractional Factorial Designs**.)

In the next section we introduce recent work on optimal foldover designs.

Optimal Foldover

Optimal Foldovers on Two-Level Fractional Factorial Designs

Two-level fractional factorial designs are among the most widely used experiments. Early works on foldover had been focused on these designs. For example, it was well known and stated in many textbooks on design of experiments [1–4] that adding a full-foldover of a resolution III design will result in a combined design that is of resolution IV. However, more recent research [5] shows that the conventional full-foldover approach does not necessarily produce the best combined design in terms of the aberration criterion.

Example 1: Suppose a 16-run design d is defined by: $5 = 12$ and $6 = 134$. This is a resolution III design, and its *word length pattern* (see **Fractional Factorial Designs**) is given by: $W(d) = (0, 0, 1, 1, 1, 0)$. If the full-foldover plan is used, that is, $\gamma = \{1, 2, 3, 4, 5, 6\}$, then the combined design is a resolution IV design, whose word length pattern is $W = (0, 0, 0, 1, 0, 0)$.

We know from [5] that the optimal foldover plan for this initial design is $\gamma^* = \{5, 6\}$. That is, if the foldover design is obtained by switching the signs of columns 5 and 6 only, then the combined design has the best aberration criterion value. Its word length pattern is $W = (0, 0, 0, 0, 1, 0)$, so that it is a resolution V design.

Using a computer search method, Li and Lin [5] obtained and tabulated optimal foldover plans for all 16-run and selected 32-run fractional factorial designs. They also proved that every foldover plan is equivalent to a *core foldover plan*, which involves only the generated factors. For example, for the design defined by $5 = 12$, $6 = 134$, there are only four core foldover plans: $\{0\}$, $\{5\}$, $\{6\}$, $\{5, 6\}$. (Following [5], we use $\{0\}$ to refer to the case where no column is chosen for sign reversal and include it as a core foldover plan.) Then all $2^6 = 64$ foldover plans are equivalent to one of the four core foldover plans. It is easy to see that use of core foldover plans results in substantial computational savings.

Optimal Foldovers on Two-Level Nonregular Designs

Fractional factorial designs are constructed through defining relations (e.g., those given in equation (1)). These designs are called *regular designs*, for which a factorial effect is either orthogonal or fully aliased with any other effect. In contrast, in *nonregular designs* factors are not added by generators and, as a result, some effects are neither orthogonal nor fully aliased, a condition known as being partially aliased (see **Minimum Aberration**). Compared with regular designs, nonregular designs have more complex aliasing structures, but they also provide more flexibility in run size and can be used to entertain more distinct models. For example, all 12-run orthogonal designs are nonregular designs. For two-level designs with 16 runs, both regular and nonregular designs exist. For a given number of factors, there are usually more nonregular designs than regular designs. For instance, among the 55 designs with 16 runs and 7 factors, there are only 6 regular designs and 49 nonregular designs [7].

Several authors discussed properties of foldover of nonregular designs. Both Diamond [8] and Miller and Sitter [9] studied the foldover of 12-run Plackett–Burman designs (see **Plackett–Burman Designs**). For commonly used 12-run and 16-run nonregular designs, Li *et al.* [7] obtained optimal foldover plans, using the generalized aberration criterion proposed by Deng and Tang [10] (see **Minimum Aberration**). Using an indicator function approach, Li *et al.* [7] defined generalized word length, which does not have to be an integer. The *generalized word length pattern* (see **Minimum Aberration**), which is also called the extended word

length pattern in [7], and generalized resolution can then be obtained for a nonregular design.

Li *et al.* [7] proved that the full-foldover plans are optimal for all 12-run and 20-run orthogonal designs in terms of the generalized aberration criterion. Optimal foldover plans for 16-run designs are usually not full-foldover plans. The complete set of optimal foldover plans is given on the web page of an author of [7] (<http://www.csom.umn.edu/~wli>).

Example 2: Consider a 16-run nonregular design d with nine factors. Its generalized word length pattern is: $W(d) = (0, 16; 14, 0; 0, 32; 0, 0)$, having 0 words of length 3, 16 words of length $3\frac{1}{2}$, and so on. A word of length $3\frac{1}{2}$ corresponds to three (out of nine) factors that are partially aliased, with a “correlation” of $1/2$. For the strict definition of a word with fractional length, see [7].

The optimal foldover plan for this design is $\gamma^* = \{9\}$. The generalized word length pattern of the combined design is given by: $W = (0, 0; 14, 0; 0, 0; 0, 0)$. This is a resolution-4 design, compared with a resolution-3.5 initial design d .

Discussion

Foldover design is a useful follow-up strategy. The foldover designs can be easily constructed. Because the resulting combined design is still orthogonal, the analysis is also straightforward. It is especially useful if the initial design is of resolution III and the follow-up experiment is needed to de-alias some main effects and two-factor interactions. In fact, early research on foldover designs such as [2] appeared to be motivated by the fact that adding a full-foldover of a resolution III design produces a resolution IV combined design. Recent studies extend the idea to the foldover on one or two columns [11] and, more generally, to the optimal foldovers on any subset of columns [5–7].

The main weakness of the foldover design is the run-size economy. For a two-level design, a foldover design always requires the *same* run size as the initial design. This can be wasteful if the main purpose of the follow-up experiment is to de-alias a small number of aliased effects. Alternative methods include adding a small optimal design in terms of commonly used criteria such as the D criterion (see **Optimal Design**) or using a Bayesian approach to

design a follow-up experiment, described in Meyer *et al.* [12]. On a related issue, while the foldover method can also be applied to general multilevel factorial designs, the run-size requirement would be an even bigger constraint. Consequently, such foldover designs are less appealing, and most existing studies on foldovers have been focused on two-level designs.

We conclude the article with two remarks. First, the optimal foldover plans described in the last section were obtained using the aberration criterion, where the goal could be to de-alias as many main effects as possible from two-factor interactions or to de-alias as many two-factor interactions as possible from each other. When the foldover design is used for a different purpose (e.g., to de-alias effects involving a specific factor), other foldover plans such as foldover on a specific column may be more appropriate. Second, since foldover design is usually used as a follow-up experiment, it is justifiable to include a **blocking** variable with one level being assigned to the initial design and another level being given to the foldover design. Ye and Li [13] proved that adding such a blocking variable would not change the order of foldover plans in terms of the aberration criterion. Thus, the optimal foldover plans given in [5] and [7], which were obtained without considering the blocking variable, are still optimal when the blocking variable is considered.

References

- [1] Montgomery, D.C. (1991). *Design and Analysis of Experiments*, 3rd Edition, John Wiley & Sons.
- [2] Box, G.E.P., Hunter, W.G. & Hunter, J.S. (1978). *Statistics for Experimenters: An Introduction to Design, Data Analysis, and Model Building*, 1st Edition, John Wiley & Sons.
- [3] Kutner, M.H., Nachtsheim, J.N., Neter, J. & Li, W. (2005). *Applied Linear Statistical Models*, 5th Edition, McGraw-Hill/Irwin.
- [4] Wu, C.F.J. & Hamada, M. (2001). *Experiments: Planning, Analysis, and Parameter Design Optimization*, 1st Edition, Wiley-Interscience.
- [5] Li, W. & Lin, D.K.J. (2003). Optimal foldover plans for two-level factorial designs, *Technometrics* **45**, 142–149.
- [6] Li, H. & Mee, R.W. (2002). Better foldover fractions for resolution III 2^{k-p} designs, *Technometrics* **44**, 278–283.
- [7] Li, W., Lin, D.K.J. & Ye, K. (2003). Optimal foldover plans for non-regular designs, *Technometrics* **45**, 347–351.
- [8] Diamond, N.T. (1995). Some properties of a foldover design, *The Australian Journal of Statistics* **37**, 345–352.

-
- [9] Miller, A. & Sitter, R.R. (2001). Using the folded-over 12-run Plackett-Burman design to consider interactions, *Technometrics* **43**, 44–55.
- [10] Deng, L.Y. & Tang, B. (1999). Generalized resolution and minimum aberration criteria for Plackett-Burman and other non-regular factorial designs, *Statistica Sinica* **9**, 1071–1082.
- [11] Montgomery, D.C. & Runger, G.C. (1996). Foldovers of 2^{k-p} resolution IV experimental designs, *Journal of Quality Technology* **28**, 446–450.
- [12] Meyer, R.D., Steinberg, D.M. & Box, G. (1996). Follow-up designs to resolve confounding in multifactor experiments (with discussion), *Technometrics* **39**, 382–390.
- [13] Ye, K. & Li, W. (2003). Some properties of blocked and unblocked 2^{k-p} designs, *Statistica Sinica* **13**, 403–408.

WILLIAM LI

Fractional Factorial Designs

Introduction

In a designed experiment that is arranged in a factorial design (*see* **Factorial Experiments**), all possible combinations of the factors' levels are run at least once. For example, a full factorial design with five factors that are each run at three levels would require at least $3^5 = 243$ runs. These designs were originally created for agricultural field experiments, in which typically only one or two experiments could be run each year; each experimental unit (a plot of land) was relatively inexpensive; and the experimenter expected a large amount of unexplained variation. Such experiments are conducive to large factorial designs. However, experiments in research, development, and production in industry are typically characterized by the ability to run many experiments each year; each experimental unit may be relatively expensive; and there is often much less unexplained variation. In such situations, it becomes useful to study a large number of factors with a relatively small number of runs, and this is the purpose of fractional factorial designs. In the situation considered above, only a subset, or fraction, of the 243 combinations would be run in a fractional factorial design.

Suppose we have a set of factors, the number of levels at which each factor will be run, and a run limit of at most n , where n is less than the number of all possible combinations. Then two key questions arise. First, which subset of at most n combinations is "best" to run? Second, what is a good strategy for analyzing the data that result from the experiment?

In industry, it is not uncommon to study seven or more factors in an experiment. To keep the number of runs at a reasonable size, the most common designs are those in which each factor is studied only at two levels, and this will be our primary interest.

Basic Idea of Fractional Factorial Design Construction

Suppose we want to study the effects of four factors, labeled A , B , C , and D , each at two levels, but we can only perform $n = 8$ runs instead of the

$2^4 = 16$ required for a full factorial design. To see the basic idea of constructing a good fractional factorial design in this case, we first consider only three factors, A , B , and C , to form a full factorial (2^3) design, and then add the fourth factor D , to form a fractional factorial design. See the columns labeled **A**, **B**, and **C** in Table 1, where we follow the common practice of listing the factor levels using ± 1 coding. These columns will sometimes be regarded as n -dimensional vectors. (Randomization (*see* **Randomization in Experimental Designs**) of the run order would normally be done as well, but this is not relevant for our discussion.)

For this design, the following is well known:

1. We may estimate the following effects from the data generated by this experiment: the main effects (*see* **Main Effects**) A , B , and C ; the two-factor interactions (*see* **Interactions**) AB , AC , and BC ; and the three-factor interaction ABC .
2. A **main effect**, say of A , may be calculated by the formula $2(\mathbf{A} \cdot \mathbf{Y})/n$, where $\mathbf{Y}$ is the column (vector) of the response data and the " $\cdot$ " denotes the inner-product operator.
3. An interaction effect, say of AB , may be calculated by first defining a new vector $\mathbf{AB}$ that is formed by element-by-element multiplication of $\mathbf{A}$ and $\mathbf{B}$, as shown in Table 1, and then applying the same rule as for a main effect; that is, the AB effect is $2(\mathbf{AB} \cdot \mathbf{Y})/n$.
4. Both theoretically and empirically, the ABC interaction is least likely to have any effect on the results. (This is an example of the *principle of effect hierarchy* – main effects are more likely to be active than two-factor interactions, which in turn are more likely to be active than three-factor interactions, etc.)

Now suppose we want to add the fourth factor, D , but are still limited to $n = 8$ runs. The fundamental strategy in creating this fractional factorial design is to run the levels of D at the same "settings" of the **ABC** column, as shown in Table 1 in the column labeled $\mathbf{D} \equiv \mathbf{ABC}$. When we run the experiment and analyze the data, the effects of D and of ABC cannot be distinguished. We say that they are *aliased* (*see* **Aliasing in Fractional Designs**) or *confounded* with each other. In fact, the effect calculated from this column turns out to be the sum of the D and the

2 Fractional Factorial Designs

Table 1 A 2^3 design augmented to a 2^{4-1} design

A	B	C	AB	AC	BC	D $\equiv$ ABC
-1	-1	-1	1	1	1	-1
1	-1	-1	-1	-1	1	1
-1	1	-1	-1	1	-1	1
1	1	-1	1	-1	-1	-1
-1	-1	1	1	-1	-1	1
1	-1	1	-1	1	-1	-1
-1	1	1	-1	-1	1	-1
1	1	1	1	1	1	1

ABC effects. This is sometimes written as $D + ABC$, where these symbols are used here to refer to the effects of the factors. However, it is very likely that the effect of ABC is either zero or very small, so it is reasonable to claim that any large effect here can be attributed to D .

In summary, one way to consider fractional factorial design construction is (a) start with enough factors (*basic factors*) to run a full factorial design and then (b) add additional factors by confounding them with interactions of the basic factors. An equivalent method would be to begin with a full factorial 2^4 design with factors $A - D$, but only keep those combinations (rows) in which $D \equiv ABC$.

Suppose we want to add a fifth two-level factor, say E , to this eight-run design. From Table 1, we must confound E with either AB , AC , or BC , and here we choose $E \equiv AB$. We are now confounding a main effect with a two-factor interaction instead of a three-factor interaction. Because two-factor interactions are more likely to have an effect than a three-factor interaction, this five-factor design is not as “good” as the four-factor design – if the effect associated with this column turns out to be large, it may often not be clear if the effect is due to E , AB , or a combination of these. We quantify these ideas in the next section.

These types of designs are most commonly known as 2^{k-p} designs. Here, k is the number of factors, the calculated value of 2^{k-p} is the number of treatment combinations, and p , known as the *fractionation element*, determines the fraction 2^{-p} of the number of runs in the design compared to a full factorial. The five-factor example would then be called a 2^{5-2} design – there are five factors, $2^{5-2} = 8$ treatment combinations, and $2^{-2} = \frac{1}{4}$ of the total number of $2^5 = 32$ from the full factorial design were run.

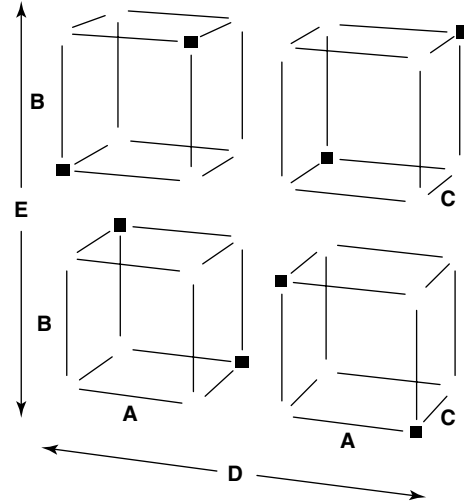


Figure 1 The 2^{5-2} design with $D \equiv ABC$ and $E \equiv AB$

The fractional nature of the 2^{5-2} design we constructed may be seen in Figure 1. That the design is a one-fourth fraction can be seen directly. The figure also shows that E is confounded with AB . For example, at the low level of E , only two of the four combinations of A and B are run; at the high level, only the other two combinations are run. Contrast this with, say, the behavior of A and C with E – at both the low level of E and the high level of E , all four combinations of A and C are run. This means that A , C , and E form a 2^3 design, a projectivity feature (see **Projectivity in Experimental Designs**) of the design. Similarly, if we collapse the design over D and E by combining all four of the “ $A/B/C$ ” cubes, we see that these three factors form a 2^3 design. This should not be surprising, because we started with this design before the addition of factors D and E .

Generators, Defining Contrasts, and Resolution

In the five-factor example, we set $D \equiv ABC$ and $E \equiv AB$. It is more common to write this as $D = ABC$ and $E = AB$, where we treat the symbols here as referring to the columns of ± 1 's. If we perform element-wise multiplication of these columns, we find that $I = ABCD = ABE$, where I , defined to be a column of $+1$'s, is the *identity element* under this column multiplication. We call the two terms

$ABCD$ and ABE the *generators* of the design. It follows that the product of these last two terms is also I , that is $I = ABCD \times ABE = ABCDABE = AABBCDE = IICDE = CDE$. The full set of such terms, $I = ABCD = ABE = CDE$, is called the *defining relation* or the *defining contrast of the design*.

The defining relation is useful because it can be used to find all the confounding in the design. For example, by multiplying all terms in the defining relation by E , we find that $E = ABCDE = AB = CD$. This shows that E is confounded not only with the two-factor AB interaction, which we already knew, but also with the CD interaction. (It is also confounded with the $ABCDE$ interaction, but the effect of this interaction can be safely assumed to be zero.) All the confounding in the design may easily be determined by finding such confounding for I and the basic factors and their interactions. Results for the 2^{5-2} example are shown in Table 2. So, if an analysis of the results found out that A was an active effect, it is technically only true that $A + BCD + BE + ACDE$ is active. In most cases, it will be safe to assume that this can be shortened to $A + BE$, or possibly even A , as we discuss below.

In general, a 2^{k-p} design contains p generators. We call each term in the defining relation a *word*, and in general a 2^{k-p} design contains 2^p words (including I) in the defining relation.

The *length* of a word is defined to be the number of letters it contains, so $ABCE$ would have length 4. The design resolution (see **Factorial Designs, Resolution of**) is defined to be the length of the shortest word (excluding I) in the defining relation. For example, the resolution of the 2^{5-2} design is 3. This is often written as III and, for the example, the design is sometimes denoted by 2^{5-2}_{III} . The practical implication of this is that in a resolution III design some main effects are confounded with one

or more two-factor interactions; in a resolution IV design some two-factor interactions are confounded with each other, but not with main effects; and neither two-factor interactions nor main effects are confounded with each other in a resolution V design. A finer distinction, known as *aberration* (see **Minimum Aberration**) is usually made among designs of equal resolution.

The 2^{k-p} designs are examples of orthogonal arrays (see **Orthogonal Arrays**). They are also examples of *geometric designs*, which are designs in which each effect is either totally confounded with or orthogonal to every other effect. For example, in the 2^{4-1} design, AB is totally confounded with CD , but is orthogonal to all other effects.

Design Construction

Consider a fixed value of the number of factors k and fractionation element p (or, equivalently, run size n). Principles such as resolution, aberration, or number of clear effects [1] can be used to compare 2^{k-p} designs to find which are optimal. However, the number of such potential designs grows rapidly with k except for special values of p , such as $p = 1$. Because theoretical results were not available aside from these special cases, optimum designs had traditionally been derived using algorithms, for example, Franklin and Bailey [2], and these in turn have been used to create tables of optimal designs. Examples may be found in Box *et al.* [3], Wu and Hamada [4], and Lorenzen and Anderson [5]. These references include designs for k as large as 5–16 and n up to 64. These or similar designs are also available in many statistical software packages.

More theoretical results have since been obtained that find optimum designs for certain (k, p) combinations without the use of algorithms. See, for example, Chen and Wu [6] and Chen [7], who found minimum aberration (MA) designs for $p = 3, 4, 5$, and Tang and Wu [8], Chen and Hedayat [9], and Butler [10], who found MA designs for $n = 2^{k-p}$ runs and $5n/16 \leq k < n$ factors.

In practice, the following points may serve as useful guides to the resolution of designs:

1. If $k \leq n/2$, a fractional factorial design of resolution IV or higher can be constructed. If $k > n/2$, the design will be resolution III.

Table 2 Confounding in the 2^{5-2} example

$I = ABCD = ABE = CDE$
$A = BCD = BE = ACDE$
$B = ACD = AE = BCDE$
$C = ABD = ABCE = DE$
$AB = CD = E = ABCDE$
$AC = BD = BCE = ADE$
$BC = AD = ACE = BDE$
$ABC = D = CE = ABDE$

4 Fractional Factorial Designs

2. Any design of resolution III may be folded over (see **Foldover Designs**) to form a design of resolution IV.
3. A 2^{k-1} design of resolution k always exists, and can be constructed by aliasing the k th factor with the interaction involving all of the other $k - 1$ factors.
4. For any design of resolution R , any subset of $R - 1$ factors or fewer form a full factorial design, possibly replicated. For example, the 2^{5-1} design with $I = ABCDE$ has resolution V, so any subset of one to four factors forms a full factorial design. This is another example of the projectivity of these designs.

Finney [11] was the first to develop a general theory of fractional factorial design construction.

Uses, Advantages, and Disadvantages

Two-level fractional factorial designs have been used extensively and productively in industry for over half a century. One use has been as the first experiment in sequential experimentation (see **Sequential Experimentation**). Here, the experimenter has many factors of potential importance but likely only a few that will actually turn out to be important. For this reason, these designs are also known as *screening designs* (see **Screening Designs**). Such designs may need to accommodate a large number of factors, so these must often be run as resolution III designs. Because follow-up designs are expected and these follow-up designs can be used to resolve uncertainties, an initial design of resolution III is acceptable. See [3], for example, for the uses of such designs.

These designs can also be effective in one-shot experiments – see Ledolter [12], for example, for a detailed case study. These designs also form the basis for experimentation both in robust designs (see **Robust Design**) and in tolerance designs (see **Tolerance Design**).

Most areas in which process complexity precludes theoretical prediction, in which more than a few factors need to be studied, and in which data may be collected are amenable to these designs. These features make these designs appealing across a broad range of fields. For example, Snelick *et al.* [13] study software code that “was described as a ‘complex benchmark’ that would be almost impossible to tune”. Using a 2_{IV}^{31-25} design, the authors found

important bottlenecks in the code and subsequently analyzed and reduced them. They point out the ability of this method to “eliminate (screen) uninteresting factors in linear time”. As another example, these designs are often used in *conjoint analysis*. This is a technique used in marketing to examine different combinations of attributes, often during the design of a new product, by presenting them to subjects. See, for example, Louviere [14].

Two advantages of these designs have already been noted: a *large number of factors* may be studied in a *relatively small number of runs*. In addition, because these designs are orthogonal, it can be shown that each factor’s effects are estimated with maximum precision, that is, the same precision that would have been obtained if only that factor was under study (e.g., only factor A in Table 1). These designs are also balanced, just as the full factorial designs from which they were derived – for example, factor A is run at both of its settings an equal number of times. Well-chosen designs usually also have the property that the maximum number of potentially active effects are unconfounded with each other.

Suppose that each factor has levels that may be varied in a continuum between the low and high settings, and that the response (aside from random error) is expected to be a smooth function of the factors’ coded levels. Then, locally, the function may be approximated by a first-order Taylor series expansion. If so, the main effects of the design are estimates of the coefficients of the first-order partial derivative terms (after division by 2, to adjust for the coded range from -1 to $+1$), based on the secant method of approximating derivatives, so the equation developed from these terms provides this first-order model. In addition, testing for active two-factor interactions can be used as a test of lack of fit (see **Lack of Fit**) to see if the first-order model is reasonable.

Other advantages of these designs stem from their close association with factorial designs. For example, large regions of design space are explored and **outliers** have a relatively small effect on the results and are also relatively easy to detect. To ensure that all of these advantages exist, these designs should be run in conjunction with other fundamental statistical design principles, such as randomization, **blocking** (which will be discussed in a later section), and other more general design principles, such as careful selection of factors and their levels.

The key disadvantage of these designs arises when the results of the analysis are unclear due to confounding. To return to the 2^{5-2} example, suppose only the effects associated with $A + BE$ and $D + CE$ were found to be active (where effects of three-factor interactions and higher were omitted from the list). Most experimenters, following the principle of simplicity (sometimes called *Occam's razor*) would declare that A and D were active. The reason is that it would be very unusual for, say, BE to be active even though neither B nor E were active. However, suppose instead that only the effects associated with $A + BE$, $B + AE$, and $E + AB + CD$ were found to be active. Then the principle of simplicity fails to point to one model: the “ A, B, E ” scenario is simple, but so are those of “ A, B, AB ”, “ A, E, AE ”, and “ B, E, BE ”.

The experimenter's prior knowledge can often reduce the likelihood of this problem. In this example, the experimenter should assign the factor least likely to be active (and so, reasonably, least likely to interact with other factors) to the letter E – this can be seen from considering various scenarios and looking only at the main effect and two-factor interactions in Table 2. More generally, this strategy is to (a) select a design with good properties for the particular values of k and p for the experiment and (b) assign the actual factors to the symbols $A, B, \dots$ on the basis of prior knowledge to minimize confounding of effects most likely to be active. Such a strategy is another advantage of fractional factorial designs – the prior knowledge of the experimenter can be used directly in the design of the experiment.

If the data have already been collected and this problem exists, another way to see which factors are active is to collect an additional set of data to eliminate the confounding. This is often done by running a foldover design (see **Foldover Designs**). A more sophisticated approach developed by Meyer *et al.* [15] may also be used, by adding those runs that allow maximum discrimination among a set of plausible models.

If confounding exists and no more data can be collected, it may still be possible to obtain a good solution. In this confounding example, if the objective is to maximize the response and each of the three effects associated with $A + BE$, $B + AE$, and $E + AB + CD$ is positive, then setting A, B , and E at their high levels from the experiment is a good strategy regardless of which of the four scenarios

above is correct. In other experiments, a “pick the winner” strategy may work – simply run at those conditions that produced the best result.

Another disadvantage sometimes seen in these designs is that it may be difficult to distinguish the active effects from the intrinsic variability in the data. However, the ability to make this distinction depends on the total number of runs made, not on the extent of confounding, so this is not a disadvantage of fractional factorial designs *per se*. A reasonable rule of thumb to overcome this shortcoming is to use at least 16 runs in an experiment and especially avoid experiments with 8 or fewer runs [16]. In our example, this would mean that the experimenter should replace the 2^{5-2} design with a 2^{5-1} design. This has the additional advantage of increasing the resolution from III to V. (This is also usually preferred to running a 2^{5-2} design followed by a foldover design, which also requires 16 runs but is only resolution IV.)

Most of our discussion assumes that the design is for a physical (real-life) experiment, that is, an experiment whose results include random variation. By contrast, designs for **computer experiments** that produce deterministic results usually employ other design strategies, such as Latin hypercube designs (see **Latin Hypercube Designs**), to allow smooth nonlinear functions to be fit in the design region. However, if a large number of factors are being studied, and most are expected to be monotonic in the response, then 2^{k-p} designs may still be useful.

Example

Our example comes from an experiment that was described in Snee *et al.* [17]. (The design and analysis have been slightly modified here for simplicity.) Viscosity measurements made in an analytical laboratory appeared to be quite variable, and it was not clear how much of this variation was due to the product itself and how much was due to the variation in the **measurement system**. A classical way to study variation in measurement systems is by a gauge repeatability and reproducibility (R&R) experiment (see **Gauge Repeatability and Reproducibility (R&R) Studies**). There, an attempt is typically made to vary systematically both product samples and those who make the measurements, while trying to hold other

6 Fractional Factorial Designs

Table 3 Factors, treatments combinations, and results of the viscosity ruggedness test

Standard order	A Sample preparation	B Moisture measurement	C Mixing speed (rpm)	D Mixing time (h)	E Healing time (h)	F Spindle	G Protective lid	Viscosity
1	M1	Volume	800	0.5	1	S1	Absent	2796
2	M2	Volume	800	0.5	1	S2	Present	2460
3	M1	Water	800	0.5	2	S1	Present	2904
4	M2	Water	800	0.5	2	S2	Absent	2320
5	M1	Volume	1600	0.5	2	S2	Present	2800
6	M2	Volume	1600	0.5	2	S1	Absent	3772
7	M1	Water	1600	0.5	1	S2	Absent	2420
8	M2	Water	1600	0.5	1	S1	Present	3376
9	M1	Volume	800	3.0	2	S2	Absent	2220
10	M2	Volume	800	3.0	2	S1	Present	2548
11	M1	Water	800	3.0	1	S2	Present	2080
12	M2	Water	800	3.0	1	S1	Absent	2464
13	M1	Volume	1600	3.0	1	S1	Present	3216
14	M2	Volume	1600	3.0	1	S2	Absent	2380
15	M1	Water	1600	3.0	2	S1	Absent	3196
16	M2	Water	1600	3.0	2	S2	Present	2340

conditions constant. Here a more active approach was taken – conditions in the analytical laboratory were intentionally varied over ranges typically used in the laboratory. This use of experimentation is known as a *ruggedness test* [18, 19].

The factors, settings, and results for this experiment are listed in Table 3. Factors *A* and *B* involved methods that were both used in the laboratory as part of an accepted protocol and factors *C*, *D*, and *E* were run over the ranges that had historically been used in the laboratory. For a 16-run design, these five factors could be run in a resolution V design. The chemists in the lab were not as interested in running factors *F* and *G* because they believed these would not be important. However, they agreed to it when they were told these could be included without the need for additional runs (although the resolution of this seven-factor design is IV). The generators for this 2^{7-3} design were chosen to be *BCDE*, *ACDF*, and *ABCG*, resulting in the defining relation $I = BCDE = ACDF = ABEF = ABCG = ADEG = BDFG = CEF G$.

In these 16 runs, factors *A*, *B*, *C*, and *D* form a set of basic factors and the data in Table 3 are listed in the “standard order” for these factors. The 15 effects from these basic factors were calculated and appear in Table 4, with aliasing up through two-factor interactions. These 15 effects were plotted on a half-normal plot (see **Half-Normal Plot**) and Lenth’s

method (see **Lenth’s Method for the Analysis of Unreplicated Experiments**) was used to add decision lines to this plot for $\alpha = 0.05$ ME (margin of error) and SME (simultaneous margin of error) values – see Figure 2.

From this figure, there is good evidence that the main effects of factors *D* (mix time), *C* (mix speed) and *F* (spindle) are active, as well as the effect of $AD + CF + EG$. Because the main effects of both *C* and *F* are large and those of *A*, *E*, and *G* are not,

Table 4 Effects calculated for the viscosity ruggedness test

Effect name	Effect
<i>A</i>	3.5
<i>B</i>	−136.5
<i>C</i>	463.5
<i>D</i>	−300.5
<i>E</i>	113.5
<i>F</i>	−656.5
<i>G</i>	19.5
<i>AB</i> + <i>CG</i> + <i>EF</i>	−28.5
<i>AC</i> + <i>BG</i> + <i>DF</i>	55.5
<i>AG</i> + <i>BC</i> + <i>DE</i>	−72.5
<i>AD</i> + <i>CF</i> + <i>EG</i>	−248.5
<i>BD</i> + <i>CE</i> + <i>FG</i>	65.5
<i>AF</i> + <i>BE</i> + <i>CD</i>	−8.5
<i>AE</i> + <i>BF</i> + <i>DG</i>	−38.5
<i>ABD</i> + other 3 fi’s	37.5

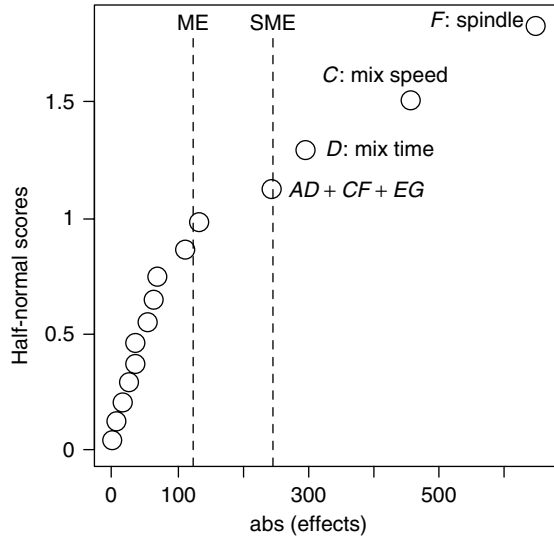


Figure 2 Half-normal plot and Lenth's method applied to the effects from the viscosity ruggedness test

it is reasonable to suspect that the interaction effect is primarily or totally due to the CF interaction. (In fact additional runs were later made to deconfound these terms, and these verified that CF was the active interaction.)

So, aside from making a decision about the $AD + CF + EG$ interaction, the analysis of this design is equivalent to that for a full factorial 2^k design. This is generally true. Also, because fractional factorial designs can accommodate a large number of factors, the study was able to include the spindle factor, and this proved to have the largest effect on the results.

Extensions

The 2^{k-p} designs are the simplest examples of fractional factorial designs. They may be extended in many ways. We consider extensions to blocked designs, **split-plot** designs, three-level and higher-level symmetric designs, nongeometric designs, **supersaturated** designs, and mixed-level designs.

Blocked Designs

Blocking (see **Blocking**) can be a very useful technique to reduce, as well as to identify, experimental error, the sources of which can be hypothesized by

the experimenter. In a full factorial design, simple blocking may be used. For example, in a 2^3 design with factors A , B , and C , suppose that variation is suspected among batches of some raw material, that 8 runs may be made per batch, and that 16 runs are planned to obtain enough precision to detect interesting effects. Then a natural design would be to run all 2^3 combinations on each of the two batches. Here, the units run in each batch are called a *block*, so this design has two blocks, each of size 8.

Now suppose we instead want to study five factors in 16 runs, but still use two blocks of size 8. We could run the same 2^{5-2} fraction in each of the two blocks, but it is more advantageous to run different fractions in each block to reduce the confounding.

One way to approach this problem is to consider a 16-run design in five factors, and then add a so-called blocking factor. Here, we could consider the generator $E = ABCD$ for the 2^{5-1} design and then define a blocking factor $b_1 = AB$. The 2^{5-1} design is then split into two blocks on the basis of the value of b_1 . The resulting defining relation is $I = ABCDE = ABb_1 = CDEb_1$. If we consider that at most main effects and two-factor interactions are a concern, then we can eliminate $ABCDE$ from consideration. Furthermore, in industrial experiments it is customary to assume that the blocking factor(s) will not interact with the treatment factors. If so, under this additional assumption we may write $I = ABb_1$. This shows that all main effects and two-factor interactions are clear of each other and of the blocking factor, except for AB , which is confounded with b_1 .

Suppose we instead need to consider splitting up this 2^{5-1} design into four blocks, each of size 4. A good design to use here is $E = ABCD$, $b_1 = AB$, $b_2 = AC$. This generates four blocks based on the four combinations of (b_1, b_2) , $\{(-1, -1), (+1, -1), (-1, +1), (+1, +1)\}$. These are associated with the 3 degrees of freedom (df) (see **Degrees of Freedom**) due to b_1 , b_2 , and b_1b_2 . The last term, sometimes called a *generalized interaction*, has the same status as b_1 and b_2 . The defining relation is $I = ABCDE = ABb_1 = CDEb_1 = ACb_2 = BDEb_2 = BCb_1b_2 = ADEb_1b_2$.

In this example, we began with $I = ABCDE$, the best design without blocking, and achieved from it the optimal blocked design. This is not always the case. For example, if we need to run a 2^{5-1} design in eight blocks instead of four, we need to start with

a four-letter generator such as $ABCD$ to create the optimal blocked design.

Sometimes trade-offs need to be considered for these designs. For example, say we want to run a 2^{5-1} design in two blocks. We could choose our earlier design, whose defining relation is $I = ABCDE = ABb_1 = CDEb_1$. The weakness of this design is that the block is confounded with the AB interaction. If we instead choose a design whose defining relation is $I = ABCE = ABD b_1 = CDE b_1$, then the blocking factor is no longer confounded with a two-factor interaction; however three pairs of two-factor interactions ($AB + CE$, $AC + BE$, and $AE + BC$) are confounded instead. A listing of two-level blocked designs that also includes this trade-off information may be found in Sun *et al.* [20].

Split-Plot Designs

In many industrial experiments, some factors have levels that are more difficult to change than others. For example, consider an experiment with eight factors, and suppose that four of them, $A - D$, have levels that are relatively difficult to change. Say they can only be run at one setting each day. Suppose the other four factors, $P - S$, are easier to adjust and can be run at two settings each day. Say we can make two runs per day and can run the experiment for 8 days. The resulting design is an example of a split-plot design (see **Split-Plot Designs**). Here, $A - D$ are called *main-plot factors* and each set of two runs per day is called a *main plot*. The factors $P - S$ are called *split-plot factors* and an individual run is called a *split plot*.

For this problem a 2^{8-4} design is suggested, but can such a design be constructed such that the levels of $A - D$ can be run in pairs, as required? This additional restriction makes the construction of these designs more difficult. In this case, such a design is possible, using basic factors A, B, C , and P , and generators $D = ABC$, $Q = BCP$, $R = ACP$, and $S = ABP$. (In fact, this design has minimum aberration (see **Minimum Aberration**) among all 2^{8-4} designs.) See Huang *et al.* [21] and Bingham and Sitter [22] for examples and listings of some designs and Bingham and Sitter [23] for more theoretical details.

More restrictive situations are also common. For example, perhaps some factors will only be varied once every 2 days, some once each day, others twice

each day, and others every hour. The resulting split-split-split-plot design may be more difficult, or even impossible, to construct. Extensive listings for such designs are not available.

Higher-Level Symmetric Designs

Although two-level designs have traditionally been the most widely used in industrial applications, higher-level fractional factorial designs are also used. These include designs in which each factor is run at the same number of levels – such designs are called *symmetric designs*. We illustrate some basic ideas with geometric, three-level, fractional factorial designs.

The confounding that was used in two-level designs, for example, $D = ABC$, cannot be directly extended to higher-level designs. In this two-level example, both the D main effect and the ABC interaction had 1 df, and this made it possible to consider confounding the two effects. However, in a three-level design, the D main effect has 2 df, but the ABC interaction has 8 df, so “ $D = ABC$ ” no longer makes sense. It is possible, however, to confound D with what is called a *2 df piece* of the ABC interaction. In fact, the ABC interaction has four such pieces, and an additional factor could be confounded with each piece. Details on how to create such designs may be found in, for example [4, 5].

Not surprisingly, the resulting confounding in 3^{k-p} designs is much more complex than in 2^{k-p} designs. So-called interaction graphs provide a useful way to examine the implications of these interactions. Sun and Wu [24] provide listings of good 3^{k-p} designs, including interaction graphs.

Higher-level designs for s^{k-p} are available when s is prime or a power of a prime. Such designs are not used often in industrial applications. The arrays used in Latin squares (see **Latin Squares and Related Experimental Designs**) and Greco-Latin squares are equivalent to s^{3-1} and s^{4-2} designs, respectively. In the original development of those designs, only one of the factors was a treatment factor, the remainder being blocking factors. However, these designs may also be used to define fractional factorials, in which all of the factors can be treatment factors.

These symmetric designs are examples of orthogonal arrays (see **Orthogonal Arrays**) of strength 2 or greater (as long as main effects are not confounded

with each other – such confounding is almost always avoided in good experiments).

Nongeometric Symmetric Designs

The 2^{k-p} and 3^{k-p} designs are geometric designs, but designs with more complex aliasing are sometimes used. The most common of these are the Plackett–Burman designs (see **Plackett–Burman Designs**). These two-level designs exist for all $N \leq 200$ whenever N modulo 4 = 0, and the designs have complex aliasing unless N is a power of 2 (in which case the design is a 2^{k-p} design). For an example of this complex aliasing, consider an experiment in which five factors, labeled $A - E$, are studied using the levels based on the first five columns of a 20-run Plackett–Burman design that is generated as shown in **Plackett–Burman Designs**. An example of the resulting aliasing is that the effect associated with column **A** is an estimate of $A + (-BC - BD - BE + CD - CE + DE)/5$. Thus the main effect of A is neither fully orthogonal to nor fully aliased with these interactions. This type of aliasing distinguishes nongeometric designs from geometric designs.

These designs are most widely used when $N = 12, 20, 24, 28$, because they fill in the gaps between the 8-, 16-, and 32-run 2^{k-p} designs.

Supersaturated Designs

For two-level designs in which N is an appropriate value, such as those noted in the previous subsection, up to $N - 1$ factors may be studied in an orthogonal design. Occasionally, consideration may be given to studying more than $N - 1$ factors in N runs, and the resulting fractional factorial design is called *supersaturated* (see **Supersaturated Designs**). These may also be constructed for designs other than two-level designs.

Mixed-Level Designs

Designs in which factors are not all run at the same number of levels are called *mixed-level designs*. Constructing fractions of such designs that have desirable properties, such as orthogonality, is often more complex than for symmetric designs.

The most common of these designs are the resolution III designs, known as *main effect plans* (see

Main Effect Designs) or *orthogonal main effect designs*, also called OMEDs. For example, it is possible to construct an OMED with 7 factors at three levels and 16 factors at two levels in 32 runs, and notation such as “ $3^7 2^{16} // 32$ ” may be used to denote this design. As another example, consider the 2^{5-1} design in four blocks that we had constructed earlier. If we replace the four-level blocking factor with a four-level treatment factor, we have created a $2^5 4^1 // 16$ design.

A listing of such designs for factors having 2–6 levels and 50 runs or fewer is available in [5]. More general orthogonal designs, known as *orthogonal arrays* (see **Orthogonal Arrays**), can also be constructed. Also, see Chen *et al.* [25] for a listing of mixed-level designs at two and three levels.

Closing Remarks

The construction of fractional factorial designs has been, and continues to be, extensively studied. A large number of good designs are available from the references cited, and are also available from numerous statistical software programs and electronic resources including the National Institute for Standards and Technology [26]. Many statistical software programs can create other designs from user requirements, using algorithms based on design-construction rules or optimal-design theory (see **Optimal Design**).

Analysis methods for designs include subjective graphical methods, Lenth’s method (for orthogonal designs), or more sophisticated techniques (see **Lenth’s Method for the Analysis of Unreplicated Experiments** for some examples).

References

- [1] Wu, C.F.J. & Chen, Y.Y. (1992). A graph-aided method for planning two-level experiments when certain interactions are important, *Technometrics* **34**, 162–175.
- [2] Franklin, M.F. & Bailey, R.A. (1977). Selection of defining contrasts and confounded effects in two-level experiments, *Applied Statistics* **26**, 321–326.
- [3] Box, G.E.P., Hunter, J.S. & Hunter, W.G. (2005). *Statistics for Experimenters*, 2nd Edition, John Wiley & Sons, New York.
- [4] Wu, C.F.J. & Hamada, M. (2000). *Experiments: Planning, Analysis, and Parameter Design Optimization*, John Wiley & Sons, New York.

-
- [5] Lorenzen, T.J. & Anderson, V.L. (1993). *Design of Experiments: A No-Name Approach*, Marcel Dekker, New York.
- [6] Chen, J. & Wu, C.F.J. (1991). Some results on s^{n-k} fractional factorial designs with minimum aberration or optimal moments, *The Annals of Statistics* **19**, 1028–1041.
- [7] Chen, J. (1992). Some results on 2^{n-k} fractional factorial designs and search for minimum aberration designs, *The Annals of Statistics* **20**, 2124–2141.
- [8] Tang, B. & Wu, C.F.J. (1996). Characterization of minimum aberration 2^{n-k} designs in terms of their complementary designs, *The Annals of Statistics* **24**, 2549–2559.
- [9] Chen, H. & Hedayat, A.S. (1996). 2^{n-1} Designs with weak minimum aberration, *The Annals of Statistics* **24**, 2536–2548.
- [10] Butler, N.A. (2003). Some theory for constructing minimum aberration fractional factorial designs, *Biometrika* **90**, 233–238.
- [11] Finney, D.J. (1945). The fractional replication of factorial arrangements, *Annals of Eugenics* **12**, 291–301.
- [12] Ledolter, J. (2002). A case study in design of experiments: improving the manufacture of viscous fiber, *Quality Engineering* **15**, 311–322.
- [13] Snelick, R., Jájá, J., Kacker, R. & Lyon, G. (1994). Synthetic-perturbation techniques for screening shared memory programs, *Software-Practice and Experience* **24**, 679–701.
- [14] Louviere, J.J. (1988). *Analyzing Decision Making: Metric Conjoint Analysis*, Sage University Paper Series on Quantitative Applications in the Social Sciences, 07-067, Sage Publications, Newbury Park.
- [15] Meyer, R.D., Steinberg, D.M. & Box, G. (1996). Follow-up designs to resolve confounding in multifactor experiments, *Technometrics* **38**, 303–313; discussion 314–332.
- [16] Hamada, M. & Balakrishnan, N. (1998). Analyzing unreplicated factorial experiments: a review with some new proposals, *Statistica Sinica* **8**, 1–41.
- [17] Snee, R.D., Hare, L.B. & Trout, J.R. (1985). *Experiments in Industry: Design, Analysis, and Interpretation of Results*, American Society for Quality Control, Milwaukee, pp. 31–33.
- [18] Youden, W.J. (1961). Experimental design and ASTM committees, *Materials Research and Standards* **1**, 862–867.
- [19] Wernimont, G. (1977). Ruggedness evaluation of test procedures, *ASTM Standardization News* **5**, 13–16.
- [20] Sun, D.X., Wu, C.F.J. & Chen, Y. (1997). Optimal blocking schemes for 2^n and 2^{n-p} designs, *Technometrics* **39**, 298–307.
- [21] Huang, P., Chen, D. & Voelkel, J. (1998). Minimum aberration split-plot designs, *Technometrics* **40**, 314–326.
- [22] Bingham, D. & Sitter, R.R. (1999). Minimum aberration two-level fractional factorial split-plot designs, *Technometrics* **41**, 62–70.
- [23] Bingham, D.R. & Sitter, R.R. (1999). Some theoretical results for fractional factorial split-plot designs, *The Annals of Statistics* **27**, 1240–1255.
- [24] Sun, D.X. & Wu, C.F.J. (1994). Interaction graphs for three-level fractional factorial designs, *Journal of Quality Technology* **4**, 297–307.
- [25] Chen, J., Sun, D.X. & Wu, C.F.J. (1993). A catalogue of two-level and three-level fractional factorial designs with small runs, *International Statistical Review* **61**, 131–145.
- [26] <http://www.itl.nist.gov/div898/handbook/>. (accessed 2007), Fractional-factorial designs are listed in the *Improve* section of this electronic document.

Related Articles

Screening Designs.

JOSEPH G. VOELKEL

Gauge Repeatability and Reproducibility (R&R) Studies

Introduction and Some Qualitative Matters

A given measurement device (henceforth called a *gauge*) is often used by more than one individual (henceforth called *operators*) to make routine measurements of some important feature of production outcomes (henceforth called *parts*). In such cases it is important to quantify measurement consistency (or lack thereof), both for a fixed operator on a fixed part, and between operators on a fixed part. The first kind of consistency or variation is commonly called the *repeatability* component and the second is termed the *reproducibility* component. Where remeasurement is possible, it is traditional to try to assess these two components of measurement variation by having each of J operators measure each of I parts m times apiece, producing a balanced $I \times J$ **factorial** dataset with responses

$$y_{ijk} = \text{the } k\text{th measurement of part } i \text{ by operator } j \quad (1)$$

$i = 1, 2, \dots, I$; $j = 1, 2, \dots, J$; and $k = 1, 2, \dots, m$. (It is obvious that where remeasurement is not possible, e.g. because measurement is destructive, there can be no direct assessment of measurement precisions unclouded by part-to-part differences. A discussion of what can be done indirectly on the basis of different sets of assumptions in the destructive testing context is provided in **Gauge Repeatability and Reproducibility (R&R) Studies, Destructive Testing.**)

It should go without saying that in the collection of gauge R&R data, some scheme of rotation should be employed so that all m observations in a part-by-operator cell are not collected one after another, creating psychological pressure on the operator to get unnaturally consistent results. Complete **randomization** of the order of making IJm measurements may not be feasible, but operators making m “repeats” one after another will produce repeatability variation that is really too small, and quite

likely reproducibility variation that is really too large.

Further, what is under investigation in an R&R study is essentially *variances* and functions thereof. For example, reproducibility variation might be conceptualized as a between-operator variance. Sound statistical intuition should be that the effective **estimation** of variances requires relatively large “sample sizes.” While in a typical gauge R&R study, there are IJm measurements, there are only J operators, and standard industry **data collection** forms and spread sheets for calculation are typically set up for J on the order of 2–5! This feature of common practice is thus in clear contradiction to well-recognized statistical realities.

It is standard to talk about gauge R&R as if measurement precision were a property of only the gauge and how it is used. This is, of course, not really the case. The same gauge used to measure different features of the same part can yield different precisions of measurement. The same gauge used to measure the same feature of otherwise “identical” parts made of two different materials can yield different precisions of measurement. (Consider, for example, using a caliper to measure “the diameter” of 3 cm nominally circular cylinders made of precision machined stainless steel and then some made of wood.) This has the implication that the extension of gauge R&R results from one specific context to another can only be on the basis of expert judgment (that the two contexts are similar enough that results from one predict those from the other).

Strictly interpreted, “R&R” has nothing to do with part-to-part variation (except as the reality above impacts observed variability in measurement). And standard analyses of R&R data do not really require the presence of more than 1 part in the study nor produce inferences about part-to-part variation. Nevertheless, typical data collection forms (and spread sheets for calculation) are typically set up for I on the order of 10. To the extent that anyone would actually look carefully at results for different parts and consider whether the models and analyses used with gauge R&R data are appropriate in light of part-to-part comparisons, this is perhaps not a bad thing. But in practice, it seems to create the illusion of having “lots of information” because IJm is large, while it is basically really *only* J that governs how much can be learned about the fundamental issue of reproducibility.

Standard Models and Analyses

Standard analyses for gauge R&R data are based on the two-way random-effects model

$$y_{ijk} = \mu + \alpha_i + \beta_j + \alpha\beta_{ij} + \epsilon_{ijk} \quad (2)$$

where all of the random effects $\alpha_i, \beta_j, \alpha\beta_{ij}, \epsilon_{ijk}$ are independent, the α_i are $N(0, \sigma_\alpha^2)$, the β_j are $N(0, \sigma_\beta^2)$, the $\alpha\beta_{ij}$ are $N(0, \sigma_{\alpha\beta}^2)$, and the ϵ_{ijk} are $N(0, \sigma^2)$. Notice that in a case where $I = 1$ (there is only one part involved in the R&R study) this is

$$y_{1jk} = \mu + \alpha_1 + \beta_j + \alpha\beta_{1j} + \epsilon_{1jk} \quad (3)$$

and with $\gamma_j = \beta_j + \alpha\beta_{1j}$ this is

$$y_{1jk} = \mu + \alpha_1 + \gamma_j + \epsilon_{1jk} \quad (4)$$

For $\sigma_\gamma^2 = \sigma_\beta^2 + \sigma_{\alpha\beta}^2$, the γ_j are $N(0, \sigma_\gamma^2)$ and this is a variant of the one-way random-effects model where it is obviously impossible to disentangle the parameters μ and σ_α^2 that are not of primary concern in this context. Rather, it is the two variances

$$\sigma_{\text{reproducibility}}^2 \equiv \sigma_\gamma^2 = \sigma_\beta^2 + \sigma_{\alpha\beta}^2 \quad (5)$$

and

$$\sigma_{\text{repeatability}}^2 \equiv \sigma^2 \quad (6)$$

that are of primary interest. The first of these can be thought of as representing a variance of operator-specific measurement biases.

The best available technologies for inference in the two-way random-effects model (2) or in the case of a single part, the usual one-way random-effects model, are

- so-called “reml” estimation as implemented, for example, in the R statistical package and described in Pinheiro and Bates [1]
- Bayes analyses using “flat prior distributions” as can, for example, be implemented using the WinBUGS program (see <http://www.mrc-bsu.cam.ac.uk/bugs/welcome.shtml>).

The most commonly used technologies for inference from gauge R&R datasets are

- range-based point estimates as described in the Automotive Industry Action Group [2]

- ANOVA-based inference, with confidence statements based on Satterthwaite approximations and various improvements thereof represented by Burdick and Graybill [3] and subsequent related work (*see Gauge Repeatability and Reproducibility (R&R) Studies, Confidence Intervals for*).

Of the commonly used methods, the ANOVA methods are expected to be superior to the range-based ones on the basis of both efficiency of point estimation and the fact that they provide confidence statements about $\sigma_{\text{reproducibility}}^2$, $\sigma_{\text{repeatability}}^2$, and $\sigma_{\text{R\&R}}^2 \equiv \sigma_{\text{repeatability}}^2 + \sigma_{\text{reproducibility}}^2$. Of the best existing technologies, Bayes analyses have the potential advantage of being both very simple to implement and also providing obvious ways to find credible limits based on posterior distributions not only for $\sigma_{\text{reproducibility}}^2$, $\sigma_{\text{repeatability}}^2$, and $\sigma_{\text{R\&R}}^2$, but also for other parametric functions as well like, e.g., the ratios $(\sigma_{\text{reproducibility}}^2 / \sigma_{\text{R\&R}}^2)$ and $(\sigma_{\text{repeatability}}^2 / \sigma_{\text{R\&R}}^2)$.

Figures of Merit for R&R Study Designs

Now consider the question of “How good is a particular choice of I, J , and m for an R&R study?” Strictly speaking, any quantitative figure of merit to be associated with a given choice (henceforth, a “design”) for a gauge R&R study should be matched to the method of estimation employed in analyzing the results. But in the interest of providing simple explicit formulas, we are here going to consider what might be done for ANOVA-based estimation.

Let us begin with an R&R study using a single part, i.e. the $I = 1$ instance of the balanced R&R data structure. It is completely standard that with the obvious meaning of treatment mean square ($MSTr$) and error mean square (MSE), for a balanced dataset $MSTr$ and MSE are multiples of independent χ^2 random variables with respective **degrees of freedom** $J - 1$ and $J(m - 1)$ and

$$EMSTr = m\sigma_\gamma^2 + \sigma^2 \text{ and } EMSE = \sigma^2 \quad (7)$$

and thus variances

$$\text{Var}MSTr = \frac{2(m\sigma_\gamma^2 + \sigma^2)^2}{(J - 1)} \text{ and}$$

$$\text{Var}MSE = \frac{2\sigma^4}{J(m-1)} \quad (8)$$

So with

$$\hat{\sigma}_{\text{reproducibility}}^2 = \frac{1}{m} (MSTr - MSE) \quad (9)$$

and

$$\hat{\sigma}_{\text{repeatability}}^2 = MSE \quad (10)$$

the vector $\begin{pmatrix} \hat{\sigma}_{\text{repeatability}}^2 \\ \hat{\sigma}_{\text{reproducibility}}^2 \end{pmatrix}$ has **covariance** matrix

$$\begin{pmatrix} \frac{2\sigma^4}{J(m-1)} & -\frac{2\sigma^4}{Jm(m-1)} \\ -\frac{2\sigma^4}{Jm(m-1)} & \frac{2(m\sigma_\gamma^2 + \sigma^2)^2}{m^2(J-1)} + \frac{2\sigma^4}{Jm^2(m-1)} \end{pmatrix} \quad (11)$$

And with

$$\hat{\sigma}_{\text{R\&R}}^2 = \frac{1}{m} MSTr + \frac{(m-1)}{m} MSE \quad (12)$$

one has

$$\begin{aligned} E\hat{\sigma}_{\text{R\&R}}^2 &= \sigma_{\text{R\&R}}^2 \text{ and} \\ \text{Var}\hat{\sigma}_{\text{R\&R}}^2 &= \frac{(m-1)2\sigma^4}{Jm^2} + \frac{2(m\sigma_\gamma^2 + \sigma^2)^2}{m^2(J-1)} \end{aligned} \quad (13)$$

It is obvious then that precision of estimation of $\sigma_{\text{reproducibility}}^2$ and $\sigma_{\text{repeatability}}^2$ and $\sigma_{\text{R\&R}}^2$ depend upon the values of the two variance components. But simple formulas like equations (11) and (13) allow one to input planning values for σ_γ^2 and σ^2 and trial values for J and m and get some idea of the quality of information about important parametric functions that could be achieved. (Further, Taylor series/propagation of error approximations applied to nonlinear functions of $(\hat{\sigma}_{\text{repeatability}}^2, \hat{\sigma}_{\text{reproducibility}}^2)$ can be used to get some rough idea of what one can expect to learn about a corresponding nonlinear function of $(\sigma^2, \sigma_\gamma^2)$ from a particular choice of design.) Vardeman and VanValkenburg [4] took the above kind of approach to the full-blown design question for an $I \times J \times m$ gauge R&R study based on means and variances of ANOVA mean squares. They found, e.g., that the determinant of the covariance matrix for linear

combinations of mean squares that are unbiased for $(\sigma^2, \sigma_\gamma^2)'$ can be written in fairly explicit terms as

$$4\sigma^8 \frac{(1+mr)^2 + m^2(I-1)t^2}{I^2 J m^2 (J-1)(m-1)}$$

for $r = \sigma_\gamma^2/\sigma^2$ and $t = \sigma_\beta^2/\sigma^2$. One might take this as a single figure of merit for a particular choice of I, J , and m where only $\sigma_{\text{repeatability}}^2$ and $\sigma_{\text{reproducibility}}^2$ are to be estimated. Note that this figure of merit does not depend upon σ_α^2 . (Recall that the variance components were defined in the section titled "Standard Models and Analyses.") Not too surprisingly, for a given set of planning values for variance components (equivalently values for r and t) and given total number of measurements IJm , this criterion causes one to prefer choices with $I = 1$ if only the narrow issue of optimizing it is considered. More details on this kind of analysis can be found in the Vardeman and VanValkenburg paper.

Conclusion

Most gauge R&R studies have the fairly simple structure of a replicated full factorial design. Nevertheless, we have here identified some ways in which they are often poorly suited to their supposed purpose, and provided some ways of evaluating whether a proposed design is likely to provide information adequate to really identify repeatability and reproducibility components of measurement variance.

Acknowledgment

This work was partially funded through NSF grant #DMS-0502347.

References

- [1] Pinheiro, J.C. & Bates, D.M. (2000). *Mixed Effects Models in S and S-PLUS*, Springer-Verlag, New York.
- [2] Automotive Industry Action Group. (2002). *Measurement Systems Analysis*, AIAG, Detroit.
- [3] Burdick, R.K. & Graybill, F.A. (1992). *Confidence Intervals on Variance Components*, Marcel Dekker, New York.
- [4] Vardeman, S.B. & VanValkenburg, E.S. (1999). Two-way random-effects analyses and gauge R&R studies, *Technometrics* **41**(3), 202–211.

STEPHEN B. VARDEMAN

Half-Normal Plot

A half-normal plot is a graphical display of effect estimates (contrasts) from a factorial experiment (see **Factorial Experiments**) to discern active contributors from inert ones. It was introduced in Daniel [1] and is a popular data analysis tool for factorial experiments involving factors at two levels, especially without replication. Many have developed alternatives to Daniel's approach including Box and Meyer [2] and Lenth [3] (see **Lenth's Method for the Analysis of Unreplicated Experiments**). See Hamada and Balakrishnan [4] for a detailed overview and comparison of this method with many others.

The basic idea is to rank order the absolute value of the effect estimates and to display them on a special probability graph. If the underlying distribution of the data is normal, the effect estimates also follow a **normal distribution**. The absolute values of these effect estimates are nonnegative, and when the true effects are zero, their corresponding distribution is a half-normal, which resembles the right half of a bell-shaped curve. When the estimates of these null effects are sorted and plotted against **quantiles** of a half-normal distribution, they should follow a straight line extending from the origin. When there are nonzero effects in the experiment, their absolute values should deviate from the straight-line pattern formed by the estimates of the null effects. Any effect that deviates from this line would be considered a significant contributor. This approach is particularly useful when there are no replicates in the experiment and no estimate of error variance to gauge the magnitude of these effects. Even though the method is designed to compare effects from the popular class of two-level factorial experiments (see **Factorial Experiments**; **Fractional Factorial Designs**; **Main Effect Designs**; **Screening Designs**; and Box *et al.* [5]), it can be adapted to compare standardized effects from experiments with more than two levels.

Define an effect estimate as the difference between the average response at the high setting and the average response at the low setting for that effect, which is also twice the regression coefficient obtained by fitting a linear regression model with the two levels of each factor coded to -1 and 1 (see **Main Effects**). Consider an experiment with $N = 2^{k-p}$ runs, where k is the number of two-level factors

and p is the fraction order. For $k = 4$ and $p = 0$, one would have a 16-run full factorial experiment (see **Factorial Experiments**). For $k = 5$ and $p = 1$, one would have a one-half fraction of a 32-run experiment, which could study five factors with 16 runs (see **Fractional Factorial Designs**). The effect estimates are computed from differences of averages of $N/2$ runs of the experiment.

Table 1 summarizes a four factor experiment from an integrated circuit manufacturing process, taken from Table 3-1 of [6]. Assume that a linear model is appropriate for this experiment.

$$Y_{ijklm} = \mu + A_i + B_j + C_k + D_l + AB_{ij} + AC_{ik} + AD_{il} + BC_{jk} + BD_{jl} + CD_{kl} + ABC_{ijk} + ABD_{ijl} + ACD_{ikl} + BCD_{jkl} + ABCD_{ijkl} + \varepsilon_{ijklm} \quad (1)$$

Y_{ijklm} represents the response and μ represents the overall average response. ε represents an independent random error term, which has a zero mean and variance σ^2 . For this full factorial experiment, all the effects can be estimated. The terms in equation (1) excluding μ and ε represent these 15 effects, and can be expressed as twice the regression coefficients if the low and high levels are coded to -1 and 1 . Since there is no replication in the above experiment, there is no estimate for residual variance and no

Table 1 A four-factor experiment regarding deposition thickness of an integrated circuit manufacturing process^(a)

Run	A	B	C	D	Response
1	-1	-1	-1	-1	14.59
2	-1	-1	-1	1	13.59
3	-1	-1	1	-1	14.24
4	-1	-1	1	1	14.05
5	-1	1	-1	-1	14.65
6	-1	1	-1	1	13.94
7	-1	1	1	-1	14.4
8	-1	1	1	1	14.14
9	1	-1	-1	-1	14.67
10	1	-1	-1	1	13.72
11	1	-1	1	-1	13.84
12	1	-1	1	1	13.9
13	1	1	-1	-1	14.56
14	1	1	-1	1	13.88
15	1	1	1	-1	14.3
16	1	1	1	1	14.11

^(a) © John Wiley & Sons, New York.

2 Half-Normal Plot

analysis of variance (*see Analysis of Variance*) to assess the significance of these effects.

Under the effect sparsity principle, namely, only a small portion of effects are likely to be active in an experiment, so one would compare the effect estimates in equation (1) against each other and identify the extremely large ones as likely contributors to the experiment.

Table 2 lists the effect estimates of the experiment in Table 1, and Figure 1 plots their absolute values against the corresponding half-normal quantiles. From the above example, the main effect (*see Main Effects*) estimate for factor A is -0.0775 because the average response at the “1” setting of A is 14.1225 and the average response at the “ -1 ” setting of A is 14.2. Their difference is -0.0775 . If levels of factor A are coded as -1 and 1, a first-order regression model would yield a slope of -0.03875 for this factor. The effect estimates for interaction terms (*see Interactions*) can be found by multiplying the columns of the corresponding factors and then taking the difference of the average response associated with the “1” setting and the average response associated with the “ -1 ” setting. For example, the effect estimate for the AB interaction is 0.0075 because the

Table 2 Effect estimates from the four-factor experiment shown in Table 1

Effect label	Estimate
A	-0.0775
B	0.1725
C	-0.0775
D	-0.4900
AB	0.0075
AC	-0.0925
AD	0.0500
BC	0.0575
BD	0.0300
CD	0.3450
ABC	0.0975
ABD	-0.0250
ACD	0.0300
BCD	-0.1100
ABCD	-0.0200

average response at the “1” setting of AB is 14.165 and the average response at the “ -1 ” setting of AB is 14.1575.

Take the absolute value of these effect estimates and assign half-normal scores associated with

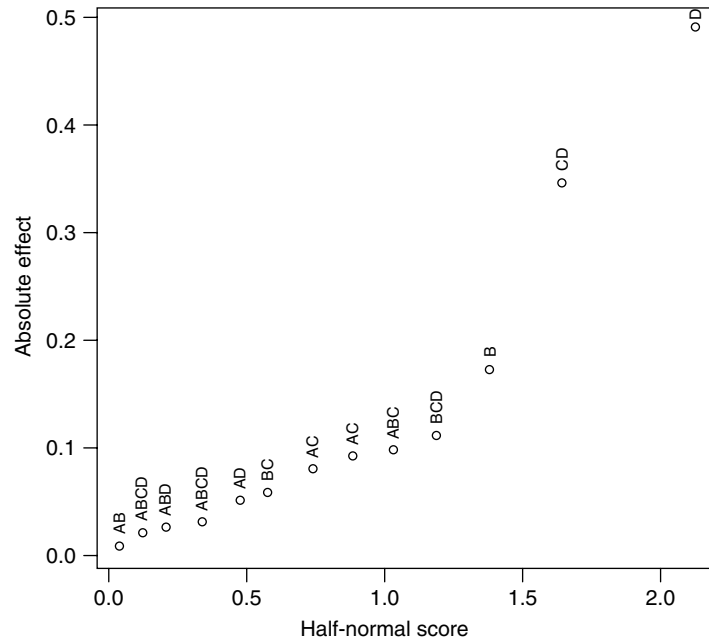


Figure 1 Half-normal plot for the example in Table 1. The main effects B and D and the interaction effect CD stand out as active contributors in this experiment

them. Let $\hat{\beta}_1, \hat{\beta}_2, \dots, \hat{\beta}_M$ be the effect estimates, and $|\hat{\beta}|_{(1)} < |\hat{\beta}|_{(2)} < \dots < |\hat{\beta}|_{(M)}$ be the ordered absolute effect estimates. $M = N - 1$ when the experiment involves only two-level factors. For the i th ordered absolute effect estimate $|\hat{\beta}|_{(i)}$, compute its half-normal score by $\Phi^{-1}(0.5 + 0.5(i - 0.5)/M)$. Then plot $|\hat{\beta}|_{(i)}$ on the vertical axis *versus* its half-normal score on the horizontal axis. $\Phi^{-1}(\cdot)$ is the inverse of the standard normal distribution function, namely, $\Phi(x) = \int_{-\infty}^x 1/\sqrt{2\pi} \exp(-z^2/2) dz$.

The use of the half-normal plot is preferred over a regular normal plot, because the visual ability to identify large contributors in a regular normal plot is obscured when these contributors are spread out in both the positive and the negative directions. They are more visible when plotted in their absolute values on a half-normal plot.

This analysis technique can also be used in fractional factorial designs (see **Fractional Factorial Designs**). The basic approach is the same as before except that the confounding structure of the experiment is labeled, so that the correct active effects are identified. For example, consider Exercise 9 in Chapter 6 of Box *et al.* [5], where an experiment with 16 runs was conducted to study the impact of five welding parameters on the heat input (y). Tables 3 and 4 display the factor description, the raw data, and the design. This is a one-half fraction of a 32-run experiment. Only 15 out of 31 effects can be estimated.

This design is a resolution- V design (see **Factorial Designs, Resolution of**) where the main effects (see **Main Effects**) are confounded with (inseparable from) the fourth-order interaction effects (see **Aliasing in Fractional Designs**). Using the defining relationship of this design, one would find the confounding pattern of all the effects. For example, the interaction effect (see **Interactions**) AB is inseparable from CDE . Hence, they are aliases of each other. Table 5 lists the effect estimates and

Table 3 Factor description for the welding example

Factor label	Description	−1	1
A	Open circuit voltage	31	34
B	Slope	11	6
C	Electrode melt-off rate	162	137
D	Electrode diameter	0.045	0.035
E	Electrode extension	0.375	0.625

Table 4 Design matrix and response of the welding experiment

Run order	A	B	C	D	E	y
12	−1	−1	−1	−1	1	3318
1	1	−1	−1	−1	−1	4141
2	−1	1	−1	−1	−1	3790
6	1	1	−1	−1	1	4061
15	−1	−1	1	−1	−1	3431
8	1	−1	1	−1	1	3425
7	−1	1	1	−1	1	3507
4	1	1	1	−1	−1	3765
11	−1	−1	−1	1	−1	2580
14	1	−1	−1	1	1	2450
3	−1	1	−1	1	1	2319
16	1	1	−1	1	−1	3067
13	−1	−1	1	1	1	1925
10	1	−1	1	1	−1	2466
5	−1	1	1	1	−1	2485
9	1	1	1	1	1	2450

their corresponding aliases (see **Aliasing in Fractional Designs**). Figure 2 exhibits these effects in a half-normal plot. On the basis of experience, many have argued that lower-order effects (main and second-order effects) are more likely to dominate in an experiment than higher-order effects (fourth- and fifth-order effects). Therefore, the effects listed in the left column of Table 5 are the likely effects when they appear to deviate from the line in the half-normal plot.

In a 2^{k-p} fractional factorial experiment (see **Fractional Factorial Designs**), the half-normal plot still consists of $N - 1$ effect estimates but some of them are confounded with other effects. It is the user's task to ensure the interpretation of active contributors with their associated aliases.

The $N - 1$ effect estimates in the two-level factorial experiments are orthogonal contrasts. To extend this approach to factorial experiments involving three- or higher-level factors, one would need to construct orthogonal contrasts and standardize them so that they have a common variance before plotting. Because there are a number of ways to construct orthogonal contrasts for three- or higher-level factors, there are multiple half-normal plots for the same experiment. The experimenter should decide which set of orthogonal contrasts is relevant to the factors.

The half-normal plot is a popular tool to identify “active” effects in a two-level factorial or fractional factorial experiment because of its ease of use, simplicity, and graphical nature. However,

4 Half-Normal Plot

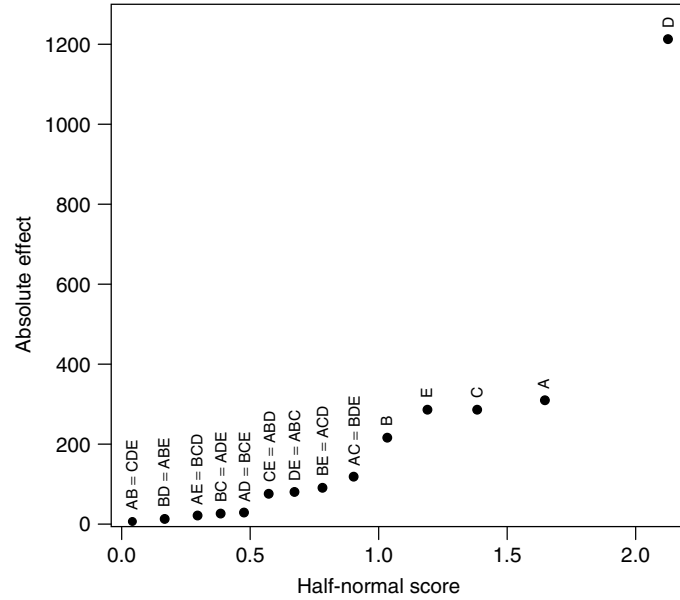


Figure 2 Half-normal plot of effect estimates from the welding experiment. The labels for two-way interaction effects include their corresponding aliases. The main effect D stands out as a significant contributor in this experiment

Table 5 Effect estimates and their aliases from the welding experiment

Effect label	Estimate	Alias
A	308.75	BCDE
B	213.50	ACDE
C	−284.00	ABDE
D	−1212.00	ABCE
E	−283.75	ABCD
AB	1.75	CDE
AC	−119.25	BDE
AD	−27.75	BCE
AE	20.50	BCD
BC	26.50	ADE
BD	11.50	ACE
BE	91.25	ACD
CD	11.50	ABE
CE	73.75	ABD
DE	−79.75	ABC

one weakness of this approach is its subjectivity in identifying active contributors. Even though Daniel [1] proposed a way to quantify significance of effects, most practitioners ignore this aspect of the method. Many have proposed ways to augment this ocular approach with more objective criteria to determine “significance”. For example, Zahn [7]

provided “guardrails” for the half-normal plots. See Hamada and Balakrishnan [4] for a list of these alternatives, technical comparisons, and recommendations. Other nontraditional fractionated designs also favor the half-normal plot as a data analysis tool. These include Plackett–Burman designs (*see Plackett–Burman Designs*) and orthogonal arrays (*see Orthogonal Arrays*).

The half-normal plot provides a means to identify large contributors in an experiment. However, additional work such as model selection, diagnostic, and confirmation study should follow to complete the analysis of the experiment.

References

- [1] Daniel, C. (1959). Use of half-normal plots in interpreting factorial two-level experiments, *Technometrics* **1**, 311–341.
- [2] Box, G.E.P. & Meyer, R.D. (1986). An analysis for unreplicated fractional factorial, *Technometrics* **28**, 11–18.
- [3] Lenth, R.V. (1989). Quick and easy analysis of unreplicated factorials, *Technometrics* **31**, 469–473.
- [4] Hamada, M.S. & Balakrishnan, N. (1998). Analyzing unreplicated factorial experiments: a review with some new proposals (with discussion), *Statistica Sinica* **8**, 1–41.

- [5] Box, G.E.P., Hunter, J.S. & Hunter, W.G. (2005). *Statistics for Experimenters: Design, Innovation, Discovery*, 2nd edition, John Wiley & Sons, New York.
- [6] Wu, C.F.J. & Hamada, M.S. (2000). *Experiments: Planning, Analysis, and Parameter Design Optimization*, John Wiley & Sons, New York.
- [7] Zahn, D.A. (1975). Modification of and revised critical values for the half-normal plot, *Technometrics* **17**, 189–200.

WINSON TAAM

Block Designs, Incomplete

Introduction

One of the Fisherian principles of good experimentation (*see* **Factorial Experiments**) is the stratification of the available experimental material into homogeneous blocks (*see* **Blocking**), so that treatments may be compared within blocks. An experiment in which v treatments are to be compared using b blocks of size k_j , $j = 1, \dots, b$, is called a *blocked experiment*, and if $k_j < v$ for at least one block, the experiment is known as an *incomplete block experiment*. The resulting allocation of treatments to the blocks is called an *incomplete block design* (IBD). For example, in an experiment to compare the tastes of six antibiotic syrups for children, it was considered inadvisable to subject any child to more than three of the antibiotics. Here $v = 6$ and the children formed the blocks with block size $k_j = 3$ for all j . Suppose that the six treatments are represented by the symbols 1, 2, $\dots$, 6. The problem is to allocate sets of three treatments to a group of b children so that all treatments could be compared. All treatment comparisons were of equal importance, so that a design formed from all possible combinations of the six treatment symbols taken three at a time would suffice. A total of ${}^6C_3 = 20$ children would be required to form the IBD shown in Table 1.

This is an example of an *unreduced balanced incomplete block design* (UBIBD); *balanced* because all treatments are equally represented and *unreduced* because all possible sets of three treatments are used. For larger v , the number of blocks needed to produce a UBIBD can become unmanageably large, this prompts one to look for designs that use fewer blocks but maintain balance or partial balance if full balance is not possible. An example of a reduced balanced incomplete block design (BIBD) for six treatments requiring only 10 blocks of three treatments per block is shown in Table 2.

Table 2 Reduced balanced incomplete block design in 10 blocks of three treatments per block

Block number									
1	2	3	4	5	6	7	8	9	10
1	1	1	1	1	2	2	2	3	4
2	2	3	3	4	3	3	4	5	5
5	6	4	6	5	4	5	6	6	6

Early Development and General Properties of Balanced Incomplete Block Designs

Yates [1, 2] first described the use of IBDs in large-scale agricultural trials to compare crop varieties, but since then an extensive literature has been produced on different construction methods and analyses for IBDs. Bose [3, 4] and Fisher [5] considered the structure and construction of BIBDs. Partially balanced incomplete block designs (PBIBD) were introduced by Bose and Nair [6]. Further extensions of the BIBDs were made by Youden [7], who introduced BIBDs for eliminating heterogeneity in two directions, extending the Latin square concept (*see* **Latin Squares and Related Experimental Designs**) to Youden “squares”.

Although for given v and k , a BIBD always exists (the unreduced design for example), there are constraints placed on the number of blocks and treatment replications for the existence of reduced BIBDs. In a BIBD, each treatment is replicated r times and each block size is k . If λ is the number of times each pair of treatments occurs together in the same block, then the five parameters, v , b , k , r , and λ (all integers), must satisfy the conditions:

$$vr = bk = n \text{ where } n \text{ is the number of experimental units} \quad (1)$$

$$\text{and } \lambda(v - 1) = r(k - 1) \quad (2)$$

Shrikhande [8] showed that these are necessary but not sufficient conditions for a reduced BIBD to exist.

Table 1 An unreduced balanced incomplete block design in 20 blocks of three treatments per block

Block number																			
1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20
1	1	1	1	1	1	1	1	1	1	2	2	2	2	2	2	3	3	3	4
2	2	2	2	3	3	3	4	4	5	3	3	3	4	4	5	4	4	5	5
3	4	5	6	4	5	6	5	6	6	4	5	6	5	6	6	5	6	6	6

2 Block Designs, Incomplete

For example, there is no BIBD with parameters $v = 21, b = 28, r = 8, k = 6$, and $\lambda = 2$, even though equations (1) and (2) are satisfied. For certain parameter combinations satisfying these equations, it is possible to use complete sets of orthogonal Latin squares (see **Latin Squares and Related Experimental Designs**) to produce BIBDs. For example, suppose that $v = t^2$ and $k = t$, where t is a prime (or power of a prime); then the complete set of $t - 1$ orthogonal Latin squares may be used to produce a BIBD with $b = t(t + 1)$, $r = t + 1$, and $\lambda = 1$. The t^2 treatments are first arranged sequentially in a $t \times t$ square. The first t blocks of the BIBD are the rows of the square, the second t blocks are the columns of the square; then for each of the $t - 1$ orthogonal Latin squares, a set of t blocks is produced corresponding to the treatments associated with the same symbol in the Latin square. For example, for $t = 3$, the $v = 9$ treatment symbols are laid out in the 3×3 square as below, together with the two orthogonal Latin squares given by Roman and Greek letters,

1	2	3	<i>a</i>	<i>b</i>	<i>c</i>	α	β	γ
4	5	6	<i>b</i>	<i>c</i>	<i>a</i>	γ	α	β
7	8	9	<i>c</i>	<i>a</i>	<i>b</i>	β	γ	α

Table 3 shows the 12 blocks, in four sets of three blocks, for the corresponding BIBD produced using this method.

This construction method necessarily produces $t + 1$ sets of t blocks where each set contains all t^2 treatments. Such BIBDs are called *resolvable designs*. For every BIBD there is a second BIBD, with the same number of treatments and blocks, formed from the treatments not in the blocks of the first BIBD. Such a BIBD is known as a *complementary BIBD*. For the BIBDs formed from the orthogonal Latin squares described above, the complementary BIBD has parameters $v = t^2, k = t(t - 1), b = t(t + 1), r = t^2 - 1$, and $\lambda = t^2 - t - 1$. Table 4 shows

Table 3 BIBD produced from a complete set of orthogonal 3×3 Latin squares ($v = 9, k = 3, b = 12, r = 4, \lambda = 1$)

Block number											
1	2	3	4	5	6	7	8	9	10	11	12
1	4	7	1	2	3	1	2	3	1	2	3
2	5	8	4	5	6	6	4	5	5	6	4
3	6	9	7	8	9	8	9	7	9	7	8

Table 4 Complementary BIBD produced from the BIBD in Table 3 ($v = 9, k = 6, b = 12, r = 8, \lambda = 5$)

Block number											
1	2	3	4	5	6	7	8	9	10	11	12
4	1	1	2	1	1	2	1	1	2	1	1
5	2	2	5	4	4	4	6	6	6	5	5
6	3	3	8	7	7	9	8	8	7	9	9
7	7	4	3	3	2	3	3	2	3	3	2
8	8	5	6	6	5	5	5	4	4	4	6
9	9	6	9	9	8	7	7	9	8	8	7

the complementary BIBD for the design given in Table 3.

Further BIBDs with $v = t^2 + t + 1$ and $k = t + 1$ may be produced from these orthogonal Latin square BIBDs by adding one of the additional $t + 1$ treatment symbols to each set of t blocks and including an additional block consisting of all $t + 1$ additional treatments. In this way, BIBDs with parameters $v = t^2 + t + 1, k = t + 1, b = t^2 + t + 1, r = t + 1$, and $\lambda = 1$ are derived. In each case, a complementary BIBD with parameters $v = t^2 + t + 1, k = t^2, b = t^2 + t + 1, r = t^2$, and $\lambda = t(t - 1)$ may also be constructed. Tables 5 and 6 give the extension of the BIBD in Table 3 and its complementary design respectively. Tables of BIBDs, produced this way and using other methods of construction, are available in Fisher and Yates Statistical Tables [9]. An easy to use catalog of designs, including BIBDs, is available at www.designtheory.org. See Khare and Federer [10] for an alternative method of constructing resolvable BIBDs.

Analysis of a General Incomplete Block Experiment

The usual method of analyzing an incomplete block experiment is to use the intrablock analysis based

Table 5 BIBD produced by extending Table 3 ($v = 13, k = 4, b = 13, r = 4, \lambda = 1$)

Block number												
1	2	3	4	5	6	7	8	9	10	11	12	13
1	4	7	1	2	3	1	2	3	1	2	3	10
2	5	8	4	5	6	6	4	5	5	6	4	11
3	6	9	7	8	9	8	9	7	9	7	8	12
10	10	10	11	11	11	12	12	12	13	13	13	13

on comparisons of treatments made within the same block, with block differences eliminated. For a general IBD, suppose that n_{ij} is the number of times treatment i , for $i = 1, 2, \dots, v$, is applied in the j th block, $j = 1, 2, \dots, b$, and that the corresponding response is y_{ijm} , for $m = 1, 2, \dots, n_{ij}$. Suppose that the model for y_{ijm} is the additive linear form

$$y_{ijm} = \mu + \tau_i + \beta_j + \varepsilon_{ijm} \quad (3)$$

where μ is a general mean level, τ_i is the effect of the i th treatment, β_j is the effect of the j th block, and ε_{ijm} is an error term with mean zero and variance σ^2 independent of any other error term. This model is overparameterized so that the parameters are not estimable without some restrictions being imposed to make them well defined. However, the objective here is to estimate contrasts in the treatment parameters, $\tau_i - \tau_j$, together with their appropriate standard errors, so that the treatments may be compared. We shall see that such contrasts may be estimable even though the individual parameters are not uniquely estimable. Equation (3) can be expressed in matrix notation as

$$\mathbf{y} = \mathbf{W}\boldsymbol{\alpha} + \boldsymbol{\varepsilon} \quad (4)$$

where $\mathbf{y}$ and $\boldsymbol{\varepsilon}$ are $n \times 1$ vectors of responses and error terms respectively, $\boldsymbol{\alpha}$ is a $p \times 1$ vector of parameters (with $p = 1 + v + b$), and $\mathbf{W}$ is the $n \times p$ design matrix of 1's and 0's identifying the treatments and blocks. The conditions on $\boldsymbol{\varepsilon}$ are that $E(\boldsymbol{\varepsilon}) = \mathbf{0}$ and $V(\boldsymbol{\varepsilon}) = \sigma^2 \mathbf{I}$, where $\mathbf{I}$ is the identity matrix. The columns of the $\mathbf{W}$ matrix can be partitioned into three submatrices corresponding to the parameter μ , the treatment effect parameters τ_i , for $i = 1, 2, \dots, v$, and the block effect parameters

β_j , for $j = 1, 2, \dots, b$. The matrix $\mathbf{W}$ and the vector $\boldsymbol{\alpha}$ can be partitioned as $\mathbf{W} = (\mathbf{1}|\mathbf{X}|\mathbf{Z})$ and $\boldsymbol{\alpha} = (\mu|\boldsymbol{\tau}'|\boldsymbol{\beta}')$, where $\mathbf{1}$ is a column of all 1's and $\mathbf{X}$ and $\mathbf{Z}$ are the $(n \times v)$ and $(n \times b)$ design matrices for treatments and blocks respectively. An alternative form for the model (4) is

$$\mathbf{y} = \mathbf{1}\mu + \mathbf{X}\boldsymbol{\tau} + \mathbf{Z}\boldsymbol{\beta} + \boldsymbol{\varepsilon} \quad (5)$$

The matrix $\mathbf{X}$ is such that $\mathbf{X}'\mathbf{1} = \mathbf{r}$, where $\mathbf{r} = (r_1, r_2, \dots, r_v)'$ is the vector of replications of each treatment and $\mathbf{X}'\mathbf{X} = \mathbf{r}^\delta$ where $\mathbf{r}^\delta$ is a diagonal matrix whose diagonal elements are those of the vector $\mathbf{r}$. Similarly, if $\mathbf{k}$ is a vector whose j th element is k_j , the number of units in the j th block, then $\mathbf{Z}'\mathbf{1} = \mathbf{k}$ and $\mathbf{Z}'\mathbf{Z} = \mathbf{k}^\delta$. Now, $\mathbf{X}'\mathbf{Z} = \mathbf{N}$, a $(v \times b)$ matrix, called the *incidence matrix* of the design, with ij th element equal to n_{ij} . If $\mathbf{T}$ and $\mathbf{B}$ are vectors of treatment and block totals respectively and if G is the overall total of the responses, then $\mathbf{X}'\mathbf{y} = \mathbf{T}$, $\mathbf{Z}'\mathbf{y} = \mathbf{B}$, and $\mathbf{1}'\mathbf{y} = G$. The least-squares estimator of the parameter vector $\boldsymbol{\alpha}$ is obtained by minimizing the sum of squared errors with respect to $\boldsymbol{\alpha}$. The normal equations are $\mathbf{W}'\mathbf{W}\hat{\boldsymbol{\alpha}} = \mathbf{W}'\mathbf{y}$. Partitioning $\mathbf{W}$ and $\boldsymbol{\alpha}$ as described above gives

$$n\hat{\mu} + \mathbf{r}'\hat{\boldsymbol{\tau}} + \mathbf{k}'\hat{\boldsymbol{\beta}} = G \quad (6)$$

$$\mathbf{r}\hat{\mu} + \mathbf{r}^\delta\hat{\boldsymbol{\tau}} + \mathbf{N}\hat{\boldsymbol{\beta}} = \mathbf{T} \quad (7)$$

$$\mathbf{k}\hat{\mu} + \mathbf{N}'\hat{\boldsymbol{\tau}} + \mathbf{k}^\delta\hat{\boldsymbol{\beta}} = \mathbf{B} \quad (8)$$

These equations do not have a unique solution because the matrix $\mathbf{W}'\mathbf{W}$ has rank $\leq v + b - 1$. The problem of finding a solution to equations (6–8) can be simplified by eliminating $\hat{\mu}$ and $\hat{\boldsymbol{\beta}}$. The resulting reduced normal equations for the intrablock analysis are

$$\mathbf{A}\hat{\boldsymbol{\tau}} = \mathbf{Q} \quad (9)$$

where $\mathbf{Q} = \mathbf{T} - \mathbf{N}\mathbf{k}^{-\delta}\mathbf{B}$ is the vector of treatment totals adjusted for blocks, and $\mathbf{A} = \mathbf{r}^\delta - \mathbf{N}\mathbf{k}^{-\delta}\mathbf{N}'$ is the intrablock treatment information matrix of the design. Equation (9) has a solution $\hat{\boldsymbol{\tau}} = \boldsymbol{\Omega}\mathbf{Q}$ for any generalized inverse $\boldsymbol{\Omega}$ of $\mathbf{A}$, i.e. for any matrix $\boldsymbol{\Omega}$ such that $\mathbf{A}\boldsymbol{\Omega}\mathbf{A} = \mathbf{A}$. If the overall mean μ is estimated using the sample mean $\bar{y}$, the block parameters $\hat{\boldsymbol{\beta}}$ are $\hat{\boldsymbol{\beta}} = \mathbf{k}^{-\delta}\mathbf{B} - \mathbf{k}^{-\delta}\mathbf{N}'\hat{\boldsymbol{\tau}} - 1\bar{y}$. Different generalized inverses, $\boldsymbol{\Omega}$, give different estimates, $\hat{\boldsymbol{\tau}}$, of the specific treatment parameters, but, provided the design is connected, treatment differences are

Table 6 Complementary BIBD produced from the BIBD in Table 5 ($v = 13, k = 9, b = 13, r = 9, \lambda = 6$)

Block number												
1	2	3	4	5	6	7	8	9	10	11	12	13
4	1	1	2	1	1	2	1	1	2	1	1	1
5	2	2	3	3	2	3	3	2	3	3	2	2
6	3	3	5	4	4	4	5	4	4	4	5	3
7	7	4	6	6	5	5	6	6	6	5	6	4
8	8	5	8	7	7	7	8	8	7	8	7	5
9	9	6	9	9	8	9	8	9	8	9	9	6
11	11	11	10	10	10	10	10	10	10	10	10	7
12	12	12	12	12	12	11	11	11	11	11	11	8
13	13	13	13	13	13	13	13	13	12	12	12	9

4 Block Designs, Incomplete

uniquely estimable. A block design is *connected* if it is possible to compare any two treatments directly within a block, or indirectly across several blocks; otherwise it is disconnected. For example, the two-block design in four treatments, two treatments per block {1, 2} and {3, 4}, is disconnected. The condition for connectedness is $\text{rank}(\mathbf{A}) = v - 1$. Provided the design is connected, a set of s linear combinations of the treatment effects can be written as $\mathbf{C}\boldsymbol{\tau}$, where $\mathbf{C}$ is an $s \times v$ contrast matrix with each row holding the coefficients, which sum to zero, of a treatment comparison. In this case, $E(\mathbf{C}\hat{\boldsymbol{\tau}}) = \mathbf{C}\boldsymbol{\tau}$ and $V(\mathbf{C}\hat{\boldsymbol{\tau}}) = \mathbf{C}\boldsymbol{\Omega}\mathbf{C}'\sigma^2$.

Intrablock Analysis of Variance for Connected Designs

The further assumption that the error terms $\boldsymbol{\varepsilon}$ are multivariate normally distributed is required when carrying out significance tests or setting **confidence intervals**. The total **sum of squares**, $\mathbf{y}'\mathbf{y}$, for the general linear model in equation (5) can be partitioned into three components

$$\mathbf{y}'\mathbf{y} = \mathbf{B}'\mathbf{k}^{-\delta}\mathbf{B} + \hat{\boldsymbol{\tau}}'\mathbf{Q} + (\mathbf{y} - \mathbf{W}\hat{\boldsymbol{\alpha}})'(\mathbf{y} - \mathbf{W}\hat{\boldsymbol{\alpha}}) \quad (10)$$

where $\hat{\boldsymbol{\alpha}}$ is any solution of the normal equations. This decomposition is invariant to the choice of generalized inverse $\boldsymbol{\Omega}$. The first component is the block sum of squares, the second term is the extra sum of squares due to treatments adjusting for blocks, and the last term is the residual sum of squares. The analysis of variance (*see Analysis of Variance*) shown in Table 7 may be used to test the overall null hypothesis $H_0: \tau_1 = \tau_2 = \dots = \tau_v$, using $F = \hat{\boldsymbol{\tau}}'\mathbf{Q}/(v-1)/s^2$ as an $F_{v-1, n-b-v+1}$ variable, where s^2 is the residual mean square.

A specific set of contrasts $\mathbf{C}\boldsymbol{\tau} = \mathbf{0}$ may be tested using the sum of squares $(\mathbf{C}\hat{\boldsymbol{\tau}})'(\mathbf{C}\boldsymbol{\Omega}\mathbf{C}')^{-1}(\mathbf{C}\hat{\boldsymbol{\tau}})$ with **degrees of freedom** equal to the rank of $\mathbf{C}$.

Analysis of a Balanced Incomplete Block Design

This general analysis simplifies if the design is a BIBD, since in this case, $\mathbf{r} = r\mathbf{1}_v$, $\mathbf{k} = k\mathbf{1}_b$, $\mathbf{r}^\delta = r\mathbf{I}_v$, and $\mathbf{k}^\delta = k\mathbf{I}_b$, where $\mathbf{1}_a$ is a vector of a 1s and $\mathbf{I}_a$ is the $a \times a$ identity matrix. In this case, $\mathbf{A} = \mathbf{r}^\delta - \mathbf{N}\mathbf{k}^{-\delta}\mathbf{N}'$ becomes $\mathbf{A} = r\mathbf{I}_v - \frac{1}{k}\mathbf{N}\mathbf{N}'$, where $\mathbf{N}\mathbf{N}'$ is the *concurrency matrix* of the design with

$$\mathbf{N}\mathbf{N}' = \begin{pmatrix} r & \lambda & \dots & \lambda \\ \lambda & r & \dots & \lambda \\ \vdots & \vdots & \ddots & \vdots \\ \lambda & \lambda & \dots & r \end{pmatrix} = (r - \lambda)\mathbf{I}_v + \lambda\mathbf{J}_v \quad (11)$$

where $\mathbf{J}_v$ is a $v \times v$ matrix of 1's.

Therefore $\mathbf{A} = \frac{\lambda}{k}[\mathbf{v}\mathbf{I}_v - \mathbf{J}_v]$, with g inverse $\boldsymbol{\Omega} = \frac{k}{\lambda v}\mathbf{I}_v$, so that a solution of the normal equations is

$$\hat{\boldsymbol{\tau}} = \boldsymbol{\Omega}\mathbf{Q} = \frac{k}{\lambda v}\mathbf{Q} = \frac{k}{\lambda v}\left(\mathbf{T} - \frac{1}{k}\mathbf{N}\mathbf{B}\right) \quad (12)$$

For any contrast matrix, $\mathbf{C}$, $\mathbf{C}\hat{\boldsymbol{\tau}}$ is the least-squares estimator of $\mathbf{C}\boldsymbol{\tau}$ with

$$V(\mathbf{C}\hat{\boldsymbol{\tau}}) = \mathbf{C}\boldsymbol{\Omega}\mathbf{C}'\sigma^2 = \frac{k}{\lambda v}\mathbf{C}\mathbf{C}'\sigma^2 \quad (13)$$

The difference between two treatment effects τ_{i_1} and τ_{i_2} ($i_1 \neq i_2$) is estimated by $\mathbf{l}'\hat{\boldsymbol{\tau}}$ with $\mathbf{l}' = (0, \dots, 1, \dots, -1, \dots, 0)$, that is $\hat{\tau}_{i_1} - \hat{\tau}_{i_2}$ with $\text{var}(\hat{\tau}_{i_1} - \hat{\tau}_{i_2}) = \frac{k}{\lambda v}\mathbf{l}'\mathbf{l}\sigma^2 = \frac{2k}{\lambda v}\sigma^2$. This is the same regardless of which two treatments are being compared, a consequence of the balance of the design. For a complete block design with r replicates of each treatment, the variance of the estimated treatment difference is $\text{var}(\hat{\tau}_{i_1} - \hat{\tau}_{i_2}) = 2\sigma^2/r$, giving the efficiency of a BIBD as $\lambda v/(rk)$.

An Example of the Analysis of an Incomplete Block Design

The reduced BIBD with $v = 6$ treatments in $b = 10$ blocks of $k = 3$ treatments per block, shown in

Table 7 Intrablock analysis of variance for the general incomplete block experiment

Source of variation	Degrees of freedom	Sum of squares	Mean squares
Treatments eliminating blocks	$v - 1$	$\hat{\boldsymbol{\tau}}'\mathbf{Q}$	$\hat{\boldsymbol{\tau}}'\mathbf{Q}/(v - 1)$
Blocks ignoring treatments	$b - 1$	$\mathbf{B}'\mathbf{k}^{-\delta}\mathbf{B} - G^2/n$	—
Residual	$n - b - v + 1$	by subtraction	s^2
Total (after fitting mean)	$n - 1$	$\mathbf{y}'\mathbf{y} - G^2/n$	—

Table 8 BIB design, shown in bold, and taste data for the comparison of six antibiotic formulations with each subject receiving three treatments

Subject (block)	1	2	3	4	5	6	7	8	9	10
	1	5	1	3	1	2	1	5	1	3
	2	5	2	3	4	4	5	5	3	3
	3	5	4	5	1	6	5	6	4	6
Total	15	10	7	15	12	5	8	9	10	9

Table 8, was used as the basic design in a larger comparative study of antibiotic formulations for children, described by Prescott *et al.* [11]. Table 8 also shows, for the first 10 children, the responses relating to taste, as indexed by the selection of an appropriate facial expression ranging from a sad face (1) to a smiley face (5). These data are used here as a simple illustration of the calculations involved in an analysis of a BIBD, but for the full data set nonparametric ranking methods were employed.

For this design, $r = 5$ and $\lambda = 2$, so that the incidence matrix, $\mathbf{NN}'$, the intrablock treatment information matrix, $\mathbf{A}$, and a possible generalized inverse, $\mathbf{\Omega}$, are given by $\mathbf{NN}' = (r - \lambda)\mathbf{I}_v + \lambda\mathbf{J}_v = 3\mathbf{I}_6 + 2\mathbf{J}_6$, $\mathbf{A} = \frac{\lambda}{k}[\nu\mathbf{I}_v - \mathbf{J}_v] = \frac{2}{3}[6\mathbf{I}_6 - \mathbf{J}_6]$, and $\mathbf{\Omega} = \frac{k}{\lambda\nu}\mathbf{I}_v = \frac{1}{4}\mathbf{I}_6$, so that $\hat{\tau} = \mathbf{\Omega Q} = \frac{k}{\lambda\nu}\mathbf{Q} = \frac{1}{4}(\mathbf{T} - \frac{1}{3}\mathbf{NB})$. Since the treatment and block totals are $\mathbf{T}' = (18, 13, 19, 23, 10, 17)$ and $\mathbf{B}' = (15, 10, 7, 15, 12, 5, 8, 9, 10, 9)$ respectively, this particular solution for $\hat{\tau}$ is given by $\hat{\tau}' = \frac{1}{12}(-5, -8, 11, 17, -16, 1)$. The unadjusted and adjusted treatment means are shown in Table 9. Table 10 gives the analysis of variance from which

it is evident that there are significant differences among the treatments. The variance of any pairwise treatment difference is $\text{var}(\hat{\tau}_{i_1} - \hat{\tau}_{i_2}) = \frac{2k}{\lambda\nu}\sigma^2 = \sigma^2/2$, which is estimated using $s^2/2 = 0.956/2 = 0.478$ with 15 degrees of freedom. Any pair of treatment estimates which differ by more than $0.478^{1/2} \times 2.13 = 1.47$ are significantly different. This suggests that antibiotic formulations 3 and 4 are significantly different from formulations 2 and 5 but not significantly different from formulations 1 and 6.

Further Developments and Related Topics

Kiefer [12] and John and Mitchell [13] discuss optimality (see **Optimal Design**) of IBDs when v divides n , based on minimization of the E-criterion that examines the largest variance over all (normalized) treatment contrasts. The former proposed families of criteria and provided some rigorous proofs, while the latter produced examples of designs satisfying the optimality conjecture. Further details are given in John and Williams [14]. Computer software packages are now available for the construction of optimal or near-optimal block designs, e.g. Nguyen [15].

It is possible to extend the analysis of BIBDs to include recovery of the interblock information on the treatments by considering the block effects to be random effects. Details of this extended analysis and methods of combining the intra- and interblock information are given in John [16, pp. 233–248], see also Federer [17].

Table 9 Unadjusted and adjusted treatment means

Treatment	1	2	3	4	5	6
Means (unadjusted)	3.6	2.6	3.8	4.6	2.0	3.4
Means (adjusted)	−0.417	−0.667	0.917	1.417	−1.333	0.0833

Table 10 Intrablock analysis of variance for the antibiotics example

Source of variation	Degrees of freedom	Sum of squares	Mean squares	F value	P value
Treatments eliminating blocks	5	21.00	4.200	4.40	0.012
Blocks ignoring treatments	9	31.33	3.481	3.64	0.013
Residual	15	14.33	0.956	—	—
Total (after fitting mean)	29	66.67	—	—	—

References

- [1] Yates, F. (1936). A new method of arranging variety trials involving a large number of varieties, *Journal of Agricultural Science* **26**, 424–455.
- [2] Yates, F. (1936). Incomplete randomised blocks, *Annals of Eugenics* **7**, 121–140.
- [3] Bose, R.C. (1939). On the construction of balanced incomplete block designs, *Annals of Eugenics* **9**, 353–399.
- [4] Bose, R.C. (1942). A note on the resolvability of balanced incomplete block designs, *Sankhya* **6**, 105–110.
- [5] Fisher, R.A. (1940). An examination of the different possible solutions of a problem in incomplete blocks, *Annals of Eugenics* **10**, 52–75.
- [6] Bose, R.C. & Nair, K.R. (1939). Partially balanced incomplete block designs, *Sankhya* **4**, 337–372.
- [7] Youden, W.J. (1940). Experimental designs to increase accuracy of greenhouse studies, *Contributions from Boyce Thompson Institute* **11**, 219–228.
- [8] Shrikhande, S.S. (1950). The impossibility of certain symmetrical balanced incomplete block designs, *Annals of Mathematical Statistics* **21**, 106–111.
- [9] Fisher, R.A. & Yates, F. (1963). *Statistical Tables*, Oliver and Boyd, Edinburgh.
- [10] Khare, M. & Federer, W.T. (1981). A simple construction procedure for resolvable incomplete block designs for any number of treatments, *Biometrical Journal* **23**, 121–132.
- [11] Prescott, P., Ryan, J. & Shuttleworth, P. (1994). A comparative study of several antibiotic formulations using a design based on a combination of balanced incomplete blocks and Latin squares, *Statistics in Medicine* **13**, 11–21.
- [12] Kiefer, J. (1975). Construction and optimality of generalized Youden designs, in *A Survey of Statistical Design and Linear Models*, J.N. Srivastava, ed, North-Holland, Amsterdam, pp. 333–353.
- [13] John, J.A. & Mitchell, T.J. (1977). Optimal incomplete block designs, *Journal of the Royal Statistical Society. Series B* **39**, 39–43.
- [14] John, J.A. & Williams, E.R. (1995). *Cyclic and Computer Generated Designs*, 2nd Edition, *Monographs on Statistics and Applied Probability*, Chapman & Hall, London, Vol. 38.
- [15] Nguyen, N.-K. (1993). An algorithm for constructing optimal resolvable incomplete block designs, *Communications in Statistics-Simulation and Computation* **22**, 911–923.
- [16] John, P.W.M. (1971). *Statistical Design and Analysis of Experiments*, Macmillan, New York.
- [17] Federer, W.T. (1998). Recovery of interblock, intergradient, and intervary information in incomplete block and lattice rectangular designed experiments, *Biometrics* **54**, 471–481.

PHILIP PRESCOTT

Interactions

Interaction

Interaction refers to the joint nature of the effect of two or more explanatory variables on a response variable. When the interaction involves two factors, it is referred to as a *two-way interaction*, three factors suggest a three-way interaction, and so forth. The existence of interaction effects means that the effect of one factor on the outcome depends on the levels of another factor. For example, if factors *A* and *B* are said to exhibit interaction, the effect of *A* on the response variable cannot be completely expressed unless the level of *B* is first specified. Interactions are common in many applications of experimental practice. For a discussion of the importance of interactions in experimental practice, Box [1] provides a useful illustration of interactions in both agricultural and engineering contexts. Cox [2] provides a comprehensive review and formal treatment of the definition, detection, and interpretation of interactions. Here we restrict our attention to interactions between two treatment variables. For further reading on interactions among various subsets of random (nonspecific) factors and intrinsic factors, see Cox [2].

Methodology

For concreteness, suppose that a **factorial** experiment yields a response Y_{ijk} measured on experimental unit k at level $i = 1, \dots, I$ of factor *A* and at level $j = 1, \dots, J$ of factor *B*. Further, suppose that n_{ij} experimental units are measured with *A* at level i and *B* at level j (so that $k = 1, \dots, n_{ij}$). One possible representation of a model that relates the response variable to the two explanatory variables is

$$Y_{ijk} = \mu_{ij} + \epsilon_{ijk} \quad (1)$$

where the μ_{ij} values are the *cell means* and the error terms, ϵ_{ijk} are independent and identically distributed (i.i.d.) $N(0, \sigma^2)$. This representation is convenient and main effects (*see Main Effects*) and interactions for the parameters can be expressed as contrasts of

the cell means. For example, suppose factor *A* has $I = 2$ levels and factor *B* has $J = 2$ levels. Then the cell means can be represented in a two-way layout, as in Table 1.

To define interactions in terms of cell means, we first begin with the definition of simple effects. Simple effects are defined as the contrasts between different levels of one factor at a single level of a second factor. If a factor only has two levels, this contrast becomes the difference between the two cell means. In the case above, an example of a simple effect is

$$(\mu_{11} - \mu_{12})$$

Interactions, then, are differences between simple effects. In the simple setup above, the interaction between *A* and *B* is defined as

$$(\mu_{11} - \mu_{12}) - (\mu_{21} - \mu_{22})$$

In more complex situations (situations with multiple levels on factor *A* or factor *B* or additional factors) the interaction can still be expressed as a difference of simple effects, but there are many different ways to express the simple effects as contrasts.

A second model for relating the response Y_{ijk} to factors *A* and *B* is the *effects model*, which can be expressed in the interaction form as

$$Y_{ijk} = \mu + \alpha_i + \beta_j + \gamma_{ij} + \epsilon_{ijk} \quad (2)$$

where γ_{ij} describes the interaction effect of level i of factor *A* and level j of factor *B* on the response Y_{ijk} . Implicit in the interaction model is the notion that the effect of *A* on the response variable is differentially expressed depending on the levels of factor *B*. The model as stated here is not identifiable, however, the following constraints are often used as a means of

Table 1 Two-way layout of cell means

		Factor A	
		1	2
Factor B	1	μ_{11}	μ_{12}
	2	μ_{21}	μ_{22}

2 Interactions

making the model identifiable,

$$\sum_{i=1}^I \alpha_i = 0 \quad (3)$$

$$\sum_{j=1}^J \beta_j = 0 \quad (4)$$

$$\sum_{i=1}^I \gamma_{ij}, \text{ for all } j \text{ and} \quad (5)$$

$$\sum_{j=1}^J \gamma_{ij}, \text{ for all } i \quad (6)$$

For a discussion on the identifiability of this model and the restrictions required to make this model identifiable, *see Analysis of Variance*. The effects model is the basis for most statistical software packages, and is therefore the most common choice for practitioners.

Designs for Detection of Interactions

Factorial designs (*see Factorial Experiments*) are commonly used as a design technique that allows assessment of the presence or absence of interaction

effects. The full factorial design allows for **estimation** of main effects and interactions between factors under consideration. For experimental problems in which the number of factor levels or the number of factors is high, a full factorial design is often infeasible due to the large number of experimental runs required.

In cases when full factorial designs are infeasible, fractional factorial designs (*see Fractional Factorial Designs*) may still provide insight into some interaction effects. In particular, resolution-IV and higher-resolution (*see Factorial Designs, Resolution of*) designs often have the ability to clearly estimate a subset of all possible two-factor interactions. If the design cannot distinguish between certain interactions, the interactions are said to be aliased. For a complete discussion on the aliasing of effects and the relation of aliasing to fractional factorial designs and interactions, *see Aliasing in Fractional Designs*.

Detection of Interaction Effects

An effective tool used to determine the presence or absence of a two-factor interaction is a plot of the estimated cell means, $\hat{\mu}_{ij}$, *versus* the levels of factor A. Further, the cell means associated with the same level of factor B are connected by lines of the

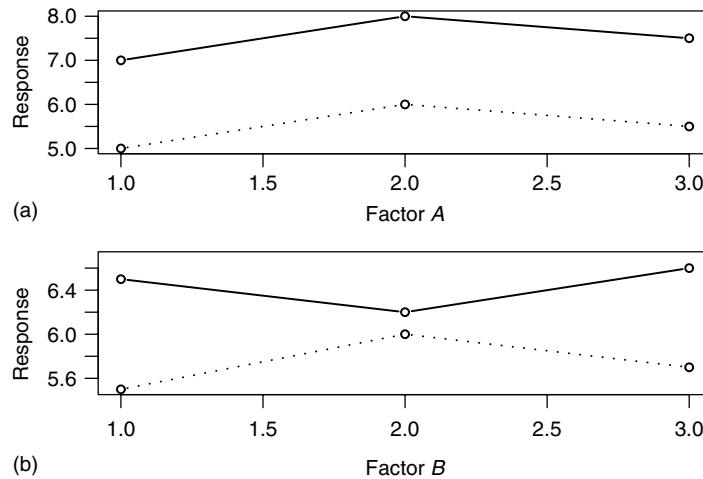


Figure 1 Two notional interaction plots. The solid lines indicate the first level of Factor B, the dotted lines indicate the second level of factor B. (a) Lines which are parallel, suggesting the absence of an interaction between factors A and B. (b) Lines which are not parallel, suggesting the presence of an interaction between factors A and B

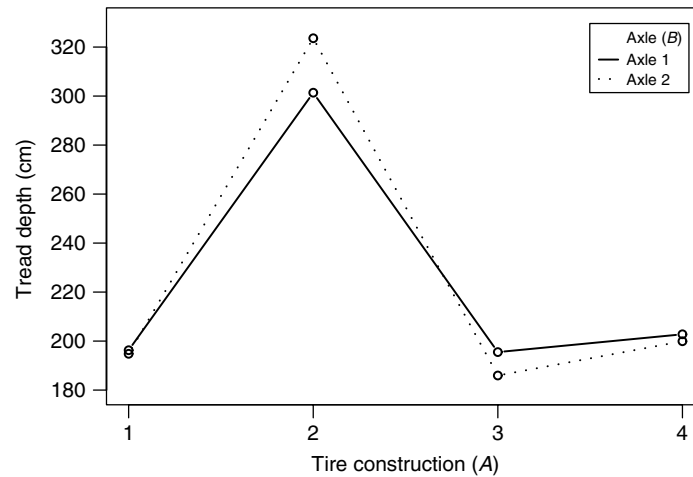


Figure 2 Interaction plot for the tread wear experiment. Depicted are the cell means for each treatment combination. The solid line represents axle configuration 1 and the dotted line represents axle configuration 2

same type. This organization of cell means by factor combinations is often referred to as an *interaction plot*. Figure 1 is a notional depiction of an interaction plot where factor *A* has three levels and factor *B* has two levels. Two possible scenarios are depicted, one in which *A* and *B* exhibit interaction, and one in which *A* and *B* exhibit no interaction. In each plot, the different line types indicate the different levels of factor *B*. Figure 1(a) reveals that the two different lines are not parallel, suggesting that factor *A* interacts with factor *B* in determining the effect on the response. More concisely stated, Figure 1 indicates an interaction between factors *A* and *B*. Figure 1(b) shows the two lines for factor *B* to be parallel, which suggests that the effect of factor *A* is consistent, regardless of the level of factor *B* (and *vice versa*).

The interaction graph provides a powerful and valuable tool for addressing the presence or absence of interaction. However, many scientists and engineers wish to validate the visual determination with a more formal assessment of the presence or absence of interactions. Using analysis of variance (ANOVA) (see **Analysis of Variance**), *F*-tests can be constructed to test the presence of an interaction between two factors. The *F*-test makes it possible to check whether an observed interaction is so large that it could not have occurred by chance. If the *F*-test is not statistically significant, there is not strong evidence against the assumption that the factors do affect the

response independently. This requires a reinterpretation of the main effects. Reinterpretation of the main effects tests is often cumbersome and many authors suggest that the main effects have no real meaning [3]. Other authors [4, 5] have suggested that interactions can be framed in terms of conditional main effects, which is the interpretation required when a significant interaction effect exists.

Example

A commercial truck tire manufacturer proposes an experimental design approach to the evaluation of tire wear for four different tire constructions. In addition, the tire manufacturer is concerned about the effect of axle configurations on tire wear. In order to address the interaction between tire construction (*A*,

Table 2 Data from the two-factor tread wear experiment

A	B	Replicate 1	Replicate 2	Replicate 3	Replicate 4
1	1	193	201	207	185
2	1	312	284	301	308
3	1	192	205	186	199
4	1	206	203	206	196
1	2	204	186	188	202
2	2	322	331	322	319
3	2	190	183	184	186
4	2	198	205	206	190

4 Interactions

Table 3 ANOVA table for the tread wear experiment

	$df^{(a)}$	$SS^{(b)}$	$MS^{(c)}$	F value	$Pr(>F)$
A	3	82026	27342	428.5817	0.0000
B	1	35	35	0.5516	0.4649
Interaction	3	1157	386	6.0478	0.0032
Error	24	1531	64	–	–

^(a)df: degrees of freedom.

^(b)SS: sum of squares.

^(c)MS: mean square.

with four levels) and axle configuration (B , with two levels), a full factorial design was constructed with four randomly selected tires of each tire construction assigned to each of the two axle configurations. The tread depth as measured in centimeters for each of the four tires was recorded after 25 000 mi. The data are shown in Table 2.

A graphical display of the cell means is depicted in Figure 2. The horizontal axis contains the different tire construction types (1, 2, 3, or 4). The vertical axis is the tread depth measurement in centimeters. The cell means for each of the two different axle configurations are connected by lines, where the solid line represents the cell means for axle configuration 1 and the broken line represents the cell means for axle configuration 2. Because these two lines are not parallel, suspicion of a potential interaction might be raised. In particular, it appears that while the measured tread is deeper (larger values of tread depth) for tire construction 2 at axle configuration 2, the opposite effect is observed for tire construction 3. That is, for tire construction 3, axle configuration 1 appears to provide a deeper tread measurement than axle configuration 2. It appears that tire construction and axle configuration interact with each other.

To formalize this visual inspection of the interaction effect, the ANOVA table for this experiment is shown in Table 3. The p value for the interaction effect is small (0.0032), confirming the visual inspection of Figure 2. While the **main effect** p value is also small, caution should be taken when interpreting a main effect in the presence of a significant interaction. However, further analysis of Figure 2 reveals that this small p value associated with the main effect of tire construction is most likely driven by the significantly different tread depths for tire construction 2.

Higher-Order Interactions

In an investigational study, it is often of interest to examine the effect of three or more factors on a response variable. This may be particularly true in an early phase of experimentation where there is an understanding of the many factors which *potentially* relate to a response variable, but the exact nature of those relationships is unclear. In this situation, higher-order interactions (three factor, four factor, etc.) may exist. These higher-order interactions are fairly simple to detect using an F -test. However, while these effects may exist, in practice it is often difficult to provide a tractable interpretation of interaction effects above three-way interactions. Visualization of four-factor and higher-order interactions using interaction plots is infeasible. In general, interactions between more than three factors are difficult to interpret and the associated **degrees of freedom** from such effects may be more effectively used to more precisely estimate the error. This pooling of higher-order interactions is discussed extensively in Hocking [6], Dean and Voss [7], and Box *et al.* [8].

References

- [1] Box, G.E.P. (1990). Do interactions matter? *Quality Engineering* **2**, 365–369.
- [2] Cox, D.R. (1984). “Interaction” (with discussion), *International Statistical Review* **52**, 1–29.
- [3] Petersen, R.G. (1985). *Design and Analysis of Experiments*, Marcel Dekker, New York.
- [4] Wu, C.F.J. & Hamada, M. (2000). *Experiments: Planning, Analysis, and Parameter Design Optimization*, John Wiley & Sons, New York.
- [5] Kuehl, R.O. (1994). *Statistical Principles of Research Design and Analysis*, Duxbury Press, Belmont.
- [6] Hocking, R.R. (1996). *Methods and Applications of Linear Models*, John Wiley & Sons, New York.

- [7] Dean, A. & Voss, D. (1999). *Design and Analysis of Experiments*, Springer, New York.
- [8] Box, G.E.P., Hunter, J.S. & Hunter, W.G. (2005). *Statistics for Experimenters: Design, Innovation, and Discovery*, 2nd Edition, John Wiley & Sons, New York.

Further Reading

Box, G.E.P. & Cox, D.R. (1964). An analysis of transformations, *Journal of the Royal Statistical Society. Series B* **26**, 211–246.

Related Articles

Aliasing in Fractional Designs; Analysis of Variance; Blocking; Factorial Designs, Resolution of; Factorial Experiments; Fractional Factorial Designs; Main Effects.

C. SHANE REESE

Latin Hypercube Designs

Introduction

Latin hypercubes were introduced by McKay *et al.* [1] in 1979, and provide useful designs for **computer experiments** and numerical integration. Evans and Swartz [2] offer a detailed discussion on various numerical integration methods. Santner *et al.* [3] provide a careful account of the design and analysis of computer experiments (*see Computer Experiments*). A recent book on computer experiments is that of Fang *et al.* [4].

Latin hypercubes are easy to generate, making them especially attractive for experts who do not specialize in design theory. Latin hypercubes achieve the maximum uniformity in each of the univariate margins of the design region, thus allowing the experimenter to use models that are capable of capturing the complex dependence of the response variable on the input variables. Another reason that contributes to the popularity of Latin hypercubes is that they have no repeated runs. Repeated runs do not bring about extra information in computer experiments as running a deterministic computer code twice yields identical output.

Latin hypercubes are a very large class of designs. A design may not perform well in terms of other criteria such as those of orthogonality or space filling, simply because it is a Latin hypercube. A design that spreads its points evenly in the design region is referred to as a *space-filling design*; such a design allows information to be collected from all portions of the design region. Many authors have contributed to this area of selecting good Latin hypercubes, and some of the ideas will be briefly described in this article.

Definition and Basic Properties

An $n \times m$ matrix $D = (d_{ij})$ is called a *Latin hypercube* of n runs for m factors if each column of D is a permutation of $1, \dots, n$. What makes Latin hypercubes distinctly different from other designs is that every factor in a Latin hypercube has the same number of levels as the run size. Because of this property, Latin hypercubes might be totally impractical for a field or laboratory experiment but are easy to

implement in the context of a computer experiment or a numerical integration.

Let $y = f(x_1, \dots, x_m)$ be a real-valued function with m variables. For convenience, all the variables are now coded to the unit interval $[0, 1]$, and therefore the design region is $[0, 1]^m$. The design region will be allowed to take a more general form when we discuss Latin hypercube sampling later in this section. The function represents the deterministic computer model in the case of computer experiments or the integrand in the case of numerical integration. There are two natural ways of generating a *Latin hypercube design* based on a given Latin hypercube. The first is through

$$x_{ij} = \frac{(d_{ij} - 0.5)}{n} \quad (1)$$

with the n points given by $(x_{i1}, \dots, x_{im})$ with $i = 1, \dots, n$. The other is through

$$x_{ij} = \frac{(d_{ij} - u_{ij})}{n} \quad (2)$$

with the n points given by $(x_{i1}, \dots, x_{im})$ with $i = 1, \dots, n$, where u_{ij} are independent random variables with a common uniform distribution on $(0, 1]$. The difference between the two methods can be seen as follows. When projected onto each of the m variables, both methods have the property that one and only one of the n design points fall within each of the n small intervals defined by $[0, 1/n), [1/n, 2/n), \dots, [(n-1)/n, 1]$. The first method gives the midpoints of these intervals while the second gives the points that are uniformly distributed in their corresponding intervals.

Example 1 The following are two Latin hypercubes of $n = 5$ runs for $m = 2$ factors.

D_1		D_2	
1	3	1	2
2	2	2	5
3	1	3	3
4	4	4	1
5	5	5	4

Their graphical representations are given in Figure 1. Although they are both Latin hypercubes, design D_2 provides a better coverage of the design region than design D_1 . This raises the need of

2 Latin Hypercube Designs

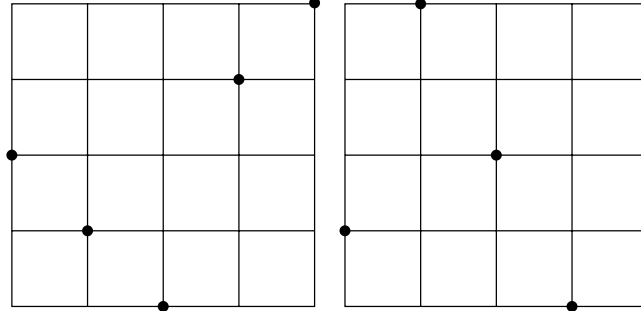


Figure 1 The graphical representations of the two Latin hypercubes in Example 1

developing methods for selecting better Latin hypercubes. We will return to this issue in the section titled “Selecting Good Latin Hypercubes”.

Latin hypercubes are very easy to generate. Simply combining several permutations of $1, \dots, n$, one obtains a Latin hypercube. There is no restriction whatsoever on the run size n and the number m of factors. Since a Latin hypercube has n distinct levels in each of its factors, it achieves the maximum uniformity in each univariate margin. Two useful properties follow from this simple fact. Because of the maximum number of levels, a Latin hypercube presents the experimenter with the opportunity of modeling the complex dependence of the response variable on each of the input variables. Because there are no repeated levels in each factor, a Latin hypercube necessarily has no repeated runs. In computer experiments, a physical mechanism is modeled using a complex deterministic function which is implemented *via* a computer code. Running the computer code twice at the same setting of input variables produces the same output, and therefore using repeated runs in computer experiments is a waste of resources.

By definition, a Latin hypercube does not promise any property in two- or higher-dimensional margins, as clearly demonstrated in Example 1. It is therefore up to the user to find the “right permutations” so that the resulting design has certain desirable properties in two or higher dimensions. One simple strategy is to use a random Latin hypercube in which the permutations are selected randomly. This helps eliminate the possible systematic patterns in the resulting design but there is no guarantee that the design will perform well in terms of other useful design criteria. The next section will look at some of the useful ideas for finding “good” Latin

hypercubes.

Latin hypercube sampling is a sampling method based on Latin hypercube designs. When some of the input factors are environmental variables and therefore follow some continuous distributions, it is desirable that the selected inputs mimic these distributions. The methods as given in equations (1) and (2) are well suited if these environmental variables follow a uniform distribution. If the j th factor follows a continuous distribution with a cumulative distribution function $F_j(x)$, then the inputs for that factor can be selected *via* the familiar quantile transformation $F_j^{-1}(x_{ij})$, where x_{ij} is given by either equation (1) or (2). The design region in this case depends on $F_j(x)$ and takes a more general form than $[0, 1]^m$. In addition to matching the marginal distributions, effort can be made to induce desired correlations among the input variables [5].

When used in numerical integration, Latin hypercube sampling achieves a variance reduction compared with the simplest Monte Carlo integration. The problem here is to find $\mu = \int_{[0,1]^m} f(x_1, \dots, x_m) dx_1 \cdots dx_m$. Let $X_1, \dots, X_m$ be independent uniform random variables on $[0, 1]$. Then $\mu = E(Y)$, where $Y = f(X_1, \dots, X_m)$. The simplest Monte Carlo method is to use $\bar{Y} = n^{-1} \sum_{i=1}^n Y_i$ to approximate μ , where $Y_i = f(X_{i1}, \dots, X_{im})$ with X_{ij} 's being mn independent uniform random variables on $[0, 1]$. Clearly, $E(\bar{Y}) = \mu$ and $\text{Var}(\bar{Y}) = n^{-1} \text{Var}(Y)$. Latin hypercube sampling also uses $\bar{Y} = n^{-1} \sum_{i=1}^n Y_i$ to estimate μ , where $Y_i = f(X_{i1}, \dots, X_{im})$ but this time X_{ij} is given by $X_{ij} = n^{-1}(d_{ij} - U_{ij})$ with $D = (d_{ij})$ being a random Latin hypercube and U_{ij} 's are independent uniform random variables on $[0, 1]$. Stein [6] showed that under Latin hypercube sampling, the estimator $\bar{Y}$ is still unbiased and has a

variance given by

$$\text{Var}(\bar{Y}) = n^{-1}\text{Var}(Z) + o(n^{-1}) \quad (3)$$

where $Z = Y - \mu - \sum_{j=1}^m [E(Y|X_j) - \mu]$. Since $\text{Var}(Y) = \sum_{j=1}^m \text{Var}[E(Y|X_j)] + \text{Var}(Z)$, Latin hypercube sampling has a smaller variance than the simplest Monte Carlo.

Selecting Good Latin Hypercubes

Latin hypercubes are a very rich class of designs. Although all Latin hypercubes enjoy an attractive property in terms of uniformity of the projected design points on each univariate margin, some Latin hypercubes have better properties in higher dimensions while others do not. It is therefore highly desirable to develop methods for selecting Latin hypercubes that perform well in terms of other criteria such as those of orthogonality or space filling. We briefly discuss two sets of ideas that are useful for this purpose. The first set of ideas aims at Latin hypercubes with better space-filling properties. The second set of ideas focuses on finding orthogonal or nearly orthogonal Latin hypercubes.

Space-Filling Latin Hypercubes

A Latin hypercube will provide a good coverage of the design region if all the points are farther apart, that is, no two points are too close to each other. This idea can be formally developed using the maximin distance criterion, according to which designs should be selected by maximizing $\min_{i \neq j} d(p_i, p_j)$, where $d(p_i, p_j)$ denotes the distance between design points p_i and p_j . Euclidean distance is commonly used but other distance measures are also useful. See Johnson *et al.* [7] for a general discussion on maximin and related minimax distance designs. The use of the maximin distance criterion for selecting Latin hypercubes was thoroughly explored by Morris and Mitchell [8].

Orthogonal arrays (see **Orthogonal Arrays**) are very useful designs for physical factorial experiments with and without blocking. A comprehensive treatment of orthogonal arrays is given by Hedayat *et al.* [9]. When used to construct Latin hypercubes, orthogonal arrays give rise to a class of good Latin hypercubes, namely orthogonal array

(OA)-based Latin hypercubes. An OA-based Latin hypercube is a Latin hypercube that becomes an orthogonal array after the levels are properly grouped [10]. To be more specific, an OA-based Latin hypercube $D = (d_{ij})$ of n runs has the property that $A = (a_{ij})$ where $a_{ij} = g(d_{ij})$ is an s -level orthogonal array of $n = ks$ runs, where $g(x) = j$ for $(j-1)k + 1 \leq x \leq jk$. OA-based Latin hypercubes inherit the t -dimensional space-filling property of orthogonal arrays of strength t . To illustrate, we first generate a random Latin hypercube with $n = 25$ runs and $m = 3$ factors, and then generate a random OA-based Latin hypercube also with $n = 25$ runs and $m = 3$ factors, using a five-level orthogonal array of strength 2. To save space, we omit the design matrices, and provide only the pairwise plots of the two designs in Figure 2. The OA-based Latin hypercube as in Figure 2(b) clearly outperforms the Latin hypercube in Figure 2(a) in terms of filling two-dimensional projection spaces.

Orthogonal Latin Hypercubes

A Latin hypercube is said to be orthogonal if the correlation between any two of its columns is zero and nearly orthogonal if the pairwise correlations are all small. The idea of considering nearly orthogonal Latin hypercubes dates back to Iman and Conover [5], who proposed a method for finding nearly orthogonal Latin hypercubes using Cholesky decomposition. Owen [11] developed a more effective algorithm for this purpose on the basis of the Gram–Schmidt orthogonalization. Tang [12] extended Owen’s algorithm to the situations where the correlations of higher orders are to be minimized.

Latin hypercubes that are exactly orthogonal were first formally studied by Ye [13], who constructed a class of orthogonal Latin hypercubes with $n = 2^k$ runs and $m = 2k - 2$ factors. For example, when $k = 4$, his construction gives an orthogonal Latin hypercube of 16 runs for six factors. Note that the ratio m/n becomes very small even for moderately large k . The orthogonal Latin hypercubes constructed by Steinberg and Lin [14] have a large ratio m/n ; the price is the more severe restriction on the run size n , which now has the form $n = 2^{2^k}$, where k is an integer. Related work in this connection is that of Butler [15], who constructed Latin hypercubes that are orthogonal with respect to models based on trigonometric functions.

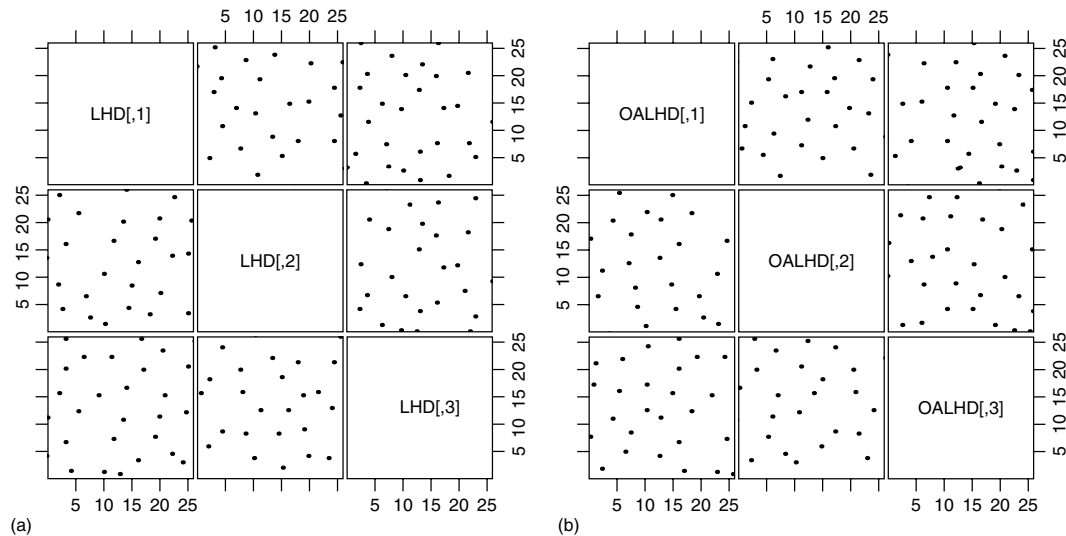


Figure 2 (a) Pairwise plots of a random Latin hypercube of 25 runs for three factors, and (b) corresponds to a random OA-based Latin hypercube of the same sizes. The underlying orthogonal array for the OA-based Latin hypercube has 25 runs, five levels, and three factors, and is of strength 2

Acknowledgments

The research is supported by the Natural Sciences and Engineering Research Council of Canada.

References

- [1] McKay, M.D., Beckman, R.J. & Conover, W.J. (1979). A comparison of three methods for selecting values of input variables in the analysis of output from a computer code, *Technometrics* **21**, 239–245.
- [2] Evans, M. & Swartz, T.B. (2000). *Approximating Integrals via Monte Carlo and Deterministic Methods*, Oxford University Press, New York.
- [3] Santner, T.J., Williams, B.J. & Notz, W.I. (2003). *The Design and Analysis of Computer Experiments*, Springer, New York.
- [4] Fang, K.T., Li, R. & Sudjianto, A. (2006). *Design and Modeling for Computer Experiments*, Chapman & Hall, New York.
- [5] Iman, R.L. & Conover, W.J. (1982). A distribution-free approach to inducing rank correlation among input variables, *Communications in Statistics Part B: Simulation and Computation* **11**, 311–334.
- [6] Stein, M. (1987). Large-sample properties of simulations using Latin hypercube sampling, *Technometrics* **29**, 143–151.
- [7] Johnson, M.E., Moore, L.M. & Ylvisaker, D. (1990). Minimax and maximin distance designs, *Journal of Statistical Planning and Inference* **26**, 131–148.
- [8] Morris, M.D. & Mitchell, T.J. (1995). Exploratory designs for computational experiments, *Journal of Statistical Planning and Inference* **43**, 381–402.
- [9] Hedayat, A.S., Sloane, N.J.A. & Stufken, J. (1999). *Orthogonal Arrays: Theory and Applications*, Springer, New York.
- [10] Tang, B. (1993). Orthogonal array based Latin hypercubes, *Journal of the American Statistical Association* **88**, 1392–1397.
- [11] Owen, A.B. (1994). Controlling correlations in Latin hypercube samples, *Journal of the American Statistical Association* **89**, 1517–1522.
- [12] Tang, B. (1998). Selecting Latin hypercubes using correlation criteria, *Statistica Sinica* **8**, 965–977.
- [13] Ye, K.Q. (1998). Orthogonal column Latin hypercubes and their application in computer experiments, *Journal of the American Statistical Association* **93**, 1430–1439.
- [14] Steinberg, D.M. & Lin, D.K.J. (2006). A construction method for orthogonal Latin hypercube designs, *Biometrika* **93**, 279–288.
- [15] Butler, N.A. (2001). Optimal and orthogonal Latin hypercube designs for computer experiments, *Biometrika* **88**, 847–857.

BOXIN TANG

Latin Squares and Related Experimental Designs

Introduction

A Latin square is an arrangement of v copies of v symbols into a $v \times v$ square so that (a) each symbol occurs once in each row and (b) each symbol occurs once in each column. Three distinct Latin squares of order $v = 4$ are shown in Example 1. Other than for small v , the number of distinct (nonidentical as matrices) Latin squares is not generally known, though it is known that it grows rapidly with v ; see Laywine and Mullen [1, Chapter 1].

Example 1. Three 4×4 Latin squares. These squares are mutually orthogonal

1	2	3	4
2	1	4	3
3	4	1	2
4	3	2	1

1	2	3	4
3	4	1	2
4	3	2	1
2	1	4	3

1	2	3	4
4	3	2	1
2	1	4	3
3	4	1	2

Our interest in Latin squares is due to their many uses in designing experiments. Suppose you have the task of comparing the wear of drill bits of four different compositions when used in industrial presses. Each day each of your four press operators will

be given a new bit, with wear to be measured at the day's end; this will be repeated for 4 days. A Latin square can be employed to block out variation associated with operators and with days. Consider the first square in Example 1, and identify rows with operators, columns with days, and symbols with bit compositions. Then operator 1 uses composition 1 on day 1, operator 2 uses composition 3 on day 4, and so on. This design displays a basic fairness of allocation, in that each composition is exposed once to each operator and once to each day. Statistically this implies orthogonal **estimation** of the operator, day, and composition effects. This is but one example of the many adaptations of Latin squares for obtaining good experimental designs that will be explored in this chapter. Before proceeding, one more basic definition is needed. Two Latin squares of order v are said to be *orthogonal* (sometimes called a *Graeco-Latin pair*) if, on superimposing one square upon the other, every ordered pair of symbols occurs in the cells exactly once. Check Example 1 and you will see that every pair of squares has this property; the superimposition of all three squares is displayed in Example 2(I). The set of three squares is *mutually orthogonal*. This combinatorial orthogonality of Latin squares translates into statistical orthogonality of various treatment and **blocking** factors when exploited in an experimental design. Latin square designs are valuable because they provide orthogonal estimation, and because they do so with relatively little experimental material for given numbers of factors and levels.

Example 2. Various designs based on Latin squares

I:	1,1,1	2,2,2	3,3,3	4,4,4
	2,3,4	1,4,3	4,1,2	3,2,1
	3,4,2	4,3,1	1,2,4	2,1,3
	4,2,3	3,1,4	2,4,1	1,3,2

II:	1,5,9	2,6,10	3,7,11	4,8,12
	2,7,12	1,8,11	4,5,10	3,6,9
	3,8,10	4,7,9	1,6,12	2,5,11
	4,6,11	3,5,12	2,8,9	1,7,10

III:	1,2	2,3	1,3
	1,3	1,2	2,3
	2,3	1,3	1,2

IV:	(1,4,2,3,3),(3,3,3,3,3)	(2,2,3,4,1),(4,1,2,4,1)	(3,1,1,2,4),(1,2,4,2,4)	(4,3,4,1,2),(2,4,1,1,2)
	(2,3,1,4,4),(4,4,4,4,4)	(1,1,4,3,2),(3,2,1,3,2)	(4,2,2,1,3),(2,1,3,1,3)	(3,4,3,2,1),(1,3,2,2,1)
	(3,2,4,1,1),(1,1,1,1,1)	(4,4,1,2,3),(2,3,4,2,3)	(1,3,3,4,2),(3,4,2,4,2)	(2,1,2,3,4),(4,2,3,3,4)
	(4,1,3,2,2),(2,2,2,2,2)	(3,3,2,1,4),(1,4,3,1,4)	(2,4,4,3,1),(4,3,1,3,1)	(1,2,1,4,3),(3,1,4,4,3)

2 Latin Squares and Related Experimental Designs

Table 1 ANOVA skeletons for various applications of Latin squares (see text)

(I)	Source		df
	Rows		$v - 1$
	Columns		$v - 1$
	Symbols in square 1		$v - 1$
	$\vdots$		$\vdots$
	Symbols in square s		$v - 1$
	Error		$(v - 1)(v - s - 1)$
	Total		$v^2 - 1$
(II)	Source		df
	Rows		$av - 1$
	Columns		$bv - 1$
	Treatments		$v - 1$
	Error		$(av - 1)(bv - 1) - (v - 1)$
	Total		$abv^2 - 1$
(III)	Source		df
	Rows		$n - 1$
	Columns		$n - 1$
	Treatments		$nk - 1$
	Error		$(nk - 2)(n - 1)$
	Total		$n^2k - 1$
(IV)	Source		df
	Cells		$n^2 - 1$
	Treatments (adj)		$nk - 1$
	Error		$(nk - n - 1)(n - 1)$
	Total		$n^2k - 1$
(V)	Source		df
	rows		$k + av - 1$
	columns		$tv - 1$
	treatments (adj)		$v - 1$
	error		$(tv - 1)(k + av - 1) - (v - 1)$
	Total		$tv(k + av) - 1$
(VI)	Source		df
	Squares		$m - 1$
	Rows (squares)		$m(v - 1)$
	Columns (squares)		$m(v - 1)$
	Symbols in set 1		$v - 1$
	$\vdots$		$\vdots$
	Symbols in set s		$v - 1$
	Error		$(v - 1)(v - s - 1) + (m - 1)(v - 1)^2$
	Total		$mv^2 - 1$

Latin Squares for Blocking out Two Sources of Nuisance Variation

The drill bit experiment in the section titled “Introduction” illustrates the application of a Latin square as a row–column design for eliminating two sources of nuisance variation. Given a Latin square, let rows denote levels of one blocking factor, columns the levels of a second blocking factor, and symbols the levels of your treatment factor. The result is simultaneously a complete block design (*see Block Designs, Incomplete*) for treatments with respect to the row blocking factor and with respect to the column blocking factor.

Let $d[i, j]$ be the treatment appearing in row i , column j of Latin square d (e.g., $d[2, 4] = 3$ for the first square of Example 1). Denote by Y_{ij} the observation for the run in row i , column j (cell i, j). Analysis of the data is based on this model:

$$Y_{ij} = \mu + \rho_i + \gamma_j + \tau_{d[i, j]} + E_{ij} \quad (1)$$

inducing the skeleton analysis of variance or **ANOVA** (put $s = 1$) in Table 1(I). The design with data ($mm \times 10^{-1}$) and the completed ANOVA table for the drill bit experiment is shown in Table 2. In accordance with a proper **randomization** – for a randomly selected Latin square, randomly permute the rows, then the columns, then the symbols – blocking factors should not be tested. Several randomizers for Latin squares can be found on the web, for example www.edgarweb.org.uk. (For ANOVA basics, *see Analysis of Variance.*)

Use of a Latin square means that each pair of rows, columns, and treatments is statistically orthogonal, one consequence of which is that treatment comparisons are based on unadjusted, or raw, data means. So in the analysis of the drill bit data, the raw means 5.5, 3.9, 5.7, 7.7 are the basis for all treatment comparisons and can be used in any standard multiple comparisons procedure (Tukey, Student–Newman–Keuls (SNK), etc.). The drawbacks to this basic Latin square design arise from the requirement that numbers of rows, columns, and

Table 2 Analysis of a 4×4 Latin square experiment

Drill bit design (data)				Treatment	Mean	Source	df	SS	MS	F ratio	p value
1 (5.3)	2 (3.4)	3 (3.6)	4 (4.9)	1	5.5	Operators	$v - 1 = 3$	14.38	4.793	–	
2 (3.8)	1 (6.8)	4 (7.9)	3 (5.3)	2	3.9	Days	$v - 1 = 3$	12.48	4.160	–	
3 (5.6)	4 (8.3)	1 (5.7)	2 (2.8)	3	5.7	Bits	$v - 1 = 3$	29.12	9.707	12.83	0.0051
4 (9.7)	3 (8.3)	2 (5.6)	1 (4.2)	4	7.7	Error	$(v - 1)(v - 2) = 6$	4.54	0.757		
				Total			$v^2 - 1 = 15$	60.52			

treatments are all equal. What can one do, maintaining orthogonality, if this restriction cannot be met, or if it leaves too few **degrees of freedom** for error? These questions are addressed in the following subsections.

More Rows and Columns

More levels of the row factor and/or column factor can be accommodated by juxtaposing several Latin squares. Given ab Latin squares of order v , arrange them into a rows of b squares each, then “paste” them together to yield an $av \times bv$ row–column design. Any ab Latin squares can be used (even ab copies of the same square). This technique (a) allows some flexibility in numbers of rows and columns while (b) preserving the statistical orthogonality properties of a single Latin square and (c) providing more degrees of freedom for error. If now d denotes the $av \times bv$ design, then the model is still equation (1), and the skeleton ANOVA becomes that in Table 1(II). It is also possible to extract degrees of freedom for treatment $\times$ row and/or treatment $\times$ column interaction, but how much can be done is design specific, that is, depends on the Latin squares that are pasted together.

More Squares

Another possibility for gaining error degrees of freedom is multiple Latin squares. If the drill bit experiment is to be run at m sites, each with its own operator and day-to-day variability, then m Latin squares (identical or different) can be employed. The resulting skeleton ANOVA is in Table 1(VI) (put $s = 1$). Again, the orthogonality properties of a single Latin square are preserved, so that treatment comparisons are based on means.

More EUs per Cell: Semi-Latin Squares and k -depth Latin Squares

The basic Latin square has one experimental unit (EU or run), and thus only one treatment evaluated, per cell. If we can procure $k > 1$ EUs per cell in an $n \times n$ row–column blocking layout, then $v = nk$ treatments can be accommodated in each row and column. The design is a *semi-Latin square*: there are k symbols (treatments) in each cell so that each symbol appears once in each row and once in each column.

Semi-Latin squares can be found as superimpositions of k Latin squares of order n , each square having its own set of distinct symbols. Choice of the k squares depends on which of two reasonable models are employed. For Y_{ijl} the measurement on unit l in cell (i, j) , they are

$$Y_{ijl} = \mu + \rho_i + \gamma_j + \tau_{d[i,j,l]} + E_{ijl} \quad (2)$$

and

$$Y_{ijl} = \mu + \beta_{ij} + \tau_{d[i,j,l]} + E_{ijl} \quad (3)$$

Model (2) is the direct extension of (1), accounting for noise through additive effects for the row and column nuisance factors, while the more detailed model (3) fits a separate noise effect for each cell. The respective ANOVA skeletons are Table 1(III, IV). With model (2), treatment analysis is based on means, a consequence of the orthogonality of treatments to rows and columns. Because treatments need not be orthogonal to cells, treatment analysis with model (3) is adjusted for the cell blocks. There are also fewer error degrees of freedom with model (3). These are the trade-offs for the further variance reduction afforded by the finer blocking structure.

With model (2), any k Latin squares may be superimposed to form the semi-Latin square, for all such choices provide orthogonality of rows, columns, and symbols and so are equally efficient. With model (3), the driving consideration is not the rows or the

4 Latin Squares and Related Experimental Designs

columns, but the cells. The goal is to maximize the efficiency of the block design formed by the cell blocks. So long as $k < n$, the best choice is to superimpose k MOLS, resulting in a *Trojan square*. Example 2(II) is a Trojan square (from Example 2(I) with symbols distinguished). Details on maximal k for $n \times n$ Trojan squares are given in the section titled “Designs based on MOLS”. For $n = 6$, the best semi-Latin square for $k = 2$ is reported in Bailey and Royle [2]. Bailey [3] gives a thorough treatment of semi-Latin squares.

Latin and semi-Latin squares use each treatment once in each row and each column. With k EUs per cell, orthogonality in an $n \times n$ row-column design is maintained if each treatment appears nk/v times in each row and column. A k -depth Latin square is an arrangement of vk copies of v symbols into a $v \times v$ square so that each cell has k symbols, and each symbol occurs k times in each row and in each column.

If when superimposing k Latin squares in the semi-Latin construction above, one takes the symbols to be the same in each square, then one gets a k -depth Latin square (see Example 2(I)). This particular class of designs, while fully efficient for model (2), is inefficient for (3) due to replication of the same treatment within cell blocks. The skeleton ANOVA for model (2) is Table 1(I) with $s = 1$ and an additional $(k - 1)v^2$ degrees of freedom added to error and total.

For model (3), start with an $n \times n$ Latin square and a BIBD for v treatments in n blocks of size k . Place block i in each cell of the Latin square containing symbol i for $i = 1, \dots, n$. The resulting $n \times n$ row-column design with k runs per cell is fully efficient for both models (2) and (3). It is a k -depth Latin square if $n = v$. It is also a *doubly resolvable BIBD*, others of which are cataloged in Preece [4]. Example 2(III) displays $n = v = 3$, $k = 2$. The standard BIBD analysis applies.

Other Numbers of Rows and Columns: Youden Designs

If one drops a single row from a Latin square, the rows of the resulting $(v - 1) \times v$ array are a complete block design, while the columns are a BIBD. This *incomplete Latin square* is one version of a class of designs known as *Youden squares*. A Youden square

is a row-column design with k rows and v columns for which each treatment appears once in each row, and columns are blocks of a BIBD. Every Youden square is an incomplete Latin square, though the converse is not true.

Rows and columns may be added by joining separate Youden designs and Latin squares. Paste columns of t Youden squares into a $k \times tv$ array, then paste at Latin squares as in the section titled “More Rows and Columns”, to produce a $(k + av) \times tv$ row-column design called a *regular generalized Youden design*. With the standard row-column model (1) the skeleton ANOVA is Table 1(V). Since treatments are orthogonal to rows, the treatment adjustment is for columns only.

Here is the general idea. Start with any **incomplete block** design, with v treatments in b blocks of size k . Thinking of the blocks as columns, paste them into a $k \times b$ row-column array. If the treatments are equally replicated and if b is a multiple of v (so $b = tv$ for some t), then the treatments can be ordered within the columns so that each treatment appears t times in each row. This is now a row-column design with treatments orthogonal to rows; treatment estimates are the same as for the original block design without the row factor. If more rows are needed, it can be extended to a $(k + av) \times tv$ row-column design using Latin squares as explained above and with the same skeleton ANOVA.

Designs Based on MOLS

The Trojan squares of the section titled “More EUs per Cell: Semi-Latin Squares and k -depth Latin Squares” are built with mutually orthogonal Latin squares, or MOLS. Here other designs based on or related to MOLS will be described. The common theme is that of adding more factors to the designs of the section titled “Latin Squares for Blocking out Two Sources of Nuisance Variation”, maintaining orthogonality without changing the amount of experimental material. But first a basic question must be addressed: how many MOLS of a given order are there? Let $N(v)$ be the largest number of MOLS of order v . Then $N(v) \leq v - 1$ for every v , and $N(v) = v - 1$ for v that is a prime number or a power of a prime number. Example 1 provides $N(4) = 3$ MOLS of order 4. While upper bounds sharper than $v - 1$ are known for some other v , and $N(v) \geq 2$

for all $v > 6$, determining the exact value of $N(v)$ is a daunting (to say the least) combinatorial problem. When $v \leq 20$ is not a prime power, it is known that $N(6) = 1$, $N(10) \geq 2$, $N(12) \geq 5$, $N(14) \geq 3$, $N(15) \geq 4$, $N(18) \geq 3$, and $N(20) \geq 4$. For a catalog of known MOLS for every $v \leq 56$ along with further results, see Abel *et al.* [5].

Row-column Designs with More Blocking and/or Treatment Factors

Let's introduce a third blocking factor, four suppliers of parts to be drilled, in the drill bit experiment described in the section titled "Introduction". The original design, with rows for operators, columns for days, and symbols for drill bits, is the first Latin square in Example 1. Let the levels of supplier be the symbols in the second square. Superimposing the two squares, we see, for instance, that operator 2 on day 4 uses drill bit composition 3 on parts from supplier 2. The pair of MOLS means that each supplier's parts are used once with each bit type, once on each day, and once by each operator. That is, the supplier factor is orthogonal to the three factors operator, day, and bit type (which, as we already know, are orthogonal to one another). The third blocking factor has been introduced without requiring any further experimental runs, and without changing the fact that treatment comparisons are based on treatment averages.

The example shows that two superimposed MOLS provide a design for comparing v treatments using v^2 runs subject to three v -level blocking factors. In the same way, s superimposed MOLS provide a design comparing v treatments using v^2 runs subject to $s + 1$ v -level blocking factors; effects of all factors are orthogonal to one another. For the drill bit example, the symbols in the third square of Example 1, used as levels of a fourth blocking factor, produce the design in Example 2(I).

Another advantage of the orthogonality is that any of the blocking factors could alternatively be used as a treatment factor, so long as estimation of treatment interactions is not required. Thus s MOLS are designs for f v -level treatment factors using v^2 runs subject to $s + 2 - f$ v -level blocking factors, for any $1 \leq f \leq s + 2$. Table 1(I) shows the skeleton ANOVA. Each of the first $s + 2$ rows is identified with either a blocking factor or a treatment factor. Because only main effects, and no interactions, of treatment factors

are estimable, these are called (blocked) *main-effects plans*. The case $f = s + 2$ demonstrates that s MOLS are equivalent to an **orthogonal array** of strength 2 with v^2 runs on $s + 2$ v -level factors.

More Rows and Columns: Orthogonal F -rectangles

The designs of the section titled "More Rows and Columns" can also incorporate more blocking (or treatment) factors through orthogonal mates. An $av \times bv$ array with one symbol per cell is an F rectangle (F square if $a = b$) if each of v symbols occurs b times in each row and a times in each column; the designs of the section titled "More Rows and Columns" are F rectangles. Two F rectangles are orthogonal if superimposition produces each ordered pair of symbols in ab cells. Mutually orthogonal F rectangles admit the same applications as described for MOLS in the section title "Row-Column Designs with More Blocking and/or Treatment Factors", with now the row and column factors having av and bv levels (and thus $av - 1$ and $bv - 1$ df), respectively. A set of s mutually orthogonal F rectangles is equivalent to a mixed-level orthogonal array of strength 2 with $s + 2$ columns. For examples of mutually orthogonal F squares, and for further references, see Laywine and Mullen [6].

More Squares: Orthogonal Collections of Latin Squares

Orthogonal mates can likewise be found for sets of m Latin squares in the section titled "More Squares". A collection of s ordered sets, each containing m Latin squares of order v , is an (s, m) *orthogonal collection* if, on superimposition of the m corresponding squares of any two sets, each ordered pair of symbols occurs m times. MOLS are the special case $m = 1$.

The symbols in a given set of m squares, the same for each of these m squares, are levels of additional v -level treatment factors. With $s = 2$ and $m = 3$, for example, the drill bit experiment could be run at 3 sites, each with its own days and operators, but evaluating levels of two four-level treatment factors: drill bit composition and (say) lubricant. The skeleton ANOVA is Table 1(VI). Designs and further details may be found in Morgan [7], where in particular it is seen that four orthogonal six-level factors can be accommodated in two 6×6 squares. As seen above

($N(6) = 1$), only one such factor is possible in one 6×6 square.

More EUs per cell

It is not possible to add orthogonal v -level treatment factors to semi-Latin squares, for they have less than v_2 EU's. However, this can be done for the other designs in the section titled "More EUs per Cell: Semi-Latin Squares and k -depth Latin Squares".

Starting with an (s, k) -orthogonal collection of Latin squares of order v , we can make a $v \times v$ row-column design with k units per cell for s orthogonal treatment factors. The symbol in cell (i, j) of square u ($= 1, \dots, k$) in set w ($= 1, \dots, s$) of the collection is the level of factor w that appears on unit u in cell (i, j) of our design. If $s = 1$ this is a k -depth Latin square of the section titled "More EUs per Cell: Semi-Latin Squares and k -Depth Latin Squares"; if $k = 1$ it is the design of the section titled "Row-Column Designs with More Blocking and/or Treatment Factors". Relative to the k -depth square, $s - 1$ additional orthogonal treatment factors have been incorporated; relative to a set of MOLS, incorporating k EUs per cell has increased error degrees of freedom by $(k - 1)v^2$. The skeleton ANOVA is Table 1(I) with an additional $(k - 1)v^2$ degrees of freedom added to error and total. Example 2(IV) has five four-level factors in a 4×4 design with $k = 2$ units per cell. Like a single k -depth Latin square, this design is fully efficient for model (2) but inefficient for model (3).

The design in the last paragraph of the section titled "More EUs per Cell: Semi-Latin Squares and k -depth Latin Squares" can accommodate orthogonal treatment factors by replacing the BIBD with an orthogonal set of BIBDs. Orthogonal sets of BIBDs may be found in Morgan and Uddin [8]. The orthogonality of treatment factors is after adjusting for the cell blocking factor in model (3).

Orthogonal Youden designs

It is sometimes possible to add orthogonal treatment factors to regular generalized Youden designs. Here, too, the orthogonality of treatment factors is after adjusting for the column blocks (see Morgan and Uddin [8]).

Discussion

This paper provides an overview of the chief uses of Latin squares for designs for industrial and manufacturing experiments. The designs combine orthogonal estimation with an economy of experimental units, providing for both blocking and treatment factors. While the variants presented offer flexibility in numbers of rows, columns, factors, and so on, there are many experimental possibilities that are not amenable to the inherent restrictions of Latin squares. In such cases, orthogonality is sacrificed.

Specialized Latin squares have been employed in experimental situations not discussed above. *Row-complete Latin squares*, including *Williams squares*, are used as crossover designs (also called *changeover* or *repeated measures designs*); see Hinkelmann and Kempthorne [9, Chapter 19]. *Quasi-complete Latin squares* are row-column designs with nondirectional neighbor balance as discussed in Bailey [10] and Morgan [11]. Hedayat [12] examines designs based on *self-orthogonal Latin squares*. *Gerechte squares* [13] are Latin squares with a third blocking factor orthogonal to treatments but not to rows or columns. When $v = 9$ they include what are popularly known as *completed Sudoku squares*.

References

- [1] Laywine, C.F. & Mullen, G.L. (1998). *Discrete Mathematics Using Latin Squares*, John Wiley & Sons, New York.
- [2] Bailey, R.A. & Royle, G. (1997). Optimal semi-Latin squares with side six and block size two, *Proceedings of the Royal Society of London. Series A* **453**, 1903–1914.
- [3] Bailey, R.A. (1992). Efficient semi-Latin squares, *Statistica Sinica* **2**, 413–437.
- [4] Preece, D.A. (1967). Cyclic generation of Robinson's balanced incomplete block designs, *Biometrics* **23**, 574–578.
- [5] Abel, R.J.R., Colbourn, C.J. & Dinitz, J.H. (2007). Mutually orthogonal Latin squares (MOLS), in *Handbook of Combinatorial Designs*, 2nd Edition, C.J. Colbourn & J.H. Dinitz, eds, CRC Press, Boca Raton.
- [6] Laywine, C.F. & Mullen, G.L. (2007). Frequency squares and hypercubes, in *Handbook of Combinatorial Designs*, 2nd Edition, C.J. Colbourn & J.H. Dinitz, eds, CRC Press, Boca Raton.
- [7] Morgan, J.P. (1998). Orthogonal collections of Latin squares, *Technometrics* **40**, 327–333.

- [8] Morgan, J.P. & Uddin, N. (1996). Optimal blocked main effects plans with nested rows and columns and related designs, *Annals of Statistics* **24**, 1185–1208.
- [9] Hinkelmann, K. & Kempthorne, O. (2005). *Design and Analysis of Experiments, Vol. 2: Advanced Experimental Design*, John Wiley & Sons, New York.
- [10] Bailey, R.A. (1984). Quasicomplete Latin squares: construction and randomization, *Journal of the Royal Statistical Society. Series B* **46**, 323–334.
- [11] Morgan, J.P. (1988). Balanced polycross designs, *Journal of the Royal Statistical Society. Series B* **50**, 93–104.
- [12] Hedayat, A. (1973). Self orthogonal Latin square designs and their importance, *Biometrics* **29**, 393–396.
- [13] Bailey, R.A., Kunert, J. & Martin, R.J. (1990). Some comments on Gerechte designs I: analysis for uncorrelated errors, *Journal of Agronomy and Crop Science* **165**, 121–130.

J. P. MORGAN

Lenth's Method for the Analysis of Unreplicated Experiments

Introduction

Consider a 2^{7-4} experimental design (see **Fractional Factorial Designs**) involving seven factors, each at two levels, and only eight experimental runs. This is an example of a saturated experiment, in that we must estimate eight parameters (seven main effects plus the mean), leaving no **degrees of freedom** for error.

The analysis of such experiments has historically been rather subjective. A popular method has been that of Daniel [1] (see **Half-Normal Plot**), whereby the absolute values of the seven main-effect estimates are plotted against half-normal scores (i.e., equally spaced **quantiles** of the half-normal distribution). If all the true main effects are zero, these points should fall near a straight line; however, if there is a small number of “active” (nonzero) effects while the remaining ones are inactive, the inactive ones should follow a straight line and the active ones should look like **outliers**.

The half-normal-plot method requires one to make a judgment about what is an outlier and what is not. In an effort to remove this element of subjectivity, Lenth [2] proposed a method for estimating a standard-error-like quantity, called the pseudo standard error or PSE. The effects are then judged relative to the PSE to decide whether there is enough evidence for them to be deemed active.

Both the Daniels method and the Lenth method rely on the concept of “effect sparsity” – that only a few effects are active and the rest are inactive. This concept is discussed more rigorously in Box and Meyer [3], along with another objective method for analyzing such experiments.

The Method

Lenth's method assumes that we have m independent contrast or effect estimates and that they all have the same variance τ^2 . In an orthogonal two-level experiment on N observations, for example, each

contrast has the form $c = \bar{y}_+ - \bar{y}_-$, where $\bar{y}_+$ is the average of the $N/2$ observations at the “high” level of a factor and $\bar{y}_-$ is the average of the $N/2$ observations at the “low” level. Each such contrast has variance $\tau^2 = 4\sigma^2/N$, where σ^2 is the error variance. Let $c_1, c_2, \dots, c_m$ denote the m contrast estimates. Usually, owing to the fact that the model is saturated, we have $m = N - 1$.

The computation involves two stages. First, let

$$s_0 = 1.5 \cdot \text{median} \{|c_j|\} \quad (1)$$

Then, let

$$PSE = 1.5 \cdot \text{median} \{|c_j| : |c_j| \leq 2.5s_0\} \quad (2)$$

Note that the second step is just like the first, except that we exclude those effects that exceed $2.5s_0$ in absolute value. PSE is termed the pseudo standard error, and it is an estimate of τ .

Once the PSE is obtained, one can multiply it by a factor from Table 1 to obtain a margin of error (ME) for the contrasts. Contrasts that exceed the ME in absolute value are deemed active at the 5% **significance level**. Alternatively, one may compute a simultaneous margin of error (SME) using a different factor from Table 1. The distinction is that there is at most a 5% chance that one individual inactive contrast will exceed the ME, while there is at most a 5% chance that *any* inactive contrast will exceed the SME.

The recommended practice is to display the contrasts in a bar chart or a **Pareto** chart, with decision lines added for the ME and the SME (see the example below). A contrast that extends beyond the SME line is clearly an active effect. One that exceeds the ME but not the SME should be viewed with some caution, as it may be an artifact of testing several contrasts.

Some software implementations of Lenth's method produce a **half-normal plot** of the contrasts with a reference line having a slope based on the

Table 1 Multipliers to obtain the ME and SME at the $\alpha = 0.05$ level

m	for ME	for SME	m	for ME	for SME
7	2.297	4.867	19	2.120	4.118
8	2.201	4.868	23	2.097	4.017
11	2.211	4.438	26	2.082	3.985
15	2.156	4.240	27	2.077	3.964
17	2.138	4.164	31	2.064	3.925

2 Lenth's Method-Analysis of Unreplicated Experiments

PSE; however, this is a poor alternative to the bar chart and is not recommended because it does not show the decision thresholds.

The values in Table 1 are excerpted from tables in Ye and Hamada [4] that cover several additional values of m and significance levels other than .05.

Example Box *et al.* [5] present data from an eight-run experiment involving seven two-level factors. The response variable is the time it takes to climb a particular hill on a bicycle, and the factors are such things as seat height, dynamo off or on, and so on. The data are displayed in Table 2. The factor levels shown in the headings correspond to the respective labels “–” and “+” in the table. The final row gives the contrast estimates c_j . (To improve clarity, the factor levels are altered from those in the reference.)

The calculations are as follows. The median of the absolute contrasts $\{3.5, 12, 1.0, 22.5, 0.5, 1.0, 2.5\}$ is 2.5; hence, $s_0 = 1.5 \times 2.5 = 3.75$ and $2.5s_0 = 9.375$. Since two of the seven contrasts exceed this bound, we compute the PSE from the remaining five: $PSE = 1.5 \cdot \text{median}\{3.5, 1.0, .5, 1.0, 2.5\} = 1.5 \times 1 = 1.5$. Now, using Table 1 with $m = 7$, we have $ME = 2.297 \times 1.5 = 3.45$ and $SME = 4.867 \times 1.5 = 7.30$. Figure 1 shows a bar chart of the contrasts with reference lines added for $\pm ME$ and $\pm SME$.

The contrasts for gear and dynamo exceed the SME, and so they are identified as active effects. Both are positive effects, meaning that a higher gear ratio and the presence of the dynamo both increase the expected time to climb the hill. The contrast for seat barely surpasses the ME in the negative direction, providing a hint that having the seat up might reduce the time. However, this contrast does not surpass the

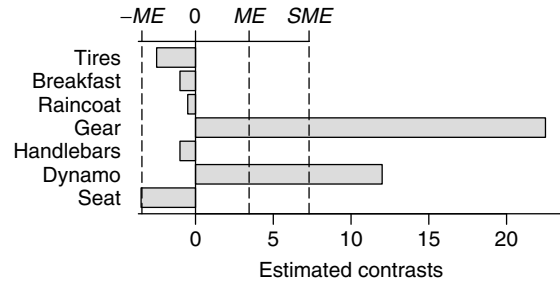


Figure 1 Bar chart of the contrast estimates for the bicycle example

SME, which means that if we adjust for multiple testing, it is not identified as active.

Discussion

Subsequent to the publication of [2], it has been found that the critical values suggested there are too conservative; hence, the tables in [4], excerpted in Table 1, are preferable. These critical values are based on the complete null case – where all the contrasts are assumed inactive. When one or more of the contrasts are actually active, that tends to inflate PSE somewhat, making the procedure somewhat conservative. It grows more conservative as the number of truly active effects increases, and is most conservative when the magnitude of the active effect is only moderate (less than 2τ or so). (Very large active effects lead to large contrasts that have a very high chance of exceeding $2.5s_0$ and thus being discarded in the second step of computation.)

The main advantage of Lenth's method is its simplicity and ease of use, making it easy to incorporate

Table 2 Data from the bicycle experiment

Time (s)	Seat down/up	Dynamo no/yes	Handlebars down/up	Gear low/med	Raincoat off/on	Breakfast no/yes	Tires soft/hard
50	–	–	–	–	–	+	+
52	–	–	+	–	+	–	–
88	–	+	–	+	–	–	–
83	–	+	+	+	+	+	+
71	+	–	–	+	+	+	–
69	+	–	+	+	–	–	+
59	+	+	–	–	+	–	+
60	+	+	+	–	–	+	–
Contrast:	–3.5	12.0	–1.0	22.5	–0.5	–1.0	–2.5

into programs or spreadsheet computations, or even to do by hand. Its main disadvantage is that it is based on a restrictive set of assumptions: independent contrast estimates that all have the same variance.

Lenth's method should not be used in replicated experiments, where the standard error can be estimated more efficiently by traditional analysis-of-variance methods.

There are other objective methods that can be used instead of Lenth's method. Most notable among these is the Bayesian method of Box and Meyer [3], which (unlike the Lenth method) is extensible to a broad class of regression variable-selection problems (see [6]). There is also a formulation in [7] of the Box-Meyer method that analyzes which factors (rather than which effects) are active. A survey and **Monte Carlo** comparison of several methods is given in [8].

References

- [1] Daniel, C. (1959). Use of half-normal plots in interpreting factorial two-level experiments, *Technometrics* **1**, 311–340.
- [2] Lenth, R. (1989). Quick and easy analysis of unreplicated factorials, *Technometrics* **31**, 469–473.
- [3] Box, G. & Meyer, R. (1986). An analysis for unreplicated fractional factorials, *Technometrics* **28**, 11–18.
- [4] Ye, K. & Hamada, M. (2000). Critical values of the Lenth method for unreplicated factorial designs, *Journal of Quality Technology* **32**, 57–66.
- [5] Box, G., Hunter, J. & Hunter, W. (2005). *Statistics for Experimenters*, 2nd Edition, John Wiley & Sons, New York, p. 245.
- [6] Chipman, H., Hamada, M. & Wu, C. (1997). A Bayesian variable-selection approach for analyzing designed experiments with complex aliasing, *Technometrics* **39**, 372–381.
- [7] Box, G. & Meyer, R. (1993). Finding the active factors in fractionated screening experiments, *Journal of Quality Technology* **25**, 94–105.
- [8] Balakrishnan, N. & Hamada, M. (1998). Analyzing unreplicated factorial experiments: a review with some new proposals. With discussions, *Statistica Sinica* **8**, 1–41.

RUSSELL V. LENTH

Main Effects

The Cell-Means Definition

In a *factorial experiment* (see **Factorial Experiments**), the *main effect* of a given factor, say F , is a measure of the average change in expected response that occurs when the level of factor F is changed. Here the “average” is taken over all other factors in the experiment. We will first develop this idea through several examples, and then follow with a formal definition.

The 2 × 2 Experiment. Here there are two factors, say A and B , each measured at two levels, which we may denote by 0 and 1. Let us suppose that the true (theoretical) mean responses to the four treatment combinations are given in the following table:

		Levels of B	
		0	1
Levels of A	0	4	6
	1	8	9

The treatment combination (0, 0) – A and B both at level 0 – is assumed to yield a mean response of 4, and similarly for the other treatment combinations (or *cells*). The change in the average expected response due to change in the level of A is

$$\frac{4 + 6}{2} - \frac{8 + 9}{2} = -3.5 \quad (1)$$

(or 3.5 if we change in the opposite direction). Thus a measure of the main effect of A is -3.5 units. Similarly, the main effect of B is -1.5 .

Note that if we were to replace the 9 by a 6 in this example, we would calculate the main effect of B to be 0. This would not mean that B has no effect on responses, but merely that it affects responses only through its *interaction* (see **Interactions**) with A . We may say in this case that B has no main effect, or that the main effect of B is absent.

A simulated experiment based on the theoretical values given above might yield the following data, with two observations per cell. Cell averages, which are least-squares estimates of the true mean responses

above, are given in parentheses:

		0	1
	0	4.75, 3.79 (4.270)	6.24, 5.79 (6.015)
	1	7.83, 7.82 (7.825)	10.37, 8.83 (9.600)

The main effect of A calculated in equation (1) could thus be estimated by

$$\frac{4.270 + 6.015}{2} - \frac{7.825 + 9.600}{2} = -3.57 \quad (2)$$

The corresponding main effect of B would be estimated as -1.76 . The *analysis of variance* in Table 1 (see **Analysis of Variance**) confirms that the main effects of A and B are significant – that is, that the estimated values are significantly different from 0.

To formulate our ideas more generally, we denote the true mean response to treatment combination $t = (i, j)$ by μ_{ij} . We have described the main effect of factor A as the difference between the average response to level “0”, or $(\mu_{00} + \mu_{01})/2$, and that to level “1” of A , namely $(\mu_{10} + \mu_{11})/2$, that is,

$$\frac{\mu_{00} + \mu_{01}}{2} - \frac{\mu_{10} + \mu_{11}}{2} \quad (3)$$

(see, e.g. [1, Chapter 10]). Of course, we could subtract in the opposite order as well. In fact, one often describes this effect by ignoring the factors of $1/2$ and writing simply:

$$\mu_{00} + \mu_{01} - \mu_{10} - \mu_{11} \quad (4)$$

Expressions such as (3) and (4) are called *contrasts*, as they contrast the responses to (in this case) the different levels of factor A . We say that A *has no main effect*, or that *the main effect of A is absent*, if any of these expressions equals 0, in which case they all do. We also note that the **sum of squares** (SS) for a contrast is not affected if we multiply the contrast

Table 1 Analysis of variance for the 2 × 2 example

Source	df	SS	MS	F	p
A	1	25.4898	25.4898	58.33	0.002
B	1	6.1952	6.1952	14.18	0.020
Interaction	1	0.0004	0.0004	0.00	0.976
Error	4	1.7479	0.4370	–	–
Total	7	33.4334			

2 Main Effects

by a nonzero constant (such as -1 or $1/2$). Thus the main effect of factor A may simply be defined by the contrast (equation (4)). We note that the coefficients in this expression are $+1$ where $i = 0$, and -1 where $i = 1$, independent of the value of j (the level of B).

In similar fashion, the main effect of factor B is defined as the contrast

$$\mu_{00} - \mu_{01} + \mu_{10} - \mu_{11} \quad (5)$$

Here the values of the coefficients are independent of the value of i (the level of A). We note that the interaction of A and B is also given by a contrast, namely

$$\mu_{00} - \mu_{01} - \mu_{10} + \mu_{11} \quad (6)$$

If this equals 0, then $\mu_{00} - \mu_{01} = \mu_{10} - \mu_{11}$, so that changing the level of B alters μ_{ij} by the same amount at each level of A . In this case we say that *interaction is absent*. We see that the coefficients of equation (6) depend on the levels of *both* factors.

The $2 \times 2 \times 2$ Experiment. Here we have three factors, say A , B , and C , each at two levels. The treatment combinations are ordered triples $t = (i, j, k)$, where k denotes the level of factor C . The contrasts for the main effects are often given *via* a table such as the following:

Cell	000	001	010	011	100	101	110	111
A	1	1	1	1	-1	-1	-1	-1
B	1	1	-1	-1	1	1	-1	-1
C	1	-1	1	-1	1	-1	1	-1

The main effect of C , for example, is defined by the contrast

$$\begin{aligned} &\mu_{000} - \mu_{001} + \mu_{010} - \mu_{011} + \mu_{100} \\ &- \mu_{101} + \mu_{110} - \mu_{111} \end{aligned} \quad (7)$$

Some authors prefer to view this as a difference of averages, analogous to equation (3), and so multiply it by $1/4$.

The 2^n Experiment. The extension to four or more two-level factors is handled analogously. Wu and Hamada [2, p. 97] give an example of a $2 \times 2 \times 2 \times 2$ or 2^4 experiment to study the growth of an epitaxial layer on polished silicon wafers. The factors studied are susceptor-rotation method (continuous or oscillating), nozzle position (2 or

6), deposition temperature (1210 or 1220 °C), and deposition time (low or high). From the observations in each treatment combination they calculate an observed average thickness (μm), and use this to estimate the main effect of factor susceptor-rotation method [2 p. 105] by calculating $(1/8)$ (sum of average thicknesses, oscillating $-$ sum of average thicknesses, continuous) $= -0.077$.

This effect turns out not to be statistically significant ($p = 0.439$).

The 2×3 Experiment. This is the simplest example of a factorial experiment with a factor at more than two levels. Here we have two factors, say A and B , such that A has two levels and B has three. The main effect of B must now be described by *two* independent contrasts (we say that it has *2 degrees of freedom* (df)). Denoting the levels of A as before and those of B by 0, 1, and 2, one choice of main effects contrasts is given in the following table, which lists the contrast coefficients:

Cell	00	01	02	10	11	12
A	1	1	1	-1	-1	-1
B	1	-1	0	1	-1	0
	1	0	-1	1	0	-1

Other choices of contrasts are possible. As before, the coefficients of the A contrast do not depend on the level of B , and those of the two B contrasts do not depend on the level of A . As we have noted, the choice of contrasts does not affect the sums of squares for main effects, which thus provide an overall summary of the intensity of each main effect. (This also applies in so-called nonorthogonal designs, although there the sums of squares must be adjusted.)

The general case. If factor F is measured at s levels, then it has $s - 1$ df. This means that its main effect is described by a set of $s - 1$ contrasts in cell means. The main effect of F is said to be absent if all these contrasts are zero. We now give a more precise definition of the contrasts belonging to a main effect.

A factorial experiment consists of n factors, the i th factor measured at s_i levels, so that there are $s_1 \times \dots \times s_n$ *treatment combinations* or *cells*. We let T be the set of treatment combinations. Each $t \in T$ is a multi-index or n -tuple (e.g., $t = (i, j, k)$ if $n = 3$), and we denote the corresponding *cell mean* by μ_t .

A *contrast* is an expression $\sum_{t \in T} c_t \mu_t$ such that the coefficients c_t add up to zero. The least-squares estimate of a contrast is calculated by replacing cell means by observed sample averages.

We say that a contrast $\sum_{t \in T} c_t \mu_t$ *belongs to the main effect of the i th factor* if the coefficients c_t do not depend on any factors other than the i th [3]. Contrasts belonging to interactions may be defined in similar fashion.

The Factor Effects Parametrization

A different way of defining main effects arises by using a different set of parameters. In the two-factor case, the means μ_{ij} are expressed by the set of equations

$$\mu_{ij} = \mu + \alpha_i + \beta_j + \gamma_{ij} \quad (8)$$

This is often accompanied by a set of side conditions of the form

$$\begin{aligned} \sum_i v_i \alpha_i &= 0, & \sum_j w_j \beta_j &= 0, \\ \sum_i v_i \gamma_{ij} &= 0 \quad \text{for each } j, \\ \sum_j w_j \gamma_{ij} &= 0 \quad \text{for each } i \end{aligned} \quad (9)$$

where the weights v_i and w_j are nonnegative and sum respectively to unity. In the presence of these side conditions, μ represents an overall level of response, α_i is the *main effect of the i th level of A* , β_j is the *main effect of the j th level of B* , and γ_{ij} is the *interaction of the i th level of A with the j th level of B* . According to Scheffé, “if the set of constants $\{\mu, \alpha_i, \beta_j, \gamma_{ij}\}$ satisfies (condition (8)) this is not sufficient for them to be the general mean, main effects, and interactions, unless the side conditions (9) are satisfied” [7, p. 92].

The model $E(Y_{ij}) = \mu + \alpha_i + \beta_j + \gamma_{ij}$ is sometimes called a *factor effects model* [5], while $E(Y_{ij}) = \mu_{ij}$ is a *cell means model* [6].

With the constraints (9) in place, one may solve equation (8) for the factor effects parameters. It is not hard to see that the resulting equations express these parameters as contrasts in the cell means μ_{ij} . But do the contrasts expressing α_i belong to the main effect for A , and those for β_j to B , in our earlier sense? The

answer is yes if, and only if, one uses equal weights ($v_1 = \dots = v_a = 1/a$ and $w_1 = \dots = w_b = 1/b$). In this case equation (9) reduces to the more familiar “zero-sum” constraints

$$\begin{aligned} \sum_i \alpha_i &= 0, & \sum_j \beta_j &= 0, \\ \sum_i \gamma_{ij} &= 0 \quad \text{for each } j, \\ \sum_j \gamma_{ij} &= 0 \quad \text{for each } i \end{aligned} \quad (10)$$

For example, solving equation (8) in the 2×3 case using the zero-sum constraints, we find that

$$\begin{aligned} \mu &= \left(\frac{1}{6}\right) \mu_{..} \\ \alpha_1 &= \left(\frac{1}{6}\right) (\mu_{.1} - \mu_{.2}) \\ \beta_1 &= \left(\frac{1}{6}\right) (2\mu_{.1} - \mu_{.2} - \mu_{.3}) \\ \beta_2 &= \left(\frac{1}{6}\right) (-\mu_{.1} + 2\mu_{.2} - \mu_{.3}) \\ \gamma_{11} &= \left(\frac{1}{6}\right) (2\mu_{11} - \mu_{12} - \mu_{13} - 2\mu_{21} + \mu_{22} + \mu_{23}) \\ \gamma_{12} &= \left(\frac{1}{6}\right) (-\mu_{11} + 2\mu_{12} - \mu_{13} \\ &\quad + \mu_{21} - 2\mu_{22} + \mu_{23}) \end{aligned} \quad (11)$$

where, as usual, $\mu_{i.} = \sum_j \mu_{ij}$, $\mu_{.j} = \sum_i \mu_{ij}$, and $\mu_{..} = \sum_i \sum_j \mu_{ij}$. The other parameters are obtained from equation (11) by applying the constraints; for example, $\alpha_2 = -\alpha_1$. It is evident that the contrasts for α_1 and α_2 belong to the main effect of A , and those of the β_j to B .

Main Effects in Fractional Factorial Experiments

A subset $S \subset T$ of the set of treatment combinations in a factorial experiment is a *fractional factorial* (see **Fractional Factorial Designs**) experiment, or *fraction*. Viewed as an *orthogonal array* (see **Orthogonal Arrays**), a fraction has a given *strength* t . If $t \geq 1$, then the restriction of each main-effect contrast to the subset S will still be a contrast, and each main effect will still be represented by the same df. For example, in the 2^3 experiment above, suppose $S = \{000, 011, 101, 110\}$, a fraction of strength 2. Then in S the contrast (equation (7)) for the main effect of C becomes

$$\mu_{000} - \mu_{011} - \mu_{101} + \mu_{110} \quad (12)$$

4 Main Effects

If in addition $t \geq 2$, then contrasts belonging to different main effects will also be mutually orthogonal. We may take these restricted contrasts as defining the main effects in the fraction [7].

One difficulty that arises in a fraction is that restricted contrasts belonging to a main effect may be the same as those belonging to an interaction. We say that the main effect is *aliased* (see **Aliasing in Fractional Designs**) with the interaction in the fraction. If one assumes in addition that interactions are absent, then we say that the main effects are *measurable* in the fraction [8].

References

- [1] Box, G.E.P., Hunter, J.S. & Hunter, W.G. (2005). *Statistics for Experimenters: Design, Innovation and Discovery*, 2nd Edition, John Wiley & Sons, New York.
- [2] Wu, C.F.J. & Hamada, M. (2000). *Experiments: Planning, Analysis, and Parameter Design Optimization*, John Wiley & Sons, New York.

- [3] Bose, R.C. (1947). Mathematical theory of the symmetrical factorial design, *Sankhya* **8**, 107–166.
- [4] Scheffé, H. (1959). *The Analysis of Variance*, John Wiley & Sons, New York, Reissued in the Wiley Classics series, 1999.
- [5] Neter, J., Kutner, M.H., Nachtsheim, C.J. & Wasserman, W. (1996). *Applied Linear Statistical Models*, 4th Edition, Irwin Publishing, Chicago.
- [6] Hocking, R.R. (1985). *The Analysis of Linear Models*, Brooks/Cole Publishing, Monterey.
- [7] Beder, J.H. (2004). On the definition of effects in fractional factorial designs, *Utilitas Mathematica* **66**, 47–60.
- [8] Rao, C.R. (1947). Factorial experiments derivable from combinatorial arrangements of arrays, *Journal of the Royal Statistical Society* Suppl 9, 128–139.

Related Articles

Main Effect Designs; Factorial Designs, Resolution of.

JAY H. BEDER

Main Effect Designs

Introduction

In the planning stages of any experimental study, an investigator must determine which experimental factors to include in the study. An experimental factor is a variable that is studied in the experiment and believed to potentially have an effect on a response variable of interest. The values of a factor at which experimental runs are conducted are called the *levels of the factor*. During an experimental study, the levels of each factor are usually experimentally changed during different runs of the experiment so that the direct effects of each factor (called *main effects* (MEs)) on the response can be assessed as well as how factors affect the response in combination (called *interactions*, see **Interactions**). However, when the list of potential experimental factors is long, it is important in the early stages of an experimental process to identify those experimental factors which do in fact have an influential effect on a response variable and to screen out those that do not. This is the primary purpose of screening experiments (see **Screening Designs**). One type of screening experiment which is useful in such situations is called a *main effects design*. ME designs focus on the identification of influential MEs rather than interactions and are particularly useful for experimental studies in which the initial list of experimental factors is large compared to the number of runs that can be made in the initial screening experiment. In this article, we consider the use and construction of ME designs.

In an ME design, a factor may be either quantitative or qualitative. Quantitative factors such as time or age can assume all values on a continuous scale. Qualitative factors such as different machine types or different crop varieties have levels which can only assume discrete values and typically cannot be rank ordered. The selection of levels for factors to include in an experiment is somewhat subjective. For example, if a quantitative factor is believed to have primarily a linear effect on a response, then only two levels of the factor may be selected whereas if it is believed that a quantitative factor may effect a response nonlinearly, then three or more levels may be selected to allow for detection of a factor curvature effect. There is usually much less choice in the selection of levels of a

qualitative factor as the levels of such a factor are determined by the qualitative nature of the variable.

The most popular screening experiments in use today are factorial experiments (see **Factorial Experiments**). A factorial experiment is one in which the treatments consist of different combinations of experimental factor levels. ME designs are one type of factorial experiment.

Further Notation and Model

In this article we consider experimental situations in which m factors are to be studied using a factorial experiment with N runs. We shall use d to denote a design which can be used in such a situation and interchangeably represent d by an $N \times m$ matrix $X_d = (\mathbf{x}_{d1}, \dots, \mathbf{x}_{dm}) = (x_{dij})$ whose i th row corresponds to run i and j th column corresponds to factor j . If factor j has s_j levels, then the entries in column j of X_d will consist of symbols $1, 2, \dots, s_j$ with $x_{dij} = t \in \{1, 2, \dots, s_j\}$ if factor j occurs at level t in run i .

The model used to analyze the data from a given design d is

$$\mathbf{Y} = \beta_0 \mathbf{1}_N + \sum_{i=1}^m X_{di} \beta_i + \boldsymbol{\varepsilon} \quad (1)$$

where $\mathbf{Y}$ is a $N \times 1$ vector of observations, β_0 is an overall mean effect, $\mathbf{1}_N$ is an $N \times 1$ vector of 1's, X_{di} is an $N \times (s_i - 1)$ matrix corresponding to factor i having s_i levels, β_i is an associated set of ME parameters, and $\boldsymbol{\varepsilon}$ is a vector of random error terms whose entries are assumed to be uncorrelated with mean *zero* and constant variance σ^2 . The actual entries of X_{di} usually depend upon whether factor i is a quantitative or qualitative factor. If factor i is a quantitative factor with s_i levels, the entries of X_{di} are usually coded (using orthogonal polynomials) so that all polynomial effects up through degree $s_i - 1$ can be estimated. For example, if factor i has $s_i = 3$ equally spaced factors all occurring in $N/3$ runs, then X_{di} has two columns with entries $-1, 0, 1$ and $1, -2, 1$, respectively, depending on whether factor i occurs at its low, middle or high level. If factor i is a qualitative variable with s_i levels then the entries of X_{di} can be coded in several different ways. We shall assume for convenience that the average effect

2 Main Effect Designs

of factor i is zero which implies one parameterization where column p of X_{di} corresponding to level p , $p = 1, \dots, s_i - 1$, has entries 1 if factor i occurs at level p , -1 if factor i occurs at level s_i and 0 otherwise. Under this parameterization, the parameter corresponding to column p of X_{di} represents the difference in the average response at level p and the overall average. Because we are only interested in estimating MEs for factors, we do not include interaction terms in model (1). We shall say that a design is saturated if $\sum_{i=1}^m (s_i - 1) + 1 = N$ or nearly saturated if $\sum_{i=1}^m (s_i - 1) + 1$ is “close” to N . When ME designs are used as **screening designs**, they are often saturated or nearly saturated. In these situations, the identification of influential effects in model (1) usually depends upon the form of the matrices X_{di} . If the columns of all the X_{di} are orthogonal to one another and the corresponding parameter estimates can be assumed to be normal random variables with equal variances, then graphical methods such as half-normal plots (see **Half-Normal Plot**), can be used to identify significant effects. However, because the use of graphical methods is somewhat subjective in nature, other, more objective, statistical methods have been developed to identify significant effects in these situations, for example, see Hamada and Balakrishnan [1] for a good discussion of available methods such as Lenth’s method (see **Lenth’s Method for the Analysis of Unreplicated Experiments**). When the columns of the X_{di} in model (1) are nonorthogonal or are orthogonal but the corresponding parameter estimates do not satisfy the assumptions required for the use of half-normal plots or other available methods, regression analysis and variable selection techniques can be used to identify influential effects.

When using model (1) to analyze the data from a given design d , a useful property for d to have is that of proportional frequencies. In particular, let factors i and j have s_i and s_j levels, respectively, and let n_{pq} be the number of runs in d at level p of factor i and level q of factor j . We say factors i and j have proportional frequencies if for all $p = 1, \dots, s_i$ and $q = 1, \dots, s_j$,

$$n_{pq} = \frac{n_{p.} \cdot n_{.q}}{N} \quad (2)$$

where $n_{p.} = \sum_{l=1}^{s_j} n_{pl}$ and $n_{.q} = \sum_{l=1}^{s_i} n_{lq}$. The property of proportional frequencies was derived by Addelman [2] and insures that any difference between

the averages of observations at two levels of factor i is orthogonal (when considered as a linear combination of the observations) to any difference between the averages of any two levels of factor j . Any design in which all pairs of factors have proportional frequencies is called an *orthogonal main effects (OME) design*. The property of proportional frequencies in an OME design thus makes the **estimation** of OME parameters in model (1) relatively easy. We shall denote an OME design d in N runs having m_i factors with s_i levels, $i = 1, \dots, k$, by $\text{OME}(N, s_1^{m_1} \dots s_k^{m_k})$.

We note that there are no interaction terms included in model (1). This implies the assumption that interaction effects are small compared to MEs. When this assumption is valid, then the model given in equation (1) is useful for identifying significant MEs of factors on a response. When this assumption is invalid, the analysis using model (1) can lead to the identification of some factors as having an influential effect on a response variable when they do not, or, what is worse, lead the experimenter to miss entirely some influential factors. In any case, it is always wise after running an ME screening design to run a follow-up experiment to explore interactions between factors and develop a more detailed model relating the response variable to experimental factors.

Orthogonal Array Main Effect (OAME) Designs

Some of the most popular ME designs used today are the two- and three-level regular **fractional factorial** design, and Plackett–Burman (PB) (see **Plackett–Burman Designs**) designs. However, because these designs are all special cases of orthogonal array (OA) (see **Orthogonal Arrays**) designs of strength two defined below and because these designs are considered elsewhere in this volume we do not consider them here. We also note that OA designs serve as the building blocks for other useful and more general ME designs. We shall henceforth denote an OA design d of strength two in N runs having m_i factors with s_i levels, $i = 1, \dots, k$ by $\text{OAME}(N, s_1^{m_1} \dots s_k^{m_k})$.

Definition 1 Let d be an ME design with $N \times m$ matrix X_d .

Table 1 Orthogonal array main effects designs

Number of observations	Design
4	2^3
6	$3^1 2^1$
8	$2^7, 4^1 2^4$
9	3^4
12	$2^{11}, 3^1 2^4, 6^1 2^2, 4^1 3^1$
15	$5^1 3^1$
16	$2^{15}, 4^1 2^{12}, 4^2 2^9, 8^1 2^8, 4^3 2^6, 4^4 2^3, 4^5$
18	$3^7 2^1, 9^1 2^1, 6^1 3^6$
20	$2^{19}, 5^1 2^8, 10^1 2^2, 5^1 4^1$
24	$2^{23}, 4^1 2^{20}, 3^1 2^{11}, 6^1 2^{14}, 4^1 3^1 2^{13}, 12^1 2^{12}, 6^1 4^1 2^1, 8^1 3^1$
25	5^6
27	$3^{13}, 9^3, 3^9$
28	$2^{27}, 7^1 2^{12}$
30	$5^1 3^1 2^1$
32	$2^{31}, 4^1 2^{28}, 4^2 2^{25}, 8^1 2^{24}, 4^3 2^{22}, 8^1 4^1 2^{21}, 4^4 2^{19}, 8^1 4^1 2^{18}, 4^5 2^{16}, 16^1 2^{16}, 8^1 4^3 2^{15}, 4^6 2^{13}, 8^1 4^4 2^{12}, 4^7 2^{10}, 8^1 4^5 2^9, 4^8 2^7, 8^1 4^6 2^6, 4^9 2^4, 8^1 4^7 2^3, 8^1 4^8$
36	$2^{35}, 3^1 2^{27}, 3^2 2^{30}, 6^1 3^1 2^{18}, 3^4 2^{13}, 6^2 2^{13}, 9^1 2^{12}, 3^{12} 2^{11}, 6^1 3^2 2^{11}, 6^1 3^8 2^{10}, 6^2 3^1 2^{10}, 6^2 3^4 2^9, 6^3 2^8, 3^{13} 2^4, 6^3 3^1 2^4, 6^1 3^9 2^3, 6^3 3^2 2^3, 6^1 3^{12} 2^2, 6^2 3^5 2^7, 6^2 3^8 2^1, 6^3 3^3 2^1, 4^1 3^{13}, 12^1 3^{12}, 6^3 3^7$

1. We say factors p and q are orthogonal if in columns p and q of X_d each level of factor p occurs equally often with each level of factor q .
2. We say d is an $\text{OAME}(N, s_1^{m_1} \dots s_k^{m_k})$ design of strength two if every pair of factors in d is orthogonal.

Comment The definition of an orthogonal array main effect (OAME) design given above can be extended to that of an OA design of strength t (see **Orthogonal Arrays**). However, because such designs typically require relatively large numbers of runs, we do not consider them here. Also, because we will only be considering OAME designs of strength two, we will henceforth omit any reference to the strength of the designs considered. An excellent reference on the properties and construction of general OA designs is the book *Orthogonal Arrays: Theory and Applications* by Hedayat et al. [3].

We note that in an ME design d , if factors p and q are orthogonal to one another, then factors p and q have the property of proportional frequencies since each level of factor p occurs equally often with each level of factor q . Thus an OAME design

d is clearly an OME design since all pairs of factors in d are orthogonal. Another useful property regarding OAME designs is that it is possible to determine a lower bound for the number of runs required to accommodate all factors (though the bound is not necessarily achievable in all cases). The following theorem can be derived using the property of orthogonality of factors in an OAME design.

Theorem 1 For an $\text{OAME}(N, s_1^{m_1} \dots s_k^{m_k})$ design d , N must be divisible by $s_i s_j$, $i, j = 1, \dots, k$.

Example 1 Suppose an experimental situation requires the use of a design having m_1 factors with three levels and m_2 factors with two levels where $m_1 \geq 2, m_2 \geq 2$. Then, using Theorem 1, any number of runs required for an $\text{OAME}(N, 3^{m_1} 2^{m_2})$ design must have N divisible by $3^2 = 9, 2^2 = 4$, and $3 \times 2 = 6$. Hence, the minimum possible value for N is $N = 36$.

A large number of OAME designs have been constructed using various methods. The largest listing and catalogue of such designs known to the author is available on the web at www.research.att.com/

4 Main Effect Designs

~njas/oadir. This website is maintained by N.J.A. Sloane. Some useful OAME designs having 36 or fewer runs are listed in Table 1. For the construction and actual form of the arrays given in that Table, the reader is referred to the website given earlier.

With regard to the usage of OAME designs, there are certain advantages and disadvantages that should be kept in mind.

Advantages

1. Economy of run size. In a number of situations, an OAME design will allow an experimenter to run a screening experiment using the minimum number of runs required to estimate all or nearly all parameter effects of interest, that is, the OAME design will be saturated or nearly saturated. For example, if a screening experiment requires 18 runs and has eight factors with seven factors having three levels and one factor having two levels, then an $\text{OAME}(18, 3^7 2^1)$ can be used, or if there are four factors having three levels each to be screened in nine runs, then an $\text{OAME}(9, 3^4)$ can be used.
2. Estimates for all ME parameters between different factors can be estimated orthogonally to one another and such estimates will be unbiased as long as higher-order interactions are negligible.
3. Flexibility. OAME designs provide the opportunity to combine factors having differing numbers of levels in the same experiment. This is particularly useful when a study involves only qualitative or both qualitative and quantitative factors requiring differing numbers of levels.

Disadvantages

1. For a given design d , we say two effects are fully aliased (*see Aliasing in Fractional Designs*) or confounded if the **correlation** between their parameter estimates under model (1) (when their corresponding terms are included in the model) has absolute value 1 and we say two effects are partially aliased or confounded if the correlation between their estimates has absolute value between 0 and 1. The primary disadvantage with OAME designs, particularly those that are nearly saturated, is that many of the estimates for ME parameters will be fully or partially confounded with low order interactions. For example, in an

$\text{OAME}(18, 3^7 2^1)$ design d , because d is nearly saturated, there are at most two lower order interactions between MEs that can be estimated. Thus, most MEs in d will be at least partially aliased with some low order interactions. Whenever a design d has a number of MEs that are at least partially aliased with a number of low order interactions, we say that the design has a complex **aliasing** structure. A complex aliasing structure can sometimes lead to difficulty in the interpretation of a significant ME parameter estimate. For example, in the $\text{OAME}(18, 3^7 2^1)$ design d , if a parameter ME is found to be significant, then it may be difficult to determine (unless all two-factor and higher-order interactions are negligible) whether the significance was due to the associated factor effect or some combination of the partially aliased low order interactions. Hamada and Wu [4] have suggested a stepwise, regression-based procedure for analyzing data from designs having complex aliasing structures and their method of analysis can be useful when only a few effects are influential and the correlation between partially aliased effects is small to moderate. However, as also pointed out in Hamada and Wu [4], their technique is not foolproof. When the assumption of two-factor and higher-order interactions being negligible is valid, the OAME designs can be of great use in identifying influential factors in a screening process. However, when the assumption is not valid, it is possible to identify some factors which are not truly significant as being influential and completely fail to identify other factors which are truly influential.

Other OME Plans

As mentioned in the section titled “Orthogonal Array Main Effect (OAME) Designs”, OAME designs are special cases of OME plans. In fact, OAME designs provide one of the basic tools for constructing OME designs to handle experimental situations for which an OAME design may not exist. The simplest method for constructing OME designs from OAME plans is called the *method of collapsing factors*. We illustrate the technique through an example.

Example 2 Consider the OAME(9, 3⁴) design d which exists and has

$$X_d = \begin{pmatrix} 1 & 1 & 1 & 1 \\ 1 & 2 & 2 & 2 \\ 1 & 3 & 3 & 3 \\ 2 & 1 & 2 & 3 \\ 2 & 2 & 3 & 1 \\ 2 & 3 & 1 & 2 \\ 3 & 1 & 3 & 2 \\ 3 & 2 & 1 & 3 \\ 3 & 3 & 2 & 1 \end{pmatrix} \quad (3)$$

Using d , we can now construct an OME(9, 3²2²) design $\tilde{d}$ using the method of “collapsing factors”. In particular, we assign all observations occurring at levels 1, 2, and 3 of factors 1 and 2 in d to the same levels for factors 1 and 2 in $\tilde{d}$. However, we also assign those observations occurring at level 1 of factors 3 and 4 in d to the same level of factors 3 and 4 in $\tilde{d}$ but assign those observations occurring at levels 2 and 3 of factors 3 and 4 in d to level 2 of factors 3 and 4 in $\tilde{d}$. Thus $\tilde{d}$ has

$$X_{\tilde{d}} = \begin{pmatrix} 1 & 1 & 1 & 1 \\ 1 & 2 & 2 & 2 \\ 1 & 3 & 2 & 2 \\ 2 & 1 & 2 & 2 \\ 2 & 2 & 2 & 1 \\ 2 & 3 & 1 & 2 \\ 3 & 1 & 2 & 2 \\ 3 & 2 & 1 & 2 \\ 3 & 3 & 2 & 1 \end{pmatrix} \quad (4)$$

and we have “collapsed” factors 3 and 4 in d having three levels each to factors 3 and 4 in $\tilde{d}$ having two levels each. In doing so, we note that level 1 in factors 3 and 4 in $\tilde{d}$ occurs in three runs whereas level 2 of factors 3 and 4 in $\tilde{d}$ occurs in six runs. However, because the levels of each pair of factors in d occur equally often together, the property of proportional frequencies in d carries over to $\tilde{d}$ during the “collapsing” process. We also note that using the method of “collapsing factors”, additional OME designs can be constructed from d such as OME(9, 3³2¹), OME(9, 3¹2³), and OME(9, 2⁴).

Using the method of “collapsing factors”, a large number of OME designs can be constructed from OAME designs such as those listed in Table 1 or the web site mentioned in the section titled “Orthogonal Array Main Effect (OAME) Designs”. A collection of

useful OME designs can be found in Dey [5]. We also note that, as in the example above, when applying the method of “collapsing factors” to a given OAME design d , the property of proportional frequencies carries over from d to the OME design $\tilde{d}$ constructed.

Similar to OAME plans, the property of proportional frequencies makes it possible to find a lower bound for the number of observations required for an OME($N, s_1^{m_1} \dots s_k^{m_k}$) design. In particular, we have the following theorem which has been proved by Jacroux [6].

Theorem 2 Suppose that an OME design d requires $m \geq 3$ factors with factor i having s_i levels, $i = 1, \dots, m$ and assume $s_1 \geq s_2 \geq \dots \geq s_m$. If $N = \tilde{s}_1 \tilde{s}_2$ is the minimum possible value for integers $\tilde{s}_1$ and $\tilde{s}_2$ satisfying

1. $\tilde{s}_1 \geq s_2$ and $\tilde{s}_2 \geq s_2$;
2. $\tilde{s}_1 \tilde{s}_2 \leq 2s_1 s_2$;
3. $s_3 \times \{lcm(\tilde{s}_1, \tilde{s}_2)\} \leq \tilde{s}_1 \tilde{s}_2$,

where $lcm(x, y)$ denotes the least common multiple between positive integers x and y , then N is a lower bound for the number of observations required for d .

A design having N satisfying the conditions of Theorem 2 is called a *minimal OME plan*.

Comment If d requires m factors having level $s_1 \geq s_2 \geq \dots \geq s_m$, then another obvious lower bound for N is $\sum_{i=1}^m s_i - m + 1$ since this is the minimal number of runs required to estimate all MEs in d .

Example 3 Consider an experimental setting in which a screening experiment is to be run and requires the use of a design having two factors with three levels and two factors with two levels. Using Theorem 1, it is seen that a lower bound for the number of runs required for an OAME design is 36. On the other hand, it is seen using Theorem 2 that the lower bound for the number of observations required for an OME design is 9. We saw in Example 2 that an OME(9, 3²2²) exists which can be used to handle this experimental situation. Thus in this case, we see that the use of an OME design can provide substantial savings in terms of runs required over the best possible OAME design. Many other such situations can be found in which the use of OME designs provides similar savings in terms of runs required over corresponding OAME designs.

For a specific experimental situation requiring m

6 Main Effect Designs

factors with $\tilde{s}_1 \geq \dots \geq \tilde{s}_m$ levels, the following steps can be used to construct a specific OME design $\tilde{d}$ for use.

1. Use Theorem 2 to find a lower bound for the number of observations required for an OME plan to accommodate all m factors.
2. Consult a table of OAME designs such as given in Table 1 or at the website mentioned previously and find an OAME design d whose number of runs is equal to or as close as possible to the lower bound found in step 1 and having factors $i, i = 1, \dots, m$ with s_i levels, $s_i \geq \tilde{s}_i, i = 1, \dots, m$.
3. Collapse the levels of the factors $i = 1, \dots, m$ in d to the appropriate number of levels for factors $i = 1, \dots, m$ in $\tilde{d}$.

We illustrate the use of this procedure in the following example.

Example 4 Consider an experimental situation requiring three three-level factors and six two-level factors. What is the smallest $\text{OME}(N, 3^3 2^6)$ that can be run? Following the stepwise procedure outlined above, we proceed as follows:

1. Using Theorem 2, a lower bound for N is seen to be $N = 9$.
2. Consulting Table 1, we see that the smallest OA design d that has nine factors with three factors having at least three levels is $\text{OAME}(16, 4^3 2^6)$.
3. Using the method of collapsing factors, we can construct the desired $\text{OME}(16, 3^3 2^6)$ design $\tilde{d}$ by taking the three four-level factors in d and collapsing their levels so that we obtain the required three factors having three levels in $\tilde{d}$.

Once an OME has been selected for use, the actual analysis of data from the design would proceed, using regression and variable selection techniques. The systematic procedure of analysis suggested by Hamada and Wu [4] may also be of use under the same conditions mentioned in the previous section, but the limitations of the method should also be kept in mind.

As in the case of OAME designs, there are certain advantages and disadvantages an experimenter should keep in mind when using OME plans as screening designs.

Advantages over OAME designs

1. Economy of run size. Because OME designs are more general than OAME plans, it is often the case that the number of runs required to construct an OME design to handle a given number of factors having given numbers of factor levels is smaller than the number required to construct an OAME design. This is clearly illustrated in Example 3.
2. Because differences between factor level averages between different factors are orthogonal, this makes the estimation of ME parameters relatively easy in OME designs.

Disadvantages

1. Because most OME designs will be nearly saturated when used as screening experiments, many of the ME parameter estimates will be at least partially aliased with low order interactions which lead to a complex aliasing structure. All the problems mentioned earlier for complex aliasing structures in OAME designs hold true for OME plans as well.
2. When obtaining an OME design $\tilde{d}$ from an OAME design d through the method of “collapsing factors”, it will be the case that some of the factors in $\tilde{d}$ will have levels which occur in differing numbers of runs. Thus some ME parameters associated with these factors will be estimated with greater precision than others. This can result in a loss of efficiency. However, the loss of efficiency can be minimized by collapsing levels of a factor in d so that the levels of the resulting factor in $\tilde{d}$ occur as equally often in runs as possible.

Nonorthogonal ME Designs

In the previous sections, we introduced designs which allow for the orthogonal estimation of ME parameters associated with different factors. In this section, we consider designs which sometimes prove useful as screening designs, but which are nonorthogonal in the sense that ME parameters cannot necessarily be orthogonally estimated. In particular, we consider what we refer to as *nearly orthogonal array main effect (NOAME) designs*. NOAME designs were introduced by Wang and Wu [7]. In OAME designs, for any two factors i and j , each level of factor i occurs equally often with each level of factor j .

This latter property implies proportional frequencies between factors which ensure that ME parameters between different factors can be orthogonally estimated. The property of proportional frequencies between factors in both OAME designs and OME designs places some restrictions on the number of factors that can be accommodated in a given set of runs. For example, in a design requiring six runs and a $3^1 2^2$ fractional factorial experiment, it can be shown that neither an OAME design nor an OME design exists. However, by relaxing the condition of orthogonality between ME estimates for different parameters, Wang and Wu [7] give a method for constructing an NOAME design which can be used for such a situation.

Wang and Wu [7] give three different methods for constructing NOAME designs. In some cases, they use their methods to find designs which minimize the number of nonorthogonal pairs of columns of certain types and in other cases, they find designs which minimize the canonical correlations between two sets of effects. A list of parameters for the actual designs constructed by Wang and Wu [7] is given in Table 2, but for the actual construction of the designs given, the reader is referred to Wang and Wu [7].

Once again, the statistical analysis of NOAME designs usually relies upon regression analysis and variable selection techniques.

Advantages of NOAME designs over OME designs

1. Economy of run size. NOAME designs can be used to handle experimental situations where OAME and OME designs do not exist for a given number of runs, factors, and factor levels. Thus such designs are used in situations where other options do not exist.

Table 2 Nearly orthogonal array main effects designs

Number of observations	Design
6	$3^1 2^3$
10	$5^1 2^5$
12	$3^4 2^3, 6^1 2^5, 3^1 2^9, 3^5 2^1, 3^4 2^2$
15	$5^1 3^5$
18	$9^1 2^8, 3^4 2^8$
20	$5^1 2^{15}$
24	$8^1 3^8, 3^8 2^7, 4^1 3^8 2^4, 3^1 2^{21}, 6^1 2^{18}, 3^{11} 2^1$ $4^7 3^1, 4^1 3^9 2^1, 3^9 2^5, 4^6 3^1 2^3$
36	$4^{11} 3^1, 4^{10} 3^2$

Disadvantages

1. The analysis of NOAME designs is not as straightforward as that of orthogonal designs since parameter ME estimates between different factors are not necessarily orthogonal. The lack of orthogonality also leads to some loss of efficiency in terms of the estimating factor MEs.
2. The alias structure associated with NOAME designs tends to be even more complex than that of OAME and OME designs. Hence, all of the issues mentioned previously, associated with complex alias structures, hold for NOAME designs.

References

- [1] Hamada, M and Balakrishnan, N. (1998). Analyzing unreplicated factorial experiments: a review with some new proposals (with discussion), *Statistica Sinica* **8**, 1–41.
- [2] Addelman, S. (1962). Orthogonal main effects plans for asymmetrical factorial experiments, *Technometrics* **4**, 21–45.
- [3] Hedayat, A., Sloane, N. & Stufken, J. (1999). *Orthogonal Arrays: Theory and Applications*, Springer, New York.
- [4] Hamada, M. & Wu, C.F.J. (1992). Analysis of designed experiments with complex aliasing, *Journal of Quality Technology* **24**, 130–137.
- [5] Dey, A. (1985). *Orthogonal Fractional Factorial Designs*, Holsted Press, New York.
- [6] Jacroux, M. (1992). A note on the determination and construction of minimal orthogonal main effects plans, *Technometrics* **34**, 92–96.
- [7] Wang, J.C. & Wu, C.F.J. (1992). Nearly orthogonal arrays with mixed levels, *Technometrics* **34**, 409–422.

Further Reading

- John, P.W.M. (1971). *Statistical Design and Analysis of Experiments*, Macmillan, New York.
- Plackett, R.L. & Burman, J.P. (1946). The design of optimum multifactorial experiments, *Biometrika* **33**, 305–325.

Related Articles

Aliasing in Fractional Designs; Factorial Experiments; Fractional Factorial Designs; Half-Normal Plot; Interactions; Main Effects; Orthogonal Arrays; Plackett–Burman Designs; Screening Designs.

MIKE JACROUX

Minimum Aberration

Introduction

Two-level **factorial** and **fractional factorial** designs are most widely used in experimental investigations. When m factors, each with two levels such as high (+1) and low (−1), are to be studied in an experiment, one may choose a full factorial design in which all possible 2^m level combinations of the factors are investigated. Such a level combination of the factors is often referred to as a *run* of the design. However, for a moderately large number of factors, a two-level full factorial design requires a large number of runs, which may be very time consuming, too expensive, and even redundant (see [1, p. 243]). Therefore, investigators often turn to a fractional factorial design, in which only a fraction of runs of a full factorial design are studied. In general, a 2^{m-p} fractional factorial design studies m factors with a 2^{-p} th fraction of 2^m runs of a full factorial design. A good fractional factorial design allows one to investigate many factors with relatively small run size, while still enabling one to estimate important effects. In order to select good fractional factorial designs for a given dimension (i.e., m and p), Box and Hunter [2] introduced the idea of design *resolution*. Fries and Hunter [3] further proposed a refined *minimum aberration* criterion to assess fractional factorial designs with the same resolution.

Minimum Aberration in Regular Designs

In this section we define regular two-level fractional factorial designs, describe how they alias certain effects and explain how the minimum aberration criterion can be helpful in selecting a good design. We illustrate the ideas by considering an experiment in which it is desired to study 7 factors in 32 ($= 2^{7-2}$) runs. Denote the seven factors by numbers 1, 2, 3, 4, 5, 6, 7, which correspond to seven columns with entries ± 1 in the design matrix. A 2^{7-2} fractional factorial design can be constructed by first writing down a basic design consisting of 32 runs associated with a full factorial design in five factors, say, 1, 2, 3, 4, 5. Then we associate the two additional factors 6 and 7 with appropriately chosen

interactions involving the factors 1, 2, 3, 4, 5. For example, we may obtain a design D_1 by associating factors 6 = 123 and 7 = 234, where the multiplication of letters is the element-wise multiplication of their corresponding columns. Hence, main effects 6 and 7 of design D_1 are aliased by its three-factor interactions 123 and 234, respectively. The assignment of factors 6 and 7 can be rewritten as $I = 1236$ and $I = 2347$, where I is a column of “+1”. One can find another relation among the factors as $I = 1236 \times 2347 = 1467$ using the fact that two same letters result in the identity column I . Altogether we have $I = 1236 = 2347 = 1467$, which forms the *defining relation* of the design. The elements 1236, 2347, and 1467 in the defining relation are referred to as *words*. The number of letters in a word is called the *length* of the word or *word length*. The words 1236 and 2347, used to define additional factors 6 and 7, are referred to as the *design generators*. In general, the design generators plus all the possible products of them constitute the complete defining relation of a design. Two-level fractional factorial designs 2^{m-p} constructed as above have the complete defining relation and are known as *regular fractional factorial designs* or *simply regular designs*.

In regular fractional factorial designs, some of the effects are completely entangled, or aliased, with one another. From the defining relation, it is straightforward to find the **aliasing** structures. For example, multiplying the above defining relation on both sides by 1 yields $1 = 236 = 467 = 12347$. That is, the **main effect** 1 is aliased with two three-factor interactions 236, 467 and a five-factor interaction 12347. Similarly, $12 = 36 = 2467 = 1347$. The **estimation** of the two-factor interaction 12 is actually the sum of effects $12 + 36 + 246 + 1347$. In the literature, when an effect A is aliased with another effect B , it is also said that the effect A is confounded with the effect B , which means these two effects are indistinguishable from each other. Design D_1 can estimate all main effects, but cannot estimate some two-factor interactions such as 12 and 36, even if we assume that the three-factor and higher-order interactions are negligible. However, all two-factor interactions involving 5 in D_1 are estimable. Therefore, the aliasing structure is an important information in planning industrial experiments.

Let $A_i(D)$ be the number of words of length i in the defining relation of a design D . The *word length*

pattern is then the vector $W(D) = [A_1(D), A_2(D), A_3(D), \dots]$. Since any two columns of a two-level fractional factorial design are orthogonal, the minimum length of a word will be 3; the maximum length possible is the total number of factors in the design. Thus, the word length pattern of design D_1 above can be shortened to $W(D_1) = [A_3(D_1), A_4(D_1), A_5(D_1), A_6(D_1), A_7(D_1)] = [0, 3, 0, 0, 0]$.

Box and Hunter [2] introduced *resolution* as a criterion for selecting regular fractional factorial designs. The resolution R of a regular fractional factorial design is defined as the length of the shortest word in the defining relation of the design. Therefore, if a regular fractional factorial design is of resolution R , then no k -factor interaction is aliased with any other interaction containing less than $R - k$ factors. The resolution of a regular design is denoted by a Roman numeral subscript. For instance, the design D_1 mentioned above is of resolution IV and is denoted as 2_{IV}^{7-2} . The **design resolution** is a criterion to measure the occurrences of worst aliases of a regular fractional factorial design. In general, a design with a higher resolution is better. However, there may be many regular fractional factorial designs with the same resolution, especially for larger designs with many factors. Resolution alone is sometimes not sufficient to distinguish among regular fractional factorial designs. Suppose a design D_2 is obtained by redefining factors 6 and 7 in design D_1 above as $6 = 1234$ and $7 = 1245$. So design D_2 has the defining relation $I = 12346 = 12457 = 3567$. Both designs D_1 and D_2 are of resolution IV, but they have rather different alias structures. If we assume that interactions involving three or more factors are negligible, which is often the case at the initial stage of the investigation, design D_2 has only three pairs of aliases (i.e., $35 = 67$, $36 = 57$, and $37 = 56$) with respect to the two-factor interactions, whereas there are nine pairs of such aliases in design D_1 . Clearly, design D_2 is a better choice for a 2_{IV}^{7-2} , if there is little knowledge about the factors.

In order to discriminate among regular fractional factorial designs with the same resolution, Fries and Hunter [3] proposed the *minimum aberration* criterion based on the word length pattern. Suppose for any two regular fractional factorial designs D and D' , r is the smallest positive integer such that $A_r(D) \neq A_r(D')$. Then design D is said to have less aberration than design D' if $A_r(D) <$

$A_r(D')$. If there exists no other regular design having less aberration than design D , then D is said to be a minimum aberration design. Minimum aberration is a criterion to compare the frequency of aliases of regular designs at different levels sequentially, and a regular design having the smallest frequency of worst aliases is the best. For example, the word length pattern for design D_2 is $W(D_2) = [A_3(D_2), A_4(D_2), A_5(D_2), A_6(D_2), A_7(D_2)] = [0, 1, 2, 0, 0]$. Since $A_3(D_1) = A_3(D_2) = 0$ and $A_4(D_2) = 1 < 3 = A_4(D_1)$, design D_2 has less aberration than design D_1 . It is easy to check that design D_1 is a minimum aberration design. Chen *et al.* [4] and Wu and Hamada [5, Chapter 4, Appendix] listed minimum aberration (and other selected) two-level regular fractional factorial designs from 8 to 128 runs. Some important contributions related to minimum aberration can be found in Franklin [6], Chen and Wu [7], Chen [8], Tang and Wu [9], Chen and Hedayat [10], Suen *et al.* [11], and Cheng *et al.* [12]. There are many additional works on minimum aberration designs; see of Wu and Hamada [5, Chapter 4, References] for details.

Nonregular Designs and Generalized Minimum Aberration

Nonregular fractional factorial designs, or simply nonregular designs, are commonly obtained from Plackett–Burman designs given in Plackett and Burman [13] or more generally from Hadamard matrices by selecting a subset of the columns. Recall that a Hadamard matrix of order n is an $n \times n$ matrix with the elements ± 1 whose columns (and rows) are orthogonal to each other. The order n is a multiple of 4. Other widely used nonregular designs are selected from orthogonal arrays. See Hedayat *et al.* [14] and Wu and Hamada [5] for further details. While regular designs have a clear defining relation and therefore simple aliasing structures, the gap of runs for regular factorial designs is becoming larger and larger (4, 8, 16, 32, 64, 128, ...). In contrast, nonregular designs have no clear defining relation and thus their aliasing structures are very complicated, but they have a very flexible run size of a multiple of 4. Hamada and Wu [15] argued that the nature of the complex aliasing structure of nonregular designs could be exploited to estimate some interactions without increasing the number of runs. They proposed an

analysis strategy to turn the liability of the complex aliasing structure into the virtue of model estimability. Therefore, good nonregular designs not only provide a much wider range of possible experiment sizes but may also have advantages in the estimation of interactions; see the illustrative example in Wang and Wu [16].

The popular minimum aberration criterion was first proposed for selecting regular designs, and it cannot be directly applied to nonregular designs. In order to assess nonregular designs, one needs to find a suitable criterion. There are several significant articles in this direction as reported by Lin and Draper [17], Wang and Wu [16], and Box and Tyssedal [18]. The criteria proposed were reasonable and heuristic. Most of them however, were highly model based. They had no apparent connection with the minimum aberration criterion or no rigorous theoretical justifications. Deng and Tang [19, 20] proposed *generalized resolution* and *generalized minimum aberration* (GMA) for ranking nonregular two-level designs in a systematic approach. These criteria are based on a generalization of the word length pattern and is useful for comparing regular/nonregular designs. We give a more detailed description next.

Let D be a two-level (regular or nonregular) factorial design with n runs and m factors. The design D can also be regarded as a set of m column vectors of length n . For any subset with k columns selected from D , say, $s = \{\mathbf{x}_1, \dots, \mathbf{x}_k\}$, define $J_k(s) = |\sum_{i=1}^n x_{i1}, \dots, x_{ik}|$, where x_{ij} is the i th component of column vector $\mathbf{x}_j$. $J_k(s)$ is the absolute value of the sum of the interaction column $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_k$.

For regular designs, $J_k(s)$ equals either 0 or n . When $J_k(s) = n$, it corresponds to a full confounding among these columns whose indices are in s . For nonregular designs, $J_k(s)$ can take values between 0 and n which means a partial confounding. For a design D , let r be the smallest integer such that $\max_{|s|=r} J_r(s) > 0$, where the maximization is over all the subsets of r distinct columns of D . Then the *generalized resolution* of D , proposed in Deng and Tang [19], is defined as $R(D) = r + [1 - \max_{|s|=r} J_r(s)/n]$. Obviously, $r \leq R(D) < r + 1$. For a regular design, $R(D) = r$ since $J_k(s)$ takes either 0 or n . Therefore, $\max_{|s|=r} J_r(s) > 0$ is equivalent to $\max_{|s|=r} J_r(s) = n$ for a regular design.

For nonregular designs, it is important to know the possible values of $J_k(s)$. It is straightforward to

see that $J_k(s)$ is a multiple of 4. Deng and Tang [20] derived a general formula for possible values of $J_k(s)$ and proved that the gap between two consecutive possible J values is a multiple of 8. We then list all possible values of $J_k(s)$ in decreasing order and denote the frequency distribution of $J_k(s)$ for all possible s out of k columns in D by $[f_{k1}, f_{k2}, \dots, f_{kg+1}]$, where f_{kj} is the frequency of $J_k(s)$ taking the j th largest possible value and $g = \lfloor \frac{n}{8} \rfloor$, the integer part of $\frac{n}{8}$. The total number of k -column subdesigns (*projections*) from an m -column design is a fixed number $\binom{m}{k}$ which is equal to $\sum_{j=1}^{g+1} f_{kj}$. Therefore, it is sufficient to consider the first g entries and let

$$F_k(D) = [f_{k1}, f_{k2}, \dots, f_{kg}] \quad (1)$$

The confounding frequency vector (or generalized word length pattern) of D , proposed by Deng and Tang [19, 20], is given by

$$W(D) = [F_3(D), F_4(D), \dots, F_m(D)] \quad (2)$$

Since only designs with orthogonal columns are considered, $F_1(D) = F_2(D) = 0$.

Note that for a regular design, f_{k1} is the number of words with length k and $f_{ki} = 0$ ($i > 1$) because $J_k(s)$ takes either 0 or n . Therefore, the corresponding confounding frequency vector reduces to the word length pattern for regular designs.

Any two fractional factorial designs D_1 and D_2 with the same run size and number of factors can be ranked by comparing their confounding frequency vectors element by element. Specifically, suppose r is the smallest integer such that $F_r(D_1) \neq F_r(D_2)$. Design D_1 is said to have less generalized aberration than design D_2 if $F_r(D_1) < F_r(D_2)$, which are compared element by element from the first entry to the last. In other words, we would sequentially compare the frequency of the components and prefer the one that minimizes the frequency of the largest value. As proposed first by Deng and Tang [19], a design is said to have *generalized minimum aberration* (also referred to as *minimum G aberration*) if there exists no other design having less generalized aberration. Clearly this criterion reduces to the traditional minimum aberration for regular designs. The GMA criterion can compare and assess any two factorial designs, whether or not they are regular.

Using the confounding frequency vector, one can rank various nonregular designs of m columns and n runs. Deng *et al.* [21] and Deng and Tang

[20] tabulated and ranked a collection of nonregular designs of run sizes 16, 20, and 24.

Other Related Criteria

Instead of counting the frequency of $J_k(s)$ in the confounding frequency vector $W(D)$, Tang and Deng [22] considered a relaxed version as $B_k(D) = n^{-2} \sum_{|s|=k} [J_k(s)]^2$ and proposed *minimum G_2 aberration* which is based on the vector $W_B(D) = [B_3(D), B_4(D), \dots, B_m(D)]$. They further developed a complementary design theory for the minimum G_2 aberration criterion. Cheng *et al.* [23] demonstrated that the criterion based on $W_B(D)$ is consistent with some model-dependent efficiency criteria under certain two-factor interaction models. Cheng and Tang [24] further developed a general theory of minimum aberration based on some statistical principles and their theory provided a unified framework for minimum aberration.

While the work on GMA shows a promising direction in the study of nonregular designs, this approach is mostly suitable for two-level designs. By studying treatment contrasts and ANOVA models, Xu and Wu [25] proposed a generalization of the minimum aberration criterion that can be used for multilevel designs. To develop an additional theory for the proposed criterion, they successfully explored the connection between the theory of factorial designs and the theory of coding. In other words, instead of studying the relationship between the factors (columns), one can investigate the relationship between the runs (rows). Following this approach, Xu [26] proposed the *minimum moment aberration* criterion, which is based on the power moments, for measuring the similarity among runs (rows). Minimizing the power moments makes runs as dissimilar as possible. Therefore, good designs should have small power moments.

Using power moments of the projection designs and counting their frequency distributions, Xu and Deng [27] proposed the *moment aberration projection* criterion to rank and classify nonregular designs (including multilevel designs). Two fractional factorial designs are *isomorphic* if one can be obtained from the other by permuting rows and columns and/or changing the signs of columns. In general, it is time consuming to perform an isomorphism check between two designs as demonstrated in Chen

et al. [4]. The confounding frequency vector and the moment aberration projection, although originally proposed to rank and assess designs, are also useful in testing for isomorphism. Clearly, two designs with different confounding frequency vectors or moment aberration projections are not isomorphic. On the other hand, two nonisomorphic designs may have the same confounding frequency vector or moment aberration projection. Xu and Deng [27] demonstrated that moment aberration projection has a better classification power than the confounding frequency vector. The ranking of various designs given by both criteria are fairly consistent, but not identical, with each other.

Fang and Mukerjee [28] found a connection analytically between the minimum aberration criterion and a uniformity criterion in the area of **uniform design**. See Fang *et al.* [29, Chapter 3] for more references.

References

- [1] Box, G.E.P., Hunter, W.G. & Hunter, J.S. (1978). *Statistics for Experimenters: An Introduction to Design, Data Analysis, and Model Building*, John Wiley & Sons, New York.
- [2] Box, G.E.P. & Hunter, J.S. (1961). The 2^{k-p} fractional factorial designs. *Technometrics* **3**, 311–351 and 449–458.
- [3] Fries, A. & Hunter, W.G. (1980). Minimum aberration 2^{k-p} designs, *Technometrics* **22**, 601–608.
- [4] Chen, J., Sun, D.X. & Wu, C.F.J. (1993). A catalogue of two-level and three-level fractional factorial designs with small runs, *International Statistical Review* **61**, 131–146.
- [5] Wu, C.F.J. & Hamada, M. (2000). *Experiments: Planning, Analysis and Parameter Design Optimization*, John Wiley & Sons, New York.
- [6] Franklin, M.F. (1984). Constructing tables of minimum aberration p^{n-m} designs, *Technometrics* **26**, 225–232.
- [7] Chen, J. & Wu, C.F.J. (1991). Some results on s^{n-k} fractional factorial designs with minimum aberration or optimal moments, *Annals of Statistics* **19**, 1028–1041.
- [8] Chen, J. (1992). Some results on 2^{n-k} fractional factorial designs and search for minimum aberration designs, *Annals of Statistics* **20**, 2124–2141.
- [9] Tang, B. & Wu, C.F.J. (1996). Characterization of minimum aberration 2^{n-k} designs in terms of their complementary designs, *Annals of Statistics* **24**, 2549–2559.
- [10] Chen, H. & Hedayat, A.S. (1996). 2^{n-m} fractional factorial designs with weak minimum aberration, *Annals of Statistics* **24**, 2536–2548.
- [11] Suen, C.-Y., Chen, H. & Wu, C.F.J. (1997). Some identities on q^{n-m} designs with application to minimum aberrations, *Annals of Statistics* **25**, 1176–1188.

-
- [12] Cheng, C.S., Steinberg, D.M. & Sun, D.X. (1999). Minimum aberration and model robustness for two-level fractional factorial designs, *Journal of the Royal Statistical Society, Series B* **61**, 85–91.
 - [13] Plackett, R.L. & Burman, J.P. (1946). The design of optimum multifactorial experiments, *Biometrika* **33**, 305–325.
 - [14] Hedayat, A.S., Sloane, N.J.A. & Stufken, J. (1999). *Orthogonal Arrays: Theory and Applications*, Springer-Verlag, New York.
 - [15] Hamada, M. & Wu, C.F.J. (1992). Analysis of designed experiments with complex aliasing, *Journal of Quality Technology* **24**, 130–137.
 - [16] Wang, J.C. & Wu, C.F.J. (1995). A hidden projection property of Plackett-Burman and related designs, *Statistica Sinica* **5**, 235–250.
 - [17] Lin, D.K.J. & Draper, N.R. (1992). Projection properties of Plackett and Burman designs, *Technometrics* **34**, 423–428.
 - [18] Box, G.E.P. & Tyssedal, J. (1996). Projective properties of certain orthogonal arrays, *Biometrika* **84**, 950–955.
 - [19] Deng, L.-Y. & Tang, B. (1999). Generalized resolution and minimum aberration criteria for Plackett-Burman and other nonregular factorial designs, *Statistica Sinica* **9**, 1071–1082.
 - [20] Deng, L.-Y. & Tang, B. (2002). Design selection and classification for Hadamard matrices using generalized minimum aberration criteria, *Technometrics* **44**, 173–184.
 - [21] Deng, L.-Y., Li, Y. & Tang, B. (2000). Catalogue of small runs nonregular designs from Hadamard matrices with generalized minimum aberrations, *Communications in Statistics: Theory and Methods* **29**, 1379–1395.
 - [22] Tang, B. & Deng, L.Y. (1999). Minimum G_2 -aberration for nonregular fractional factorial designs, *Annals of Statistics* **27**, 1914–1926.
 - [23] Cheng, C.S., Deng, L.Y. & Tang, B. (2002). Generalized minimum aberration and design efficiency for nonregular fractional factorial designs, *Statistica Sinica* **12**, 991–1000.
 - [24] Cheng, C.S. & Tang, B. (2005). A general theory of minimum aberration and its applications, *Annals of Statistics* **33**, 944–958.
 - [25] Xu, H. & Wu, C.F.J. (2001). Generalized minimum aberration for asymmetrical fractional factorial designs, *Annals of Statistics* **29**, 1066–1077.
 - [26] Xu, H. (2003). Minimum moment aberration for nonregular designs and supersaturated designs, *Statistica Sinica* **13**, 691–708.
 - [27] Xu, H. & Deng, L.Y. (2005). Moment aberration projection for nonregular fractional factorial designs, *Technometrics* **47**, 121–131.
 - [28] Fang, K.T. & Mukerjee, R. (2000). A connection between uniformity and aberration in regular fractions of two-level factorials, *Biometrika* **87**, 193–198.
 - [29] Fang, K.T., Li, R. & Sudjianto, S. (2005). *Design and Modeling for Computer Experiments*, Chapman & Hall/CRC.

LIH-YUAN DENG

Optimal Design

Introduction

With luck, and perhaps after data transformation, the analysis of a well-designed experiment may require little more than a few simple plots that reveal how the responses depend on the factors. Subsequent **estimation** of the parameters modeling this dependence usually requires the use of **least squares**. In a good experiment the variances and covariances of the estimated parameters will be small. Optimal experimental designs minimize functions of these variances and so provide good estimates of the parameters.

This article is mainly concerned with D-optimality in which the generalized variance of the parameter estimates is minimized. Many of the standard designs, such as the 2^m factorials and the series of 2^{m-f} fractional factorials, as well as standard designs for **mixture experiments** and some designs for second-order response surfaces can be justified because they are D-optimal, in addition to other justifications. However, a major feature of the methods of optimal experimental design is the ability to find good designs in nonstandard situations. Here are some advantages of the approach:

1. the availability of algorithms for the construction of designs;
2. the calculation of good designs for a specified number of trials;
3. the provision of simple, but incisive, methods for the comparison of designs;
4. the ability to divide response-surface and other designs into blocks of a specified size;
5. the determination of good response-surface designs over nonregular regions;
6. the construction of designs for mixture models, perhaps over nonregular regions;
7. the ability to find designs for nonstandard models, such as nonlinear models.

The first six of these seven points are illustrated in the sections that follow. Chapter 17 of Atkinson *et al.* [1] gives a thorough coverage of optimal designs for nonlinear models.

The article starts in the section titled “Simple Regression” with linear and quadratic regression in one variable. Several criteria of optimality are

introduced in the third section. Examples of numerically constructed D-optimal designs are given in the section titled “Second-Order **Response Surface** in Two Factors”. These include designs for an arbitrary number of trials, blocked response-surface designs and designs over nonregular design regions. The distinction between exact and continuous designs is made in the following section in “Exact and Continuous Designs” and used in the section titled “The General Equivalence Theorem and Design Efficiency” in the description of the general equivalence theorem for D- and G-optimal designs. The theorem is also used to demonstrate the optimality of designs and calculate efficiencies. Algorithms for design construction are in the section titled “Design Algorithms”. Extensions, including designs for detecting **lack of fit**, and references to the literature on optimal experimental design are in the section titled “Discussion and Literature”.

Simple Regression

Since optimal designs focus on the variances of the estimated parameters, we need to specify a model. We start with the simplest interesting case, linear regression in one variable, which serves to introduce some important ideas. There is a single response y , the expected value of which depends linearly on the value of the factor or process variable x through the relationship

$$E(y) = \beta_0 + \beta_1 x \quad (1)$$

The design problem is to choose N values of x at which measurements are to be taken. Although there are N observations, the number of distinct values of x used in the experiment, which we call n , may be much less than N . These n values of x must lie in the experimental region $\mathcal{X}$. For a single variable we can scale the values of x such that $-1 \leq x \leq 1$, where ± 1 correspond to the minimum and maximum values of the variable, which might typically be time, temperature, pressure, stirring speed, or catalyst concentration.

The model is to be fitted to the data by least squares. For this method to be appropriate, the variance of the observations should be constant and we write $\text{var}(y) = \sigma^2$. The least-squares estimators of the parameters in equation (1) are called $\hat{\beta}_0$ and $\hat{\beta}_1$ with

$$\text{var}(\hat{\beta}_1) = \frac{\sigma^2}{\sum (x_i - \bar{x})^2} \quad (2)$$

2 Optimal Design

where the sample average $\bar{x} = \sum x_i / N$. Here and in equation (2) the summations are over all N values x_i . From equation (2) the variance of $\hat{\beta}_1$ is minimized when $\sum (x_i - \bar{x})^2$ is maximized. For a fixed number of trials N this occurs when all trials are at $+1$ or -1 . If N is even, exactly half the trials will be at $x = +1$ and the other half at $x = -1$. This design has $\bar{x} = 0$, so that $\text{var}(\hat{\beta}_0) = \sigma^2 / N$ which does not depend any further on the design. This design is therefore optimal for both parameters.

This design has many of the features of an optimal design. The number of support points of the design n , here two, may be appreciably less than the number of trials N . The fine structure of the design depends on the values of N , here whether it is even or odd, in which case one more trial is put at one end of the design region; it does not matter which. A third feature is that the design covers the whole of the design region, using the most extreme values of x available. In any application of the design, the order in which the two treatments $x = -1$ and $x = 1$ are applied should be randomized (*see Randomization in Experimental Designs*), to avoid any confounding with omitted variables. For example, if experiments are performed over time, gradual fouling of catalysts may mean that yields slowly decline during the experiment. If all trials at one level of x were performed in the second half of the experiment, the decline in yield would be confounded with any effect of the factor. Joiner [2] gives further examples of such “lurking variables”.

Fitted models are often used for prediction. The prediction at the point x , not necessarily included in the data from which the parameters were estimated, is

$$\hat{y}(x) = \hat{\beta}_0 + \hat{\beta}_1 x = \bar{y} + \hat{\beta}_1 (x - \bar{x}) \quad (3)$$

with variance

$$\text{var}\{\hat{y}(x)\} = \sigma^2 \left\{ \frac{1}{N} + \frac{(x - \bar{x})^2}{\sum (x_i - \bar{x})^2} \right\} \quad (4)$$

The optimal design found in this section for parameter estimation also has a minimax property for prediction: it minimizes the maximum variance of the prediction $\hat{y}(x)$ over $\mathcal{X}$. We discuss design for prediction further in the section titled “Criteria of Optimality”.

Optimal designs are tailored to the model that is assumed; particularly when there is one factor, the

designs often have the same number of design points as there are parameters. An example is the design at two points for the two-parameter regression model. An objection to such a design is that it does not provide for checking the model, for which trials at three or more values of x are needed. This suggests that the design should also be efficient in case the relationship is a second-order model

$$E(y) = \beta_0 + \beta_1 x + \beta_2 x^2 \quad (5)$$

It is always possible that an even-higher-order model is required, but empirical evidence from the results of decades of experimentation has shown that third- or higher-order models are rarely needed. If such models seem to be required, data transformations, such as the power transformations analyzed by **Box** and Cox [3], are often called for.

Polynomial models such as equations (1) and (5) can often be thought of as Taylor series approximations of some underlying more complicated model. In such cases, relatively simple polynomial models may provide good predictions of the response, the **bias** of the approximation being offset by a low variance due to a model with few estimated parameters. The trade-off between bias and variance and its effect on the design of experiments is one of the central topics of response-surface methodology (*see Response Surface Methodology*).

The design for the three-parameter second-order model, equation (5), introduces further aspects of the dependence of the optimal design on the precise purpose of experimentation. Unlike the two-point design for the first-order model, there is now no overall optimal design for all parameters. Testing whether the first-order model is adequate is equivalent to testing whether $\beta_2 = 0$. The variance of $\hat{\beta}_2$ is minimized by a design that puts half the trials at $x = 0$ and divides the remaining half between $x = -1$ and $x = 1$. But the design that provided the best estimates of all three parameters, in the sense of D-optimality described in the next section, puts one-third of the trials at each of these three points. In a satisfactory manner, the optimal design depends on the precise purpose of the experiment.

Criteria of Optimality

This section mainly describes the criterion of D-optimality which provides designs minimizing the generalized variance of the estimated parameters.

We consider a general experiment in which there are m factors and write the linear model as

$$E(y) = F\beta \quad (6)$$

Here y is the $N \times 1$ vector of responses, β is a vector of p unknown parameters, and F is the $N \times p$ extended design matrix. The i th row of F is $f^T(x_i)$, a known function of the m explanatory variables. For the quadratic model equation (5), for which $m = 1$ and $p = 3$,

$$f^T(x_i) = (1 \quad x_i \quad x_i^2) \quad (7)$$

The design is determined by the experimental values of x . The extended design matrix reflects not only the design but in addition the model to be fitted.

We assumed in the section titled ‘‘Simple Regression’’ that the observational errors were independent with constant variance σ^2 . This value usually does not depend on the design. However, failure to block the experiment (*see Blocking*), when there are block effects, can lead to an increase in the estimate of σ^2 and a loss of power of tests about parameter values.

The least-squares estimator of the parameters is

$$\hat{\beta} = (F^T F)^{-1} F^T y \quad (8)$$

where the $p \times p$ matrix $F^T F$ is the information matrix for β . The larger $F^T F$, the greater is the information in the experiment.

With σ^2 constant, the **covariance** matrix of the least-squares estimator is

$$\text{var } \hat{\beta} = \sigma^2 (F^T F)^{-1} \quad (9)$$

The variance of $\hat{\beta}_j$ is proportional to the j th diagonal element of $(F^T F)^{-1}$ with the covariance of $\hat{\beta}_j$ and $\hat{\beta}_k$ proportional to the (j, k) th off-diagonal element. If we are interested in the comparison of experimental designs, the value of σ^2 is not relevant, since the value is the same for all proposed designs for a particular experiment.

Confidence regions for all p elements of β are of the form

$$(\beta - \hat{\beta})^T F^T F (\beta - \hat{\beta}) \leq k \quad (10)$$

In the p -dimensional space of the parameters, equation (10) defines an ellipsoid, the boundary of which is a contour of constant residual **sum of squares**.

The volume of the ellipsoid is inversely proportional to the square root of the determinant $|F^T F|$. From the expression for $\text{var } \hat{\beta}$ in equation (9), $\sigma^2 |(F^T F)^{-1}| = \sigma^2 / |F^T F|$ is called the *generalized variance of $\hat{\beta}$* . Designs which maximize $|F^T F|$ minimize this generalized variance and are called *D-optimal* (for Determinant).

Interest in the fitted model may be not only in the parameters, but also in the quality of the predictions. The predicted value of the response for simple regression is given by equation (3). For the general model with p parameters, equation (6), the prediction can be written

$$\hat{y}(x) = \hat{\beta}^T f(x) \quad (11)$$

with variance

$$\text{var}\{\hat{y}(x)\} = \sigma^2 f^T(x) (F^T F)^{-1} f(x) \quad (12)$$

Designs which minimize the maximum over $\mathcal{X}$ of $\text{var } \hat{y}(x)$ are called *G-optimal*. As the theorem of the section titled ‘‘The General Equivalence Theorem and Design Efficiency’’ indicates, as N increases, D-optimality and G-optimality give increasingly similar designs.

An alternative to this minimax approach to $\text{var } \hat{y}(x)$ is to find designs that minimize the average value of the variance over $\mathcal{X}$. Such designs are variously called *I-optimal* or *V-optimal* from Integrated and Variance. Another criterion of importance in applications, particularly for nonlinear models, is that of *c-optimality* in which the variance of the linear combination $c^T \hat{\beta}$ is minimized, where c is a $p \times 1$ vector of constants. We do not exhibit such designs in this article, but focus on D- and G-optimality. Details of these and other design criteria are in Chapter 10 of Atkinson *et al.* [1].

Second-Order Response Surface in Two Factors

Many standard designs which have been derived over the years by a variety of methods share the property that they are D- and G-optimal. One example is the series of 2^m factorial designs (*see Factorial Experiments*) and their symmetric fractions forming the 2^{m-f} fractional factorial designs (*see Fractional Factorial Designs*), which can also be used to construct D-optimally blocked 2^m designs. Some

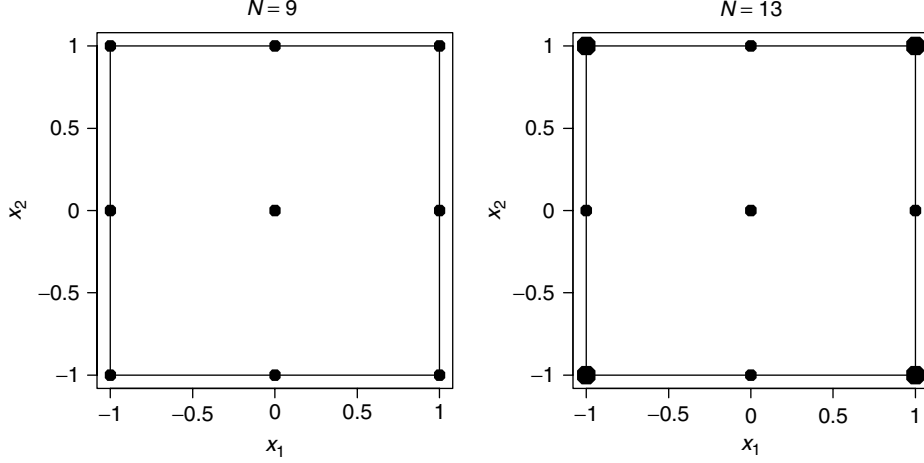


Figure 1 Second-order response surface in two factors: D-optimal designs for number of trials $N = 9$ and 13; heavy dots indicate support points with two observations

properties of the designs are discussed as Example 6. But these excellent properties do depend on the number of trials N . If this is not a power of 2 then only some of the 2^m points of the full factorial will be selected as design points or to be replicated, depending on whether N is less than or greater than 2^m ; the methods of optimal design construction will then be needed to determine these points. Similarly, the designs depend on the design region $\mathcal{X}$. If not all 2^m combinations of the high and low (+1 and -1) values of the factors are available, designs will need to be calculated over this irregular region. This section presents examples of the construction of designs for arbitrary N , for nonregular design regions and also for **blocking**. In the case of the blocked 2^m factorial, we might be forced by, for example, the size of batches of raw materials, to have blocks which are not all of size 2^{m-f} .

We now use the second-order polynomial in two factors to demonstrate some of the properties of D-optimal designs and of the versatility of the approach based on the theory of design optimality. The model is

$$E(y) = \beta_0 + \beta_1 x_1 + \beta_2 x_2 + \beta_{11} x_1^2 + \beta_{22} x_2^2 + \beta_{12} x_1 x_2 \quad (13)$$

Although optimal designs depend on the model, the designs we find also have high efficiency for models in which some of the terms in equation (13) are absent.

Initially we take $\mathcal{X}$ as the square region $[-1, 1]^2$. Numerical searches for the design maximizing $|F^T F|$ will be confined to a grid of N_C candidate points, rather than to searches over the continuous region. Very little is lost in terms of the value of $|F^T F|$, the maximization is simplified and a grid of values corresponds to much experimental practice where settings are rounded to convenient values, such as 5°C intervals of temperature. The algorithms used to find the designs are outlined in the section titled “Design Algorithms”.

Example 1: Dependence of D-Optimal Design on N When $N = 9$ the D-optimal design maximizing $|F^T F|$ is the 3^2 factorial with trials at all combinations of the values -1, 0, and 1 for x_1 and x_2 . This design is plotted in Figure 1(a). For $N = 13$, the 3^2 design is augmented by the 2^2 factorial at ± 1 , that is by replicating the corner points of the 3^2 factorial, as shown in Figure 1(b). This is distinct from the conventional central composite design (see **Central Composite Designs**) in which the center point, that is the point $x_1 = x_2 = 0$ is replicated (see **Center Points**). In either design the replicated observations can be used to provide a model free estimate of the error variance.

For this second-order model, p , the number of parameters is six, which is therefore the minimum value of N . The D-optimal design for $N = 6$ has three design points that belong to the 3^2 factorial and

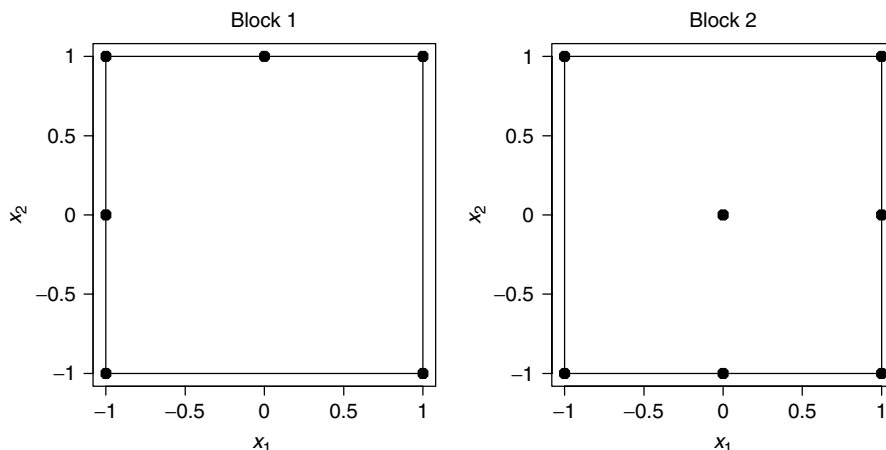


Figure 2 Second-order response surface in two factors. Optimal design for blocks of 6 and 7 trials, a total of $N = 13$ trials

three that do not. This design is only slightly better than the best design supported at six points of the 3^2 . As N increases, the number of points of the optimal design that are not factorial points decreases and the efficiency of fractions of the factorial increases. Some details, together with the definition of design efficiency, are in the section titled “The General Equivalence Theorem and Design Efficiency”.

Example 2: Blocking We now find designs for the second-order model (13) in which the trials are divided into two blocks of predetermined size. We assume additive blocking, so the model in equation (13) is augmented by an indicator variable for blocks at two levels. As an example, Figure 2 shows the division of 13 trials into two blocks, one of size 6 and the other of size 7.

This design is one of several with the same value of $|F^T F|$, including some designs from searching over the points of the 5^2 factorial in which the 3^2 factorial is augmented by points with at least one coordinate at ± 0.5 . The design of Figure 2 is symmetrical in x_1 and x_2 . But, in general, the labeling of the variables can be interchanged, as can the order of the blocks. This design has the advantage that its projection is the design for $N = 13$ in Figure 1(b). Thus the design is D-optimal whether or not blocking is required.

Example 3: Nonregular Design Region We consider a nonregular design region in which some

extreme values of x_1 and x_2 are excluded; in a chemical process these might correspond to conditions so severe that the product is degraded, or so dilute or cold that the desired reactions do not occur at a significant rate. Such a design region is shown in Figure 3. Optimal designs were found by searching over the 63 points of the 11^2 factorial that are included in the design region.

Figure 3 gives designs for $N = 9$ and 13. Although there are some other designs that are not far from these designs in efficiency, it is hard to summon up any intuition that leads to efficient designs; the algorithmic methods of optimal experimental design seem essential for the construction of efficient designs.

Finally, Figure 4 shows the D-optimal design when the 13 trials are divided into two blocks of sizes 4 and 9. In this case the projection of the two designs does not quite give the optimal unblocked design for $N = 13$ in Figure 3.

Example 4: Mixture Experiments In mixture experiments (*see Mixture Experiments*) the response depends on the proportion of the various components, but not on their amount. This imposes constraints on the values of the m factors and the resulting experimental region is a simplex in $m - 1$ dimensions.

The statistical study of mixture experiments was introduced by Scheffé [4]. Cornell [5] gives a book-length treatment. Because of the constraints on the

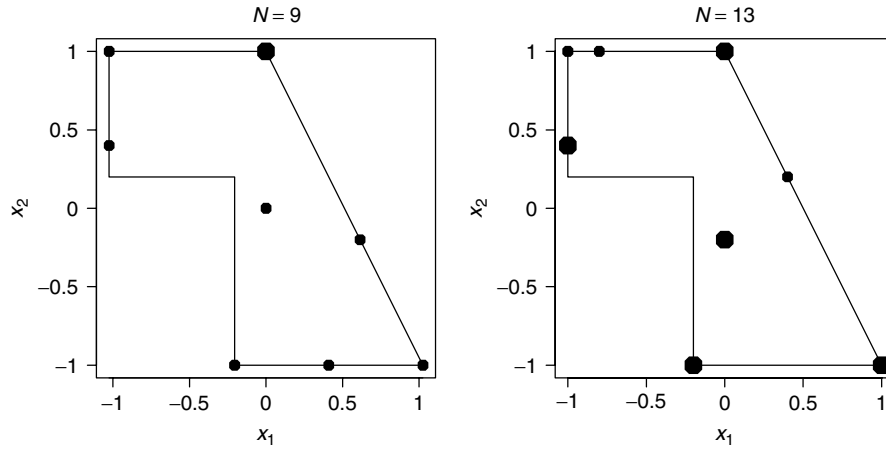


Figure 3 Second-order response surface in two factors with a nonregular design region: D-optimal designs for number of trials $N = 9$ and 13; heavy dots indicate support points with two observations

factors, not all terms in standard polynomial models are estimable. However, optimal design methods provide D-optimal designs for Scheffé's linear models. Chapter 16 of Atkinson *et al.* [1] gives the details of the designs, including blocked designs similar to those of Figure 4.

It is often necessary to introduce further constraints on the values of the factors to ensure that all mixtures in the experimental region are of interest; Martin *et al.* [6] give examples in glass manufacture. The resulting design regions can be hard to

visualize. However, once the design region has been given a mathematical description, optimal designs can readily be found by methods similar to those illustrated above for experiments involving response-surface models.

Exact and Continuous Designs

The D-optimal designs found so far have been for a specified model, design region $\mathcal{X}$ and number of trials N . As Figures 1 and 3 show, the structure of the

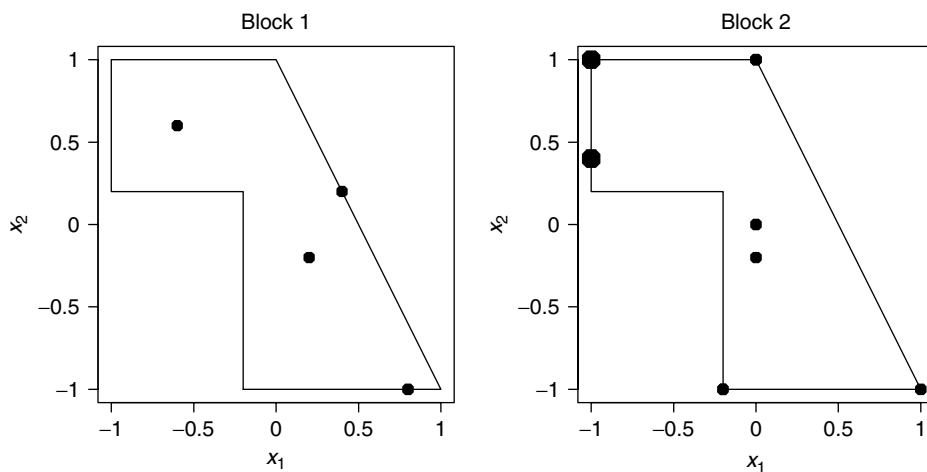


Figure 4 Second-order response surface in two factors with a nonregular design region: D-optimal design for blocks of 4 and 9 trials, a total of $N = 13$ trials; heavy dots indicate support points with two observations

designs can depend on the value of N . However, as N increases, the dependence decreases and the series of *exact* designs for specific values of N tends to a limiting *continuous* design independent of N . This section describes some uses of continuous optimal designs: they are often easier to find than an exact design for one specific N and they can readily be approximated to provide good exact designs for many values of N . The next section shows how continuous designs can be checked for optimality. Implications for algorithms for finding exact designs are given in the section titled “Design Algorithms”.

The theory was introduced in Kiefer [7] and developed in a series of further papers. Kiefer’s work on optimal design is collected in Brown *et al.* [8]. A basic idea is to write designs as probability measures.

Suppose an exact design has n points of support with r_i the integer number of trials at experimental conditions x_i . This design can be written as

$$\xi_N = \left\{ \begin{array}{cccc} x_1 & x_2 & \dots & x_n \\ r_1/N & r_2/N & \dots & r_n/N \end{array} \right\} \quad (14)$$

where $\sum r_i = N$. The measure ξ_N specifies the experimental conditions and the proportion of experimental effort at each condition. For example, for a response-surface design in m factors, the vector x_i gives the values of all m factors for the r_i replicate observations at the i th set of conditions; often for an exact design r_i will be 1.

The information matrix introduced in equation (8) can be written in terms of this measure as

$$F^T F = N M(\xi_N) \quad (15)$$

Likewise, the variance of the predicted response, introduced in equation (12), can be written in the standardized form

$$d(x, \xi_N) = \frac{N \text{var}\{\hat{y}(x)\}}{\sigma^2} = f^T(x) M^{-1}(\xi_N) f(x) \quad (16)$$

The continuous design ξ is found by replacing the fractions r_i/N in ξ_N by weights w_i . At the end of the section titled “Simple Regression” it was stated that the D-optimal continuous design for the quadratic model in one variable with $\mathcal{X} = [-1, 1]$ was

$$\xi^* = \left\{ \begin{array}{ccc} -1 & 0 & 1 \\ 1/3 & 1/3 & 1/3 \end{array} \right\} \quad (17)$$

where ξ^* , rather than ξ , indicates optimality. This design puts one-third of the trials at each of the three levels of x ; $-1, 0$, and 1 . For any N that is a multiple of three we can find the exact optimal design by putting $N/3$ trials at each of these sets of experimental conditions. For other values of N a good design is to distribute the trials as evenly as possible over these three points.

The General Equivalence Theorem and Design Efficiency

Exact D-optimal designs such as those of the section titled “Second-Order Response Surface in Two Factors”, maximize $|F^T F|$. From equation (15) D-optimal continuous designs therefore maximize $|M(\xi)|$. Kiefer and Wolfowitz [9] provide an equivalence theorem relating D- and G-optimal designs that provides a method of checking the optimality of any continuous design.

If ξ^* is a D-optimal design measure, then ξ^* also minimizes the maximum over $\mathcal{X}$ of the standardized variance of prediction $d(x, \xi)$ given by equation (16); that is ξ^* is also G-optimal. If we write this maximum as

$$\bar{d}(\xi) = \max_{x \in \mathcal{X}} d(x, \xi) \quad (18)$$

then $\bar{d}(\xi^*) = p$, the number of parameters of the linear model. For any nonoptimal design ξ the maximum over $\mathcal{X}$ of $d(x, \xi)$ is greater than p , which provides a method for checking for D-optimality.

The theorem holds for continuous optimal designs. Exact optimal designs such as that of equation (16) for the quadratic model when N is a multiple of 3 will clearly also satisfy it. But, for example, the exact D-optimal design when $N = 9$ of Figure 1 will not satisfy the theorem.

The equivalence theorem provides a **benchmark** for the comparison of designs. To compare a continuous design ξ with the optimal design ξ^* we define the D-efficiency as

$$D_{\text{eff}} = \left\{ \frac{|M(\xi)|}{|M(\xi^*)|} \right\}^{1/p} \quad (19)$$

The comparison of information matrices for designs that are measures removes the effect of the number of observations. Taking the p th root of the ratio of the determinants in equation (19) results in an

efficiency measure which has the dimensions of a variance, irrespective of the dimension of the model. Since the variance of estimated regression coefficients is inversely proportional to the number of observations, two replicates of a design measure for which $D_{\text{eff}} = 50\%$ would be as efficient as one replicate of the optimal measure.

Example 5: Quadratic Model in One Variable

The D-optimal continuous design for the quadratic model in one variable with $\mathcal{X} = [-1, 1]$ is given in equation (16). The claim that this design is D-optimal can be checked using the general equivalence theorem through calculation of the variance of the predicted response.

For this design $d(x, \xi^*)$ is the quartic given by

$$d(x, \xi^*) = 3 - \frac{9x^2}{2} + \frac{9x^4}{2} \quad (20)$$

which has a maximum over $\mathcal{X}$ of 3 at $x = -1, 0$, or 1 , the three design points. Further, this maximum value is equal to the number of parameters p . The plot in Figure 5 exhibits the optimality of the design since the values of three at the design points are indeed the optima over $\mathcal{X}$.

Example 6: 2^m Factorial and Fractions The first-order model in m factors with main effects and all

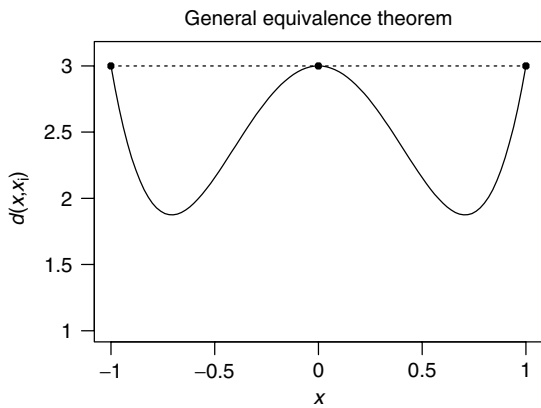


Figure 5 Quadratic model in one variable; $d(x, \xi^*)$ over $\mathcal{X}$ for the design of equation (16) which puts equal weights at $x = -1, 0$, and 1 . Since the maximum value is three, which is the number of parameters in the model, the design is D-optimal

possible interactions may be written as

$$E(y) = \beta_0 + \sum_{i=1}^m \beta_i x_i + \sum_{i=1}^{m-1} \sum_{j=i+1}^m \beta_{ij} x_i x_j + \cdots + \beta_{ijk \dots} x_i x_j x_k \cdots \quad (21)$$

When $\mathcal{X}$ is cuboidal, the D-optimal continuous design is the equireplicated 2^m factorial, which has $n = 2^m$ points of support at all combinations of values of ± 1 for each factor. If the full model (21) is fitted $p = N$ and $M(\xi^*)$ is the identity matrix. Then $d(x, \xi^*) = p$ at the design points and the equivalence theorem for continuous D-optimality is satisfied. Because of the orthogonality of the design, reflected in the diagonal nature of $M(\xi^*)$, the same design is optimal if terms are omitted from equation (21) or if the design is projected by the omission of one or more factors. Provided all the terms in the model can be estimated, the regular fractions of the design, that is the series of 2^{m-f} fractional factorials, are also the continuous D-optimal designs.

It is clear when analyzing the properties of this design that the maximum values of $d(x, \xi^*)$ occur at the points $x_i = \pm 1$. However, with more complicated examples, it is prudent to search $\mathcal{X}$ more carefully. Suppose the design used was a shrunken 2^m factorial with $x_i = \pm c$, $c < 1$. Then $d(x, \xi)$ would again equal one at the points of support of the design. However, this would not be the optimal design, since these would not be the maximum values of $d(x, \xi)$ over $\mathcal{X}$.

Example 7: Second-Order Response Surface in m Factors The second-order polynomial in m factors is

$$E(y) = \beta_0 + \sum_{j=1}^m \beta_j x_j + \sum_{j=1}^{m-1} \sum_{k=j+1}^m \beta_{jk} x_j x_k + \sum_{j=1}^m \beta_{jj} x_j^2 \quad (22)$$

Farrell *et al.* [10] show that the D-optimal continuous design for a cuboidal experimental region is supported on subsets of the points of the 3^m factorial, with the members of each subset having the same number of nonzero coordinates. Only three subsets are required, over each of which a specified design weight is uniformly distributed. For $m \leq 5$ the subsets have support $\{[0], [m-1], [m]\}$, that is the

center point, the midpoints of edges, and the corner points of the 3^m factorial. For larger values of m , other subsets also support optimal designs. It is interesting to note that the central composite designs ([11]; see **Central Composite Designs**), which belong to the family ($[0]$, $[1]$, $[m]$), cannot provide the support for a D-optimal continuous design when $m > 2$.

The number of support points of these continuous D-optimal designs increases rapidly with m , being 113 for $m = 5$. Exact designs with drastically reduced support are therefore required. Since the D-optimal continuous designs have support on the points of the 3^m factorial, it is reasonable to expect that good exact designs for small N can be found by searching over the points of these factorials. Table 11.7 of Atkinson *et al.* [1] summarizes designs found by such searches.

The minimal design has $N = p$. For $2 \leq m \leq 5$, these design all have efficiencies above 85% but designs with appreciably higher efficiencies can be found at the cost of only a few more design points. For $m = 3$, when $p = 10$, one design with $N = 14$ has an efficiency of 97.59%. When $m = 4$, $p = 15$ and a design with $N = 18$ has an efficiency of 93.11%. Finally, for $m = 5$ and $p = 21$, the best design with $N = 26$ is 95.19% efficient. The addition of one or a few trials above $N = p$ causes an appreciable increase in the efficiency of the design, in addition to the reduced variance of parameter estimates coming from a larger design with larger N .

Section 11.5 of Atkinson *et al.* [1] tabulates both continuous and exact D-optimal response-surface designs including designs for spheroidal regions. Chapter 15 of the same book discusses the blocking of response-surface designs.

Example 8: Second-Order Response Surface in Two Factors The designs in Figure 1 for the two-factor response surface model when $N = 9$ and 13 illustrate the dependence of exact optimal designs on the value of N . The continuous D-optimal design, like these two designs, is supported on the points of the 3^2 factorial with weight 0.1458 at each corner point $(\pm 1, \pm 1)$, weight 0.0802 at each center points of the edges $(0, \pm 1; \pm 1, 0)$ and weight 0.0960 at the center point of the design. Relative to this continuous design the exact design for $N = 9$ has a D-efficiency of 97.40%, whereas the value for $N = 13$ is 99.77%.

The smallest design, the exact D-optimal design for $N = 6$, has three support points that are not those of the 3^2 factorial and a D-efficiency of 89.15% relative to the continuous optimal design. The best design on six points of the 3^2 factorial has a very slightly lower efficiency of 88.49%. The best designs of both types for $N = 7$ and 8 already have efficiencies close to, or above, 95%. The conclusion is that, even for these very small designs, fractions of the 3^2 factorial provide efficient designs, as they do for higher values of m in Example 7. Application of the methods of optimal experimental design allows us to quantify the loss involved in experimenting at only these three levels.

Design Algorithms

Algorithms for the construction of continuous designs can either use relatively straightforward function maximization or be based on the general equivalence theorem through the properties of the variance $d(x, \xi)$. These show that the optimal design contains only points with the highest possible value of $d(x, \xi)$, namely p .

In the sequential construction of continuous D-optimal designs, observations are added at the value of x in $\mathcal{X}$ for which this variance is a maximum; $\mathcal{X}$ can either be treated as a continuous region or discretized to give a set of N_C candidate points. For an N -trial design written as the measure ξ_N , let the measure $\bar{\xi}_N$ put unit mass at the point where this maximum occurs. Then, provided $\bar{d}(\xi) > p$, we take

$$\xi_{N+1} = \frac{(N\xi_N + \bar{\xi}_N)}{(N+1)} \quad (23)$$

If $\bar{d}(\xi) = p$, the equivalence theorem of the section titled “The General Equivalence Theorem and Design Efficiency” shows that the design is D-optimal and the algorithm stops.

Convergence is not monotonic. As an instance, starting from the nine-trial design of Figure 1 for the two-factor response-surface model, we find that the variance is maximum at the corners of the region and one trial would be added at one of these corners. Continuing the procedure with $N = 10$ would lead to the addition of a trial at another corner point, and so on, until the design for $N = 13$ is obtained when $\bar{d}(\xi_{13})$ is very close to p . Addition of one further trial must destroy the symmetry of this design and

there will be a slight decrease in D-efficiency and an increase in the value of $\bar{d}(\xi)$, with $\bar{d}(\xi_{14}) > \bar{d}(\xi_{13})$. However, continuing with the sequential addition of trials in this way will eventually lead to a design with the weights given in the first paragraph of Example 8.

The convergence of such algorithms can be slow, but ultimate convergence is usually not required. Starting from an arbitrary design, the algorithm typically adds mass at a restricted number of conditions in $\mathcal{X}$. These conditions suggest the support of the optimal design. Once this pattern becomes clear, either the arbitrary starting design can be rejected and the algorithm restarted or the information gained can be used to guide maximization of the function $|M(\xi)|$. The advantage of the sequential algorithm is that it reduces the search for an optimal design to a sequence of optimizations in m dimensions.

Exact designs for N trials are usually found using exchange algorithms. It now follows from the equivalence theorem that the n support points of the design should have high values of $d(x, \xi)$ and the $N_C - n$ unused candidate points should have low values. Exchange algorithms explore the replacement of one of the N trials of the design by any of the other $N_C - 1$ candidate points, thereby allowing for a possible increase in replication of points already in the design. In the original algorithm [12] all $(N_C - 1) \times n$ exchanges were considered and the one leading to the largest increase in $|M(\xi_N)|$ made. At this point n may change, although N is constant. The process was then repeated until no further improvement was found. More recent algorithms [13, 14] independently order the values of $d(x, \xi_N)$ for the candidate and support points and make the first exchanges in the ordered lists that increases $|M(\xi_N)|$. Again the process is repeated until no further improvement is found. Whatever algorithm is used, the search is repeated several times from random starting points.

Computer procedures for finding D-optimal designs of the kind given here, such as Example 3, accordingly require the ability to specify a model and a list of candidate points. Commercial packages with this flexibility include Design Expert, JMP, SAS, and WebDOE. The R package AlgDesign is downloadable from <http://cran.r-project.org> and the Gosset package of Hardin and Sloane [15] is at <http://www.research.att.com/~njas/gosset/>.

Discussion and Literature

The motivation for optimal designs in the introduction to this article was that they provide a means of minimizing functions of the variances of the estimated parameters. Box [16, Chapter 9] argues that it is doubtful if single criterion optimal designs are useful in locating designs satisfying the compromises between objectives required in a practical experiment. Although the focus of this article is on D-optimality, the theory and algorithms extend straightforwardly to compound designs in which a weighted combination of criteria can be employed to find a design that is, for example, simultaneously good for parameter estimation, establishing a response transformation, and for prediction in a specified region. A general approach to compound designs is given in Chapter 21 of Atkinson *et al.* [1].

The list of 14 characteristics of a good experimental design in Box and Draper [17] includes the ability to check models for lack of fit. The early stages of an experimental program often use 2^m factorials and their fractions to screen factors. To check whether a second-order model such as equation (13) is required, one or more trials are included at the center of the experimental region. An informal extension of this procedure is to augment the optimal design by inclusion of a few extra trials in unused regions of $\mathcal{X}$. A more formal approach to designs for model checking, using the methods of optimal design, is given by DuMouchel and Jones [18], who specify both a primary model, for which good estimation of the parameters is required, and secondary terms, which need to be detected if they are nonzero. It is hoped that the primary model will provide an adequate fit. Examples of designs and the exploration of the relationship with compound D-optimality are given in Chapter 20 of Atkinson *et al.* [1].

The first book in English on optimal experimental design was Fedorov [12]. Pukelsheim [19] focuses on theory, as does the shorter book of Fedorov and Hackl [20]. Atkinson *et al.* [1] is more concerned with applications and provides SAS code for the construction of many designs. In addition to the linear models covered in this article, methods are given for the optimal design of experiments for correlated data, generalized linear models and, in particular, in Chapter 17, for the design of experiments for nonlinear models. The papers in Berger and Wong [21] focus on recent applications of optimal design,

including some in the pharmaceutical industry. As with any topic, a large amount of information can be obtained by searching the Web. About 75% of the references on Google are to “optimal design” while 25% use the term *optimum design*.

References

- [1] Atkinson, A.C., Donev, A.N. & Tobias, R.D. (2007). *Optimum Experimental Designs, with SAS*, Oxford University Press, Oxford.
- [2] Joiner, B.L. (1981). Lurking variables: some examples, *American Statistician* **35**, 227–233.
- [3] Box, G.E.P. & Cox, D.R. (1964). An analysis of transformations (with discussion), *Journal of the Royal Statistical Society. Series B* **26**, 211–246.
- [4] Scheffé, H. (1958). Experiments with mixtures, *Journal of the Royal Statistical Society. Series B* **20**, 344–360.
- [5] Cornell, J.A. (1990). *Experiments with Mixtures*, 2nd Edition, John Wiley & Sons, New York.
- [6] Martin, R.J., Bursnall, M.C. & Stillman, E.C. (2001). Further results on optimal and efficient designs for constrained mixture experiments, in *Optimal Design 2000*, A.C. Atkinson, B. Bogacka & A. Zhigljavsky, eds, Kluwer Academic Publishers, Dordrecht, pp. 225–239.
- [7] Kiefer, J. (1959). Optimum experimental designs (with discussion), *Journal of the Royal Statistical Society. Series B* **21**, 272–319.
- [8] Brown, L.D., Olkin, I., Sacks J. & Wynn, H.P. (eds) (1985). *Jack Carl Kiefer Collected Papers III*, John Wiley & Sons, New York.
- [9] Kiefer, J. & Wolfowitz, J. (1959). Optimum designs in regression problems, *Annals of Mathematical Statistics* **30**, 271–294.
- [10] Farrell, R.H., Kiefer, J. & Walbran, A. (1968). Optimum multivariate designs, in *Proceedings of the 5th Berkeley Symposium*, University of California Press, Berkeley, Vol. 1, pp. 113–138.
- [11] Box, G.E.P. & Draper, N.R. (1963). The choice of a second order rotatable design, *Biometrika* **50**, 335–352.
- [12] Fedorov, V.V. (1972). *Theory of Optimal Experiments*, Academic Press, New York.
- [13] Cook, R.D. & Nachtsheim, C.J. (1980). A comparison of algorithms for constructing exact optimal designs, *Technometrics* **22**, 315–324.
- [14] Atkinson, A.C. & Donev, A.N. (1989). The construction of exact D-optimum experimental designs with application to blocking response surface designs, *Biometrika* **76**, 515–526.
- [15] Hardin, R.H. & Sloane, N.J.A. (1993). A new approach to the construction of optimal designs, *Journal of Statistical Planning and Inference* **37**, 339–369.
- [16] Box, G.E.P. (2006). *Improving Almost Anything*, 2nd Edition, John Wiley & Sons, New York.
- [17] Box, G.E.P. & Draper, N.R. (1975). Robust designs, *Biometrika* **62**, 347–352.
- [18] DuMouchel, W. & Jones, B. (1994). A simple Bayesian modification of D-optimal designs to reduce dependence on an assumed model, *Technometrics* **36**, 37–47.
- [19] Pukelsheim, F. (1993). *Optimal Design of Experiments*, John Wiley & Sons, New York.
- [20] Fedorov, V.V. and Hackl, P. (1997). *Model-Oriented Design of Experiments, Lecture Notes in Statistics, Vol. 125*, Springer-Verlag, New York.
- [21] Berger, M. & Wong, W.-K. (eds) (2005). *Applied Optimal Designs*, John Wiley & Sons, New York.

ANTHONY C. ATKINSON

Orthogonal Arrays

Introduction and Motivating Example

We wish to conduct a simple experiment to determine the effects of two factors on a quantitative response Y . An experiment to investigate the effects of several factors on a response is referred to as a *factorial experiment* (see **Factorial Experiments**). Assume that each factor is to be observed at only two levels (two values), which we denote -1 and $+1$. If a factor is quantitative, but we decide to observe it at only two of its possible values, -1 is used to denote the smaller of the two levels.

Suppose that we decide to take a total of four observations (runs). At what combinations of values (called the *treatments*) should we observe the factors? We will represent the treatments that we observe for each run by a design matrix or array. Columns represent factors and rows runs. The entry in a given column and row indicates the level of the factor represented by the column that is observed for the run represented by the row. Table 1 gives three of several possibilities. Notice that, in all three designs, the first run consists of the treatment with factor 1 at level 1 and factor 2 at level 1.

A minimal requirement of any design is that it should allow us to estimate all the factor effects (main effects and interactions) that we believe are present. In this experiment, suppose that only main effects are present. Then a design should allow us to estimate the main effects of both factors. To investigate this, for each design we generated observations Y with value equal to the sum of the levels of the two factors plus random error (taken to be normal with mean 0 and standard deviation 0.1). We analyzed the data with statistical software using a main-effects-only analysis

of variance model (see **Analysis of Variance**) and the type-III sums of squares (SS). For more about this type of analysis, see [1, 2] or [3]. The results are summarized in Table 2.

For the first design, the software does not provide estimates of the sum of squares due to the main effects of factors 1 and 2. These sums of squares are estimates of the variation in the response due to these main effects. On reflection, it is not surprising that we do not obtain estimates. The effect of a factor is observed by changing its level and seeing how the response changes. For design 1, whenever we change factor 1 we also change factor 2. Hence, we cannot determine if observed changes in the response are due to factor 1, factor 2, or the sum of their effects. In this case, we say factors 1 and 2 are *completely aliased* (see **Aliasing in Fractional Designs**). In general, this occurs when the columns in a design corresponding to two different factors are perfectly correlated (a **correlation** of either 1 or -1). For design 1, the correlation between the four pairs of values is 1.

In the second and third designs, we obtain estimates of the main effects of factors 1 and 2. However, additional analysis reveals a problem with the second design. Suppose we had assumed that there were no main effects of factor 2 and fit a main-effects model that included only factor 1. The results (for both designs, 2 and 3) are summarized in Table 3.

As should be the case, for both designs the sum of squares for factor 2 is incorporated (added) into the error sum of squares and the total sum of squares is unchanged. However, for design 2 the effect of factor 1 has changed (increased). For design 3 the effect of factor 1 is unchanged.

The problem is that for design 2 the levels of factors 1 and 2 are correlated (the correlation is 0.5774). This is known as *collinearity* in regression and as *partial aliasing* in factorial experiments. When

Table 1 The design matrices for three designs for studying two factors with four runs

Run	Design matrix 1 factors		Design matrix 2 factors		Design matrix 3 factors	
	1	2	1	2	1	2
1	1	1	1	1	1	1
2	1	1	1	1	1	-1
3	-1	-1	-1	1	-1	1
4	-1	-1	-1	-1	-1	-1

Table 2 Degrees of freedom and type-III sums of squares for simulated data for the three designs

Source	Design matrix 1		Design matrix 2		Design matrix 3	
	df	SS	df	SS	df	SS
Factor 1	0	—	1	2.1236	1	4.8323
Factor 2	0	—	1	2.2881	1	3.8384
Error	2	0.0148	1	0.0015	1	0.000027
Total	3	16.3689	3	10.4370	3	8.6706

2 Orthogonal Arrays

Table 3 Degrees of freedom and type-III sums of squares for simulated data for designs 2 and 3 when only factor 1 is included

Source	Design matrix 2		Design matrix 3	
	<i>df</i>	<i>SS</i>	<i>df</i>	<i>SS</i>
Factor 1	1	8.1473	1	4.8323
Error	2	2.2896	2	3.8384
Total	3	10.4370	3	8.6706

partial aliasing is present the effects of the two factors cannot be completely separated. It is manifested when the sum of squares for one effect (factor 1 in the example) changes when a second effect (factor 2 in the example) is removed from the model.

In design 3, the levels of factors 1 and 2 are uncorrelated (the correlation is 0). In this case, factor effects are independent of each other and the sum of squares for the **main effect** of factor 1 is the same whether or not the main effect of factor 2 is included in the model. When the levels of two factors are uncorrelated, the factors have a property known as *orthogonality*. Orthogonality is formally defined in the next section. Orthogonality is also equivalent (provided the levels are coded properly) to perpendicularity of vectors. When viewed as 4×1 column vectors, the two columns of design 3 are perpendicular.

This simple experiment illustrates the fact that to avoid aliasing (complete or partial), and thus obtain independent estimates of the effects of factors, it is desirable that every pair of factors in a design be orthogonal. This is formalized in the notion of an orthogonal array.

Definitions and Examples

In a factorial experiment, if all the level combinations of two factors appear in the same number of runs, the two factors are said to be *orthogonal*. A design is called *orthogonal* if all pairs of factors in the design are orthogonal. More generally, in a matrix two columns are said to be orthogonal if all possible combinations of symbols in the two columns appear equally often together. If all pairs of columns of a matrix are orthogonal, the matrix is said to be orthogonal. For the case where all factors have s

levels, we now define a rich class of designs that have this feature and even stronger properties.

Definition 1 Let S be a set of s symbols or “levels” denoted $0, 1, \dots, s-1$. An $N \times k$ matrix A with entries from S is called an *orthogonal array* with s levels, strength t ($0 \leq t \leq k$), and index λ if every $N \times t$ submatrix of A contains each t -tuple of symbols from S exactly λ times as a row. We use the notation $OA(N, k, s, t)$ to denote such an array.

Because there are s^t possible t -tuples of elements from S , an immediate consequence of this definition is that

$$N = \lambda s^t \quad (1)$$

In particular, equation (1) shows that λ is a function of N , s , and t so the parameters N , k , s , and t are sufficient to define an orthogonal array.

Orthogonal arrays as defined above are sometimes referred to as *symmetrical orthogonal arrays* (see for example [4]). Orthogonal arrays were first introduced by Rao [5, 6]. In the context of two-level factorial experiments, it is not uncommon to denote the elements of S -1 , $+1$ rather than 0 , 1 .

Example 1 An $OA(8, 4, 2, 3)$ is given in Table 4. Rows represent experimental runs and columns factors. Notice that every 8×3 submatrix (every selection of three columns of the matrix) has the property that all eight, three-tuples of elements of $S = \{0, 1\}$ occur exactly $\lambda = 1$ times as a row.

Example 2 An $OA(12, 11, 2, 2)$ with elements of S denoted by -1 and $+1$ is given in Table 5. Notice that every 12×2 submatrix has the property that all four two-tuples of elements of $S = \{-1, 1\}$ occur exactly $\lambda = 3$ times as a row.

Table 4 An $OA(8, 4, 2, 3)$

0	0	0	0
0	0	1	1
0	1	0	1
0	1	1	0
1	0	0	1
1	0	1	0
1	1	0	0
1	1	1	1

Table 5 An OA(12, 11, 2, 2)

+1	+1	-1	+1	+1	+1	-1	-1	-1	+1	-1
-1	+1	+1	-1	+1	+1	+1	-1	-1	-1	+1
+1	-1	+1	+1	-1	+1	+1	+1	-1	-1	-1
-1	+1	-1	+1	+1	-1	+1	+1	+1	-1	-1
-1	-1	+1	-1	+1	+1	-1	+1	+1	+1	-1
-1	-1	-1	+1	-1	+1	+1	-1	+1	+1	+1
+1	-1	-1	-1	+1	-1	+1	+1	-1	+1	+1
+1	+1	-1	-1	-1	+1	-1	+1	+1	-1	+1
+1	+1	+1	-1	-1	-1	+1	-1	-1	+1	-1
-1	+1	+1	+1	-1	-1	-1	+1	-1	+1	+1
+1	-1	+1	+1	+1	-1	-1	-1	+1	-1	+1
-1	-1	-1	-1	-1	-1	-1	-1	-1	-1	-1

In some experiments it may not be possible or desirable to use the same number of levels for each factor. An extension of orthogonal arrays to the case where factors may have different numbers of levels is the following.

Definition 2 A *mixed orthogonal array*, denoted $OA(N, s_1^{k_1} s_2^{k_2}, \dots, s_v^{k_v}, t)$ of strength t is an $N \times k$ matrix where $k = k_1 + k_2 + \dots + k_v$ is the total number of factors, in which the first k_1 columns have symbols from $\{0, 1, \dots, s_1 - 1\}$, the next k_2 columns have symbols from $\{0, 1, \dots, s_2 - 1\}$, and so on, with the property that in any $N \times t$ submatrix every possible t -tuple occurs an equal number of times as a row.

Example 3 An OA(12, 2^4 , 3, 2) is given in Table 6.

Mixed orthogonal arrays are sometimes referred to as *mixed-level orthogonal arrays* or *asymmetrical orthogonal arrays*. Also, an $OA(N, s_1^{k_1}, s_2^{k_2}, \dots, s_v^{k_v}, t)$ is sometimes denoted by $L_N(s_1^{k_1} s_2^{k_2}, \dots, s_v^{k_v})$, when

Table 6 An OA(12, 2^4 , 3, 2)

0	0	0	0	0
0	1	0	1	0
1	0	1	1	0
1	1	1	0	0
0	0	1	1	1
0	1	1	0	1
1	0	0	1	1
1	1	0	0	1
0	0	1	0	2
0	1	0	1	2
1	0	0	0	2
1	1	1	1	2

the strength is assumed to be two. Mixed orthogonal arrays were studied as early as 1961 by Addelman and Kempthorne [7] and were called *asymmetrical orthogonal arrays* by Rao [8].

Construction and Tables of Orthogonal Arrays

A variety of techniques are available for constructing orthogonal arrays. For example, an s^{k-p} **fractional factorial** design of resolution $k - p + 1$ is an $OA(s^{k-p}, k, s, k - p)$. The N -run Plackett–Burman designs (see **Plackett–Burman Designs**), which are designs for experiments with all factors at two levels, can be generated cyclically given the first row. The second row is obtained from the first by shifting all entries to the right by one position and then placing the last entry in the first position. The third row is generated from the second row in the same way and this process continues until $N - 1$ rows are generated. All entries in the last row are -1 . You can verify that the 12-run Plackett–Burman design given previously is generated in this way. Table 7.5 of [4] lists the first row for 12-, 20-, 24-, 36-, and 44-run Plackett–Burman designs.

Most techniques involve the use of certain combinatorial objects (linear codes, mutually orthogonal Latin squares, Hadamard matrices) or abstract algebra (Galois fields, projective geometry). Consult [9] for details, as well as a wealth of information about orthogonal arrays. Extensive tables of orthogonal arrays exist (see for example [9, 10]). N. J. A. Sloane maintains an on-line library of orthogonal arrays at the web site <http://www.research.att.com/~njas/oadir/index.html>. One can also find Plackett–Burman designs at this web site.

Applications

Fractional Factorial Experiments

Some General Comments. Orthogonal arrays with two levels provide useful designs for two-level fractional factorial experiments with k factors (see **Fractional Factorial Designs**). In particular, an $OA(N, k, 2, t)$ yields a design matrix for an N -run experiment in which no factorial effect of order $w \leq t$ is aliased with any other factorial effect of

order $t + 1 - w$ or less. A design with this property is said to have *resolution* $t + 1$ (see **Factorial Designs, Resolution of**). While it is true that a standard 2^{k-p} fractional factorial design of resolution $k - p + 1$ is an $\text{OA}(2^{k-p}, k, 2, k - p)$, it is not true that every $\text{OA}(2^{k-p}, k, 2, k - p)$ necessarily be a standard 2^{k-p} fractional factorial design. In other words, an arbitrary $\text{OA}(2^{k-p}, k, 2, k - p)$ may exhibit complex aliasing, i.e., some higher-order interactions may be partially aliased (rather than fully aliased) with lower-order effects.

Data from an $\text{OA}(N, k, 2, t)$ can be analyzed as follows. Use a model with k factors that includes only main effects and all t and fewer factor interactions of the k factors. Fit the model using standard statistical software. When the number of runs is equal to the number of effects in the model use half-normal plots (see **Half-Normal Plot**) or Lenth's method (see **Lenth's Method for the Analysis of Unreplicated Experiments**). Consult [1, 2] or [3] for more on the analysis of data from a fractional factorial experiment.

More generally, symmetric orthogonal arrays with more than two levels provide useful designs for higher-level factorial experiments. Also, asymmetric orthogonal arrays are useful for constructing designs in mixed-level factorial experiments. See [4] for further discussion.

Main-Effects Plans. When two-factor and higher-order interactions are assumed to be negligible, a model that includes only main effects would be considered adequate for describing the relationship between the response and the factors (see **Main Effect Designs**). Designs in which all pairs of factors are orthogonal have desirable properties in that no main effects are aliased with each other. In 1946, Plackett and Burman [11] constructed a large collection of orthogonal designs that are $\text{OA}(4m, 4m - 1, 2, 2)$, where $4m$ is not a power of 2. The design in Table 5 is an example. Their designs are known as *Plackett–Burman designs* (see **Plackett–Burman Designs**) in the literature and it has become common practice to refer to all $\text{OA}(4m, 4m - 1, 2, 2)$ as Plackett–Burman designs. Selecting $F \leq 4m - 1$ columns of an $\text{OA}(4m, 4m - 1, 2, 2)$ yields a design with $4m$ runs that allows one to estimate the main effects of F factors without aliasing.

Screening Designs. Screening experiments (see **Screening Designs**) are used when one is interested

in investigating a large number of factors simultaneously, but only a few are thought to be *active* (have nonnegligible effects). In screening experiments it is usually only possible to run small fractions of the full factorial. As a consequence, such designs are likely to have low resolution. However, if one discovers that only a few factors are active it is desirable that the subdesign, called a *projection*, consisting only of the active factors be a full factorial. We typically do not know in advance which factors are likely to be active, so designs with the property that the projections onto all sufficiently small subsets of factors are full factorial designs are attractive. When the projection onto the active factors is a full factorial, a reanalysis of the data using only the active factors can be run as a full factorial and all main effects and interactions estimated.

Cheng [12] shows that certain orthogonal arrays on two symbols have this desirable projection property, making these orthogonal arrays useful for screening experiments. A design is said to be of *projectivity* p if, for every subset of p factors, a complete factorial design (possibly with some combinations replicated) is produced. If the two symbols in the orthogonal array are -1 and $+1$, the orthogonal array is said to have a *defining word of length* r if there exist r columns of the array with the property that the column whose i th row is the product of all the entries in the i th row of the r columns is either the vector in which all entries are 1 or the vector in which all entries are -1 . Cheng [12] shows that an $\text{OA}(N, k, 2, t)$ with $k \geq t + 1$ has projectivity $t + 1$ if and only if it has no defining word of length $t + 1$. For example, one can verify that the 12-run Plackett–Burman design (the $\text{OA}(12, 11, 2, 2)$ given in Table 5) has no defining word of length three and hence has projectivity three.

As an extension, one might also be interested in designs with the property that all projections onto a small number of factors be of resolution III (but not necessarily a complete factorial design). Such designs are said to have a hidden projection property. See [4] for more on this.

Optimal Designs for Response Surface Modeling

Suppose N units are observed in an experiment. On the i th, unit a response y_i and p explanatory variables (also called *predictors*, *regressors*, or *covariates*) $x_{i1}, x_{i2}, \dots, x_{ip}$ are measured, $1 \leq i \leq N$. Suppose

also that the response and explanatory variables are related by the following equation or *response surface model* (see **Response Surface Methodology**);

$$y_i = \beta_0 + \beta_1 x_{i1} + \cdots + \beta_p x_{ip} + \varepsilon_i \quad (2)$$

where the ε_i represent independent random errors that are assumed to be normally distributed with mean 0 and variance σ^2 . $\beta_0, \beta_1, \dots, \beta_p$ are unknown constants or parameters to be estimated. The model in equation (2) is sometimes referred to as a *first-order response surface model*. After appropriately scaling each x_{ij} , it may be plausible to assume all x_{ij} satisfy $-1 \leq x_{ij} \leq +1$. It is customary to arrange the values of the x_{ij} in an $N \times (p+1)$ model matrix

$$\mathbf{X} = \begin{pmatrix} 1 & x_{11} & \cdots & x_{1p} \\ \vdots & \vdots & \ddots & \vdots \\ 1 & x_{N1} & \cdots & x_{Np} \end{pmatrix} \quad (3)$$

If the last p columns of $\mathbf{X}$ in equation (3) are an $\text{OA}(N, p, 2, t)$, using -1 and $+1$ to denote the levels, and if $t \geq 2$, then the resulting design is optimal (see **Optimal Design**) under a wide range of standard criteria such as A-, D-, and E-optimality. For example, the $\text{OA}(12, 11, 2, 2)$ given previously is a D-optimal 12-run design for a first-order response surface model with $p = 11$.

When equation (2) includes not only all linear terms (the $x_{i1}, x_{i2}, \dots, x_{ip}$) but also all cross-products of subsets of size q or less of these terms ($q \leq p/2$), then a design for which

$$\begin{pmatrix} x_{11} & \cdots & x_{1p} \\ \vdots & \ddots & \vdots \\ x_{N1} & \cdots & x_{Np} \end{pmatrix} \quad (4)$$

is an $\text{OA}(N, p, 2, t)$, using -1 and $+1$ to denote the levels, and for which $t \geq 2q$, is again optimal under a wide range of standard criteria such as A-, D-, and E-optimality. For $p/2 < q \leq p$, an $\text{OA}(N, p, 2, p)$ is needed in equation (4) for optimality. For example, an $\text{OA}(N, p, 2, t)$ with $t \geq 4$ is a **D-optimal design** for the model

$$y_i = \beta_0 + \sum_{j=1}^p \beta_j x_{ij} + \sum_{1 \leq j < \ell \leq p} \beta_{j\ell} x_{ij} x_{i\ell} + \varepsilon_i \quad (5)$$

with main effects and two-factor interactions.

Robust Parameter Design

In robust parameter design (see **Robust Design**), factors are divided into *control factors* (factors that are fixed once they are selected by a product designer) and *noise factors* (factors that are hard to control or may vary in normal use of a product). A typical objective is to find values of the control factors that make the product insensitive to variation in the noise factors. This can be accomplished by conducting experiments in which the noise factors are varied systematically in order to represent their variation in normal use. One strategy for designing such experiments is to choose separate designs for the control and noise factors. The design for the control factors is called the *control array* and the design for the noise factors the *noise array* (see **Product Array Designs**). Then, for each setting of the control factors in the control array, all settings in the noise array are observed. Such a design is called a *cross array*. The total number of runs in the cross array is the product of the number of factor settings in the control array and the number of factor settings in the noise array. Orthogonal arrays (symmetric or asymmetric) are often used for both the control and noise arrays. A name closely associated with robust parameter design is Genichi Taguchi. For more on robust parameter design, see [4].

Computer Experiments

There are situations in which a physical experiment is not feasible. In such cases it may be possible to describe the physical process of interest by a mathematical model implemented with code on a computer. One can then experiment with the code by varying the inputs. This is called a *computer experiment* (see **Computer Experiments**). For complex code that runs slowly (for example, taking 1 day for a single run), the number of runs is limited and one must design the computer experiment carefully. One principle is to select runs that are, in some sense, spread evenly over the design space. Orthogonal arrays have been used as the basis for designs in **computer experiments**. For k predictor variables each constrained to be between 0 and 1, and for some positive integer N , let $S = \{1/(N+1), 2/(N+1), \dots, N/(N+1)\}$. For this choice of symbols, an $\text{OA}(N^t, k, N, t)$ is an N^t -run design ($1 \leq t \leq k$) with the property that for any t columns,

if the rows are taken as points, we have a set of evenly spaced points on the t -dimensional unit cube $[0, 1]^t$. Thus, an $OA(N^t, k, N, t)$ provides a design with the property that when projected onto any t , or smaller, dimensional subspace (any subset of t or fewer columns), the points are spread evenly. When $t = 1$, this produces what is known as a *Latin hypercube design* (see **Latin Hypercube Designs**). Koehler and Owen [13] provide additional examples of the use of orthogonal arrays in computer experiments.

An Example

Hunter *et al.* [14] ran a screening experiment to study the effects of seven factors; initial structure (A), bead size (B), pressure treatment (C), heat treatment (D), cooling rate (E), polish (F), and final treatment (G), on the fatigue life of weld-repaired castings. Possible designs included an 8-run, resolution-III fractional factorial and a 16-run, resolution-IV fractional factorial. The eight-run design allows one to fit only a model with main effects but no interactions, and no separate estimate for error is possible. The 16-run design allows one to add some two-factor interactions (assuming others are negligible) and allows an estimate of the error. To reduce the number of runs, avoid aliasing main effects with each other, and still accommodate some two-factor interactions (assuming others are negligible – see [4, Chapter 8] for discussion of the actual aliasing present in the design), a 12-run $OA(12, 7, 2, 2)$ can be used. This was the design actually employed. The design and response (log lifetime) are given in Table 7.

The objective was to identify the factors that affect casting lifetime. A main-effects-only model can be fit to the data. Using a normal plot or Lenth's method, only factor F is identified as significant. A traditional analysis of variance gives a P value of 0.0556 and R^2 of 0.75. A follow-up analysis can be conducted by adding various subsets of all two-factor interactions with F to a model including all main effects. This process further identifies the FG interaction as significant (the P value for F is now 0.0144 and 0.042 for FG), with a marked improvement in the model fit (R^2 increases to 0.95). The conclusion is that polish and its interaction with the final treatment affect lifetime. Examining the factor effects indicates that setting F at its highest level and G at its lowest level produces the longest

Table 7 The $OA(12, 7, 2, 2)$ for the 12-run, seven-factor experiment, with the log of the observed lifetimes

A	B	C	D	E	F	G	log lifetime
+1	+1	-1	+1	+1	+1	-1	6.058
+1	-1	+1	+1	+1	-1	-1	4.733
-1	+1	+1	+1	-1	-1	-1	4.625
+1	+1	+1	-1	-1	-1	+1	5.899
+1	+1	-1	-1	-1	+1	-1	7.000
+1	-1	-1	-1	+1	-1	+1	5.752
-1	-1	-1	+1	-1	+1	+1	5.683
-1	-1	+1	-1	+1	+1	-1	6.607
-1	+1	-1	+1	+1	-1	+1	5.818
+1	-1	+1	+1	-1	+1	+1	5.917
-1	+1	+1	-1	+1	+1	+1	5.863
-1	-1	-1	-1	-1	-1	-1	4.809

lifetime. This is consistent with what one observes in the data in Table 7.

References

- [1] Box, G.E.P., Hunter, W.G. & Hunter, J.S. (2005). *Statistics for Experimenters. An Introduction to Design, Data Analysis, and Model Building*, 2nd Edition, John Wiley & Sons, New York, p. 633.
- [2] Dean, A. & Voss, D. (1999). *Design and Analysis of Experiments*, Springer-Verlag, New York, p. 740.
- [3] Montgomery, D.C. (2004). *Design and Analysis of Experiments*, 5th Edition, John Wiley & Sons, New York, p. 660.
- [4] Wu, C.F.J. & Hamada, M. (2000). *Experiments*, John Wiley & Sons, New York, p. 630.
- [5] Rao, C.R. (1946). Hypercubes of strength "d" leading to confounded designs in factorial experiments, *Bulletin of the Calcutta Mathematical Society* **38**, 67–78.
- [6] Rao, C.R. (1947). Factorial experiments derivable from combinatorial arrangements of arrays, *Journal of the Royal Statistical Society, Supplement* **9**, 128–139.
- [7] Addelman, S. & Kempthorne, O. (1961). Orthogonal main-effect plans, Technical Report ARL 79, Aeronautical Research Lab, Wright-Patterson Air Force Base, November 1961.
- [8] Rao, C.R. (1973). Some combinatorial problems of arrays and applications to the design of experiments, in *A Survey of Combinatorial Theory*, J.N. Srivastava, ed, North-Holland, Amsterdam, pp. 349–359.
- [9] Hedayat, A.S., Sloane, N.J.A. & Stufken, J. (1999). *Orthogonal Arrays*, Springer, New York, p. 416.
- [10] Raghavarao, D. (1971). *Construction and Combinatorial Problems in Design of Experiments*, John Wiley & Sons, New York, p. 386.

-
- [11] Plackett, R.L. & Burman, J.P. (1946). The design of optimum multifactorial experiments, *Biometrika* **33**, 305–325.
- [12] Cheng, C.-S. (2006). Projection properties of factorial designs for screening experiments, in *Screening. Methods for Experimentation in Industry, Drug Discovery, and Genetics*, A. Dean & S. Lewis, eds, Springer Science+Business Media, New York, pp. 156–168.
- [13] Koehler, J.R. & Owen, A.B. (1996). Computer experiments, in *Handbook of Statistics*, S. Ghosh & C. Rao, eds, Elsevier Science B. V., Amsterdam, Vol. 13, pp. 261–308.
- [14] Hunter, G.B., Hodi, F.S. & Eager, T.W. (1982). High-cycle fatigue of weld repaired cast Ti-6Al-4V, *Metallurgical Transactions* **13A**, 1589–1594.

Related Articles

Fractional Factorial Designs; Latin Hypercube Designs; Plackett–Burman Designs; Product Array Designs; Screening Designs.

WILLIAM I. NOTZ

Plackett–Burman Designs

Construction of Plackett–Burman Designs

Plackett and Burman (PB) [1] provided us with a class of two-level orthogonal designs which now just are referred to as *PB designs*. If we restrict the name to only those designs introduced by Plackett and Burman, they exist for the number of runs $n = 4m$, $m = 1, 2, \dots, 25$ (except 23 which was given by Baumert *et al.* [2]).

Plackett and Burman used three ways of constructing these designs. One is cyclic generation. For a cyclic generated orthogonal array (see **Orthogonal Arrays**), it is possible to write down all the $n - 1$ contrast columns knowing only the sequence of ± 1 's in the first row. For instance for the 12-run PB design this string omitting the 1's and only writing down the signs is given by: $+ - + - - - + + - +$. Shifting this row cyclically one place 10 times and adding a row of minus signs produces a design matrix for the 12-run PB design. Adding an initial column of plus signs gives us the 12×12 Hadamard matrix (see **Projectivity in Experimental Designs**), shown in Table 1 where the 11 factors are denoted $x_A, x_B, \dots, x_K$.

Plackett and Burman [1] used this way of constructing orthogonal arrays for $n = 8, 12, 16, 20, 24, 32, 36, 44, 48, 60, 68, 72, 80$, and 84. For $n = 28, 52, 76$, and 100 they used block cycling and for $n = 40, 56, 64, 88$, and 96 their way of obtaining orthogonal arrays was by doubling. The technique of doubling can be explained as follows. Let H_n be a $n \times n$ Hadamard matrix with all entries in the first column equal to $+1$. The matrix $\begin{bmatrix} H_n & H_n \\ H_n & -H_n \end{bmatrix}$ is then a $2n \times 2n$ Hadamard matrix whose last $2n - 1$ columns constitute the orthogonal two-level array.

For $n = 2^k$ $n = 1, 2, 3, 4, 5$, and 6, PB designs coincide with the well known fractional factorial two-level designs (see **Fractional Factorial Designs**), which also are called *geometric designs* or sometimes *regular designs*. The rest of the PB designs are called *nongeometric designs* or sometimes *nonregular designs*. The signs for cyclic generation of the 20-, 24-, and 36-run PB designs are:

$$n = 20 : + + - - + + + - + - + - - \\ - - + + -$$

$$n = 24 : + + + + + - + - + + - - + + \\ - - + - + - - - - \\ n = 36 : - + - + + + - - - + + + + + \\ - + + + - - + - - - - + - + - \\ + + - - + -$$

How to obtain the other PB designs is described in great detail in [1]. PB designs are also available on Neil Sloane's web site: www.research.att.com/~njas/index.html#TABLES.

The rest of this article focuses on **projectivity** of PB design, methods for analyzing them, an example to show their efficiency for experimentation and some techniques that can be used to construct designs with efficient run sizes and good projective properties from PB designs.

Projectivity of Plackett–Burman Designs

PB designs were introduced as main-effects plans (see **Main Effect Designs**), and nongeometric PB designs have for a long time been considered to be important **screening designs** in situations where only main effects were expected to show up. Later (Lin and Draper [3], Box and Bisgaard [4], Cheng [5], Box and Tyssedal [6], and Samset and Tyssedal [7]) it was discovered that the nongeometric PB designs have quite remarkable projective properties compared to the geometric fractional **factorial** designs. The concept of design projectivity for two-level designs was defined by Box and Tyssedal [6]. A $n \times k$ design with n runs and k factors each at two levels is said to be of projectivity P if the design contains a complete 2^P factorial in every possible subset of P out of the k factors, possibly with some points replicated. Box and Tyssedal [6] describe such designs as (n, k, P) screens.

Box and Tyssedal [6] were able to prove three propositions for orthogonal two-level arrays that can be used for classifying PB designs with respect to being of projectivity $P = 2$ and only two or of projectivity at least $P = 3$ (see **Projectivity in Experimental Designs**). Later Samset and Tyssedal [7] discovered by a computer search that PB designs with $n = 68, 72, 80$, and 84 also are $(n, n - 1, 4)$ screens. An orthogonal design with $k = n - 1$ is called *saturated*. A complete listing of the projectivity of saturated PB designs is given in Table 2.

2 Plackett–Burman Designs

Table 1 A 12×12 Hadamard matrix where the last 11 columns constitute a 12-run PB design

	x_A	x_B	x_C	x_D	x_E	x_F	x_G	x_H	x_I	x_J	x_K
1	1	-1	1	-1	-1	-1	1	1	1	-1	1
1	1	1	-1	1	-1	-1	-1	1	1	1	-1
1	-1	1	1	-1	1	-1	-1	-1	1	1	1
1	1	-1	1	1	-1	1	-1	-1	-1	1	1
1	1	1	-1	1	1	-1	1	-1	-1	-1	1
1	1	1	1	-1	1	1	-1	1	-1	-1	-1
1	-1	1	1	1	-1	1	1	-1	1	-1	-1
1	-1	-1	1	1	1	-1	1	1	-1	1	-1
1	-1	-1	-1	1	1	1	-1	1	1	-1	1
1	1	-1	-1	-1	1	1	1	-1	1	1	-1
1	-1	1	-1	-1	-1	1	1	1	-1	1	1
1	-1	-1	-1	-1	-1	-1	-1	-1	-1	-1	-1

We observe that all saturated geometric fractional factorial designs are of projectivity $P = 2$ and only 2. However for all of them it is possible to remove $n/2 - 1$ columns and obtain a $(n, n/2, 3)$ screen (Box and Tyssedal [6]).

In their computer search Samset and Tyssedal [7] also discovered that all PB designs with projectivity $P = 3$ have the hidden projection properties (Lin and Draper [3] and Wang and Wu [8]) that all main-effects and two-factor interactions for any four factors are estimable under the assumption that higher-order interactions are negligible.

Methods for Analyzing Plackett–Burman Designs

Nongeometric designs have more complex aliasing (see **Aliasing in Fractional Designs**), than the geometric designs (Lin and Draper [9]). For instance in a 12-run PB design every **main effect** may be partially aliased with 45 two-factor interactions and a single two-factor interaction will appear in the alias pattern of all main effects not involved with this two-factor interaction. As a result standard

methods often used for analyzing unreplicated two-level designs such as normal and half-normal plots (Daniel [10]), or more quantitative methods such as **Lenth's method** (Lenth [11]) may be of little value since these methods are based on being able to separate active contrasts from contrasts estimating only noise.

However, several methods are available also for analyzing these designs. These can mainly be classified as effect based or factor based. Effect-based methods aim at identifying significant effects. A model consisting of main effects and interactions (see **Interactions**), is assumed to give an adequate approximation of the response. The principle of effect heredity; i.e., an interaction is excluded from the model unless at least one (weak heredity) or both (strong heredity) of the main effects associated with the interaction also are included, is often a precept. Examples of such methods can be found in Hamada and Wu [12] and Chipman *et al.* [13]. An effect-based method that doesn't depend on the heredity principle is given in Tyssedal and Kulahci [14]. Factor-based methods aim at identifying active factors and they are less dependent on model assumptions, heredity included (see **Projectivity in Experimental Designs**). Their drawback may be that if the factor effects are small relative to the level of noise, it may be difficult to discriminate among several candidate sets of active factors. Examples of such methods are given in Box and Meyer [15], Kulahci and Box [16] and Tyssedal and Samset [17]. For both methods it is crucial that sparsity of either effects or factors respectively can be assumed. In addition projective properties of the

Table 2 The projectivity of Plackett–Burman designs

Screens	Number of runs, n
$P = 2$	4, 8, 16, 32, 40, 56, 64, 88, 96
$P = 3$	12, 20, 24, 28, 36, 44, 48, 52, 60, 76, 100
$P = 4$	68, 72, 80, 84

design used are conclusive in order for a factor-based search to work well (see **Projectivity in Experimental Designs**).

An Example

Montgomery *et al.* [18] published a paper in quality engineering with the title: *Some Cautions in the Use of Plackett–Burman designs*. They used a 12-run PB designs, see Table 1, and data were simulated from the following model

$$y = 200 + 8x_A + 10x_D + 12x_E - 12x_Ax_D + 9x_Ax_E + \varepsilon \quad (1)$$

where ε is a normally distributed random error term with mean 0 and variance 9. The columns for the factors x_A , x_D , and x_E correspond to the ones in Table 1. We note that the model obeys both the weak and the strong heredity principle. Montgomery *et al.* [18] used normal **probability plots** in combination with the method suggested by Hamada and Wu [11] for the analysis. Due to large two-factor interactions these two methods are not very suitable for analyzing this response and their conclusion was that a sequential strategy, starting with a 2_{III}^{11-7} fractional factorial and adding two additional follow-up runs in this situation outperformed a PB design.

Now suppose we instead chose to use a factor-based method as proposed in Tyssedal and Samset [17] which we will call a projective-based search. This method exploits the projective properties of the design (Lin and Draper [3] and Box and Tyssedal [6]). In particular, the 12-run PB design projects into a 2^1 design replicated six times in any one factor. It projects into a 2^2 design replicated three times in any two factors and it projects into a 2^3 design with four runs replicated in any three factors.

Thereby assuming at most three active factors, we can investigate how well every set of k factors explain the data $1 \leq k \leq 3$, by examining the fit for replicated runs, calculating

$$\hat{\sigma}^2 = \frac{1}{\sum_{i=1}^s (r_i - 1)} \sum_{i:\text{replicated}} \sum_{j=1}^{r_i} (y_{ij} - \bar{y}_i)^2 \quad (2)$$

where s is the number of runs replicated and r_i is the number of replications for replicated run number i .

Table 3 Results from a projective-based search on the response given in equation (1)

1 active		2 active		3 active	
$\hat{\sigma}$	Factor	$\hat{\sigma}$	Factors	$\hat{\sigma}$	Factors
18.15	x_E	16.78	x_A, x_E	4.26	x_A, x_D, x_E
24.62	x_F	16.80	x_C, x_E	8.05	x_B, x_E, x_I
24.75	x_A	17.40	x_E, x_I	8.29	x_C, x_E, x_J
24.87	x_D	17.66	x_B, x_J	8.34	x_C, x_E, x_H
24.87	x_G	17.98	x_E, x_F	9.25	x_D, x_E, x_F

Sets of factors that give small values of $\hat{\sigma}^2$ are candidates for being declared active.

We simulated a new set of data from the model given in equation (1). A candidate list showing the five set of factors minimizing $\hat{\sigma}$ for $k = 1, 2$, and 3 is given in Table 3. Clearly assuming at most three active factors one would in this case end up with x_A , x_D , and x_E after 12 runs and no heredity assumptions need to be imposed on the model.

In cases where several set of factors have approximately the same $\hat{\sigma}$, backward regression on a model fully expanded with main effects and interactions will sometimes eliminate the ambiguity. For more on this method we refer to Tyssedal and Samset [17] and Tyssedal *et al.* [19] where also a successful application of the 12-run PB design is given. Finally we remark that in this case a search through all possible models consisting of main-effects and two-factor interactions for up to three factors would work equally well. The success of such an effect-based search, however, depends on the underlying model (see **Projectivity in Experimental Designs**).

Construction of Designs from PB Designs

The **foldover** of a $n \times k$ two-level design is the design obtained when all mirror images of the original design's runs are added in addition to a column where the first n entries are +1 and the last n entries are -1 (see **Projectivity in Experimental Designs**). If $\mathbf{X}$ is a nongeometric PB design, Samset and Tyssedal [7] found by computer search that the foldover is always of projectivity $P = 4$ or a $(2n, n + 1, 4)$ screen except for $n = 40, 56, 88$, and 96. Designs constructed in this way from PB designs will always have the hidden projection property that all main-effects and two-factor interactions of any five factors can be estimated (Samset and Tyssedal [7]).

A two-level design with more than $n - 1$ factors is called *supersaturated* (see **Supersaturated Designs**). Lin [20] introduced a new class of supersaturated designs using half-fractions of orthogonal arrays, i.e., PB designs. For any arbitrary column the $n/2$ rows corresponding to either $+1$ or -1 in that column will constitute a design allowing $n - 2$ factors in $n/2$ runs. If the PB designs used for construction is of projectivity $P = 4$, the supersaturated designs will have projectivity at least $P = 3$, Samset and Tyssedal [7].

PB designs have also been suggested for use in **split-plot** experimentation (see **Split-Plot Designs**), both as whole-plot and subplot arrays in order to reduce the number of individual tests (Kulahci and Bisgaard [21] and Tyssedal and Kulahci [14]).

Concluding Remarks

Due to their complex alias pattern practitioners have often been warned against using nongeometric PB designs. This warning seems to be exaggerated and may in some situations act as a barrier against efficient experimentation. Nongeometric PB designs provide us with a lot more design alternatives. If factor sparsity can be assumed their excellent projective properties compared to the geometric fractional factorials allow us to identify active factor spaces [22] under weak assumptions on the underlying model, and their hidden projection properties are very attractive if the model is adequately approximated in terms of main-effects and two-factor interactions.

One of the reasons why these designs are not used as often as they should is the lack of good available software to do the analysis. In the future more effort is needed to overcome this unnecessary obstacle.

References

- [1] Plackett, R.L. & Burman, J.P. (1946). The design of optimum multifactorial experiments, *Biometrika* **33**, 305–325.
- [2] Baumert, L.D., Golomb, S.W. & Hall, M. (1962). Discovery of a Hadamard matrix of order 92, *Bulletin of the American Mathematical Society* **68**, 237–238.
- [3] Lin, D.K.J. & Draper, N.R. (1992). Projection properties of Plackett and Burman designs, *Technometrics* **34**, 423–428.
- [4] Box, G.E.P. & Bisgaard, S. (1993). What can you find out from 12 experimental runs, *Quality Engineering* **5**, 663–668.
- [5] Cheng, C.S. (1995). Some projection properties of orthogonal arrays, *The Annals of Statistics* **23**, 1223–1233.
- [6] Box, G.E.P. & Tyssedal, J. (1996). Projective properties of certain orthogonal arrays, *Biometrika* **83**(4), 950–955.
- [7] Tyssedal, J. & Samset, O. (1999). *Two-Level Designs with Good Projection Properties*, Preprint, Statistics No. 12/1999, Norwegian University of Science and Technology, Norway.
- [8] Wang, J.C. & Wu, C.F.J. (1995). A hidden projection property of Plackett–Burman and related designs, *Statistica Sinica* **5**, 235–250.
- [9] Lin, D.K.J. & Draper, N.R. (1993). Generating alias relationships for two-level Plackett and Burman designs, *Computational Statistics and Data Analysis* **15**, 147–157.
- [10] Daniels, C. (1959). Use of half-normal plots in interpreting factorial two-level experiments, *Technometrics* **1**, 311–340.
- [11] Lenth, R.V. (1989). Quick and easy analysis of unreplicated factorials, *Technometrics* **31**(4), 469–473.
- [12] Hamada, H. & Wu, C.F.J. (1992). Analysis of designed experiments with complex aliasing. *Journal of Quality Technology* **24**(3), 130–137.
- [13] Chipman, H., Hamada, H. & Wu, C.F.J. (1997). A Bayesian variable-selection approach for analyzing designed experiments, *Technometrics* **39**, 372–381.
- [14] Tyssedal, J. & Kulahci, M. (2005). Analysis of split-plot designs with mirror image pairs as sub-plots. *Quality and Reliability Engineering* **21**(5), 539–551.
- [15] Box, G.E.P. & Meyer, R.D. (1993). Finding the active factors in fractionated screening experiments, *Journal of Quality Technology* **25**, 94–105.
- [16] Kulahci, M. & Box, G.E.P. (2003). Catalysis of discovery and development in engineering and industry, *Quality Engineering* **15**(3), 509–513.
- [17] Tyssedal, J. & Samset, O. (1997). *Analysis of the 12 Run Plackett–Burman Design*, Preprint, Statistics No. 8/1997, Norwegian University of Science and Technology, Norway.
- [18] Montgomery, D.C., Borror, C.M. & Stanley, J.D. (1997). Some cautions in the use of Plackett–Burman designs. *Quality Engineering* **10**(2), 371–381.
- [19] Tyssedal, J., Grinde, H. & Røstad, C.C. (2006). The use of a 12 run Plackett–Burman design in the injection moulding of a technical plastic component, *Quality and Reliability Engineering* **22**, 651–657.
- [20] Lin, D.K.J. (1993). A new class of supersaturated designs, *Technometrics* **35**, 28–31.
- [21] Kulahci, M. & Bisgaard, S. (2005). The use of Plackett–Burman designs to construct split-plot designs, *Technometrics* **47**(4), 495–501.
- [22] Box, G.E.P. & Tyssedal, J. (2001). Sixteen run designs of high projectivity for screening, *Communications in Statistics: Simulation and Computation* **30**(2), 217–228.

JOHN TYSSEDAL

Projectivity in Experimental Designs

Concept and Usefulness

Projectivity concerns the properties of a **factorial** design when it is restricted to a subset of experimental factors. Box and Tyssedal [1] defined projectivity of two-level designs as follows: A $n \times k$ design with n runs and k factors each at two levels is said to be of projectivity P if the design contains a complete 2^P factorial in every possible subset of P out of the k factors, possibly with some points replicated. Box and Tyssedal [1] describe such designs as (n, k, P) screens. In general a $n \times k$ design with all factors at s levels is said to be of projectivity P if every subset of P factors out of possible k contains a complete s^P factorial design, possibly with some points replicated. This article will present results for two-level designs only.

The concept of projectivity is related to strength of orthogonal arrays (Rao [2], *see* **Orthogonal Arrays**), but differs in that complete copies of a 2^P factorial are not required in order to have projectivity P . Projectivity is also related to the resolution of a design (Box *et al.* [3], *see* **Factorial Designs, Resolution of**). A design is of resolution R if no p -factor effect is confounded with any other effect containing less than $R - p$ factors. For geometric fractional factorial designs (*see* **Fractional Factorial Designs**); that is a $1/2^s$ fraction of a 2^k factorial it is always the case that $P = R - 1$ [1].

The rationale for considering projection properties of **screening designs** was first pointed out by Box *et al.* [3] and further explained in Box and Tyssedal [4]. In factor screening a model that often sufficiently approximates reality is that out of a larger number k of tested factors only a small subset, typically two or three, maybe four, are expected to really affect a particular response. This is known as *factor sparsity* and tells us that out of the whole k dimensional factor space there is an active subspace of lower dimension within which most of the changes in the measured response occur. When choosing an appropriate design for the experimental investigation, it is then of importance to consider projections of the design on to small subsets of factors.

Statistically, projectivity P implies that all main effects and all interactions of any P factors are estimable with no **bias** if the other factors are inert, but its usefulness for factor screening goes beyond that of **estimation** capacity. If our candidate set under investigation contains the true active factors, replicated runs have the same expected value. Thereby a model independent estimate of the error variance may be obtained regardless of any functional relationship, and the ability of a subset of factors to explain the variation in the response may be evaluated with rather weak assumptions on the underlying model. A pretty common assumption for analyzing a screening experiment is to assume that a model consisting of main effects and interactions gives an adequate approximation of the response. Also the principle of effect heredity (*see* **Plackett–Burman Designs**), is often a precept. These assumptions cannot be relied upon in every situation, Box and Tyssedal [4]. The importance of taking projective properties into account is that it may enable us to identify the active subspace without necessarily imposing restrictive assumptions on the underlying model. A more thorough investigation of the real relationship between design factors and responses can then be performed once the active subspace is identified.

The rest of this article will focus on projective properties of two-level designs, an example to illustrate how such properties can be taken advantage of when a factor screening is performed and some techniques that can be used to obtain designs with good projection properties in various experimental situations.

Projective Properties of Two-Level Orthogonal Arrays

Orthogonal two-level designs may be constructed from Hadamard matrices. A Hadamard matrix is a $n \times n$ matrix with entries ± 1 where rows and columns are pairwise orthogonal. Given a Hadamard matrix, H , an equivalent matrix with all elements in the first column equal to $+1$ can be obtained by multiplying by -1 each element in every row of H whose first element is -1 . The remaining $n - 1$ columns must have half 1 's and half -1 's and constitute an orthogonal two-level design that can accommodate $n - 1$ factors in n runs. Orthogonal designs with $k = n - 1$ are called *saturated*.

2 Projectivity in Experimental Designs

When $n = 2^k$ one type of two-level design can be written down as follows. In the first column -1 and $+1$ are alternated. The second column consists of alternating pairs of -1 's and $+1$'s. In general the length of the string of -1 's and $+1$'s doubles for each column up to column k which consists of 2^{k-1} -1 's followed by 2^{k-1} $+1$'s. The rest of the columns are then all possible interaction columns between these k columns. Designs written down in this way are known as the *fractional factorials* (see **Fractional Factorial Designs**), sometimes called *geometric designs* or *regular designs*. The rest of the orthogonal two-level arrays are called *nongeometric designs* or sometimes *nonregular designs*. A necessary condition for these designs to exist is that $n \equiv 0 \pmod{4}$.

Now let H_n be a $n \times n$ Hadamard matrix with all elements in the first column equal to $+1$. Then a $2n \times 2n$ Hadamard matrix may be constructed by the following technique: $H_{2n} = \begin{pmatrix} H_n & H_n \\ H_n & -H_n \end{pmatrix}$, called *doubling*. The last $2n - 1$ columns constitute a $2n \times 2n - 1$ orthogonal array. Another way of obtaining orthogonal arrays is by cyclic generation (see **Plackett–Burman Designs**). Box and Tyssedal [1] were able to prove three propositions that are useful for determining the projectivity of an orthogonal two-level array.

- A saturated design obtained from a doubled $n \times n$ Hadamard matrix is always of projectivity $P = 2$ and only 2.
- A saturated design obtained from a cyclic orthogonal array is either a geometric factorial orthogonal array with $P = 2$ and only 2, or else has projectivity at least $P = 3$.
- Any saturated two-level design obtained from an orthogonal array containing $n = 4m$ runs, with m odd, is of projectivity at least $P = 3$.

The last result is also in agreement with a published result in Cheng [5].

All saturated fractional factorials can be obtained from a doubled $n \times n$ Hadamard matrix and are therefore projectivity $P = 2$ and only two designs. However for all of them it is possible to remove $n - 1$ columns and obtain a $(2n, n, 3)$ screen [1]. If the saturated design obtained from H_n is of projectivity $P = 3$, the design obtained from doubling H_n and then removing the column with $+1$ for the first

n runs and -1 for the last n runs will be a $(2n, 2n - 2, 3)$ screen.

It is interesting to note that for $n = 8$ it is only possible to obtain a $(8, 4, 3)$ screen whereas the orthogonal array for $n = 12$ is cyclic generated and provides us with a $(12, 11, 3)$ screen. Hall [6] showed it was possible to construct five 16-run orthogonal arrays. One provides us with the well-known fractional factorial design from which a $(16, 8, 3)$ screen can be constructed. For the four others $(16, 12, 3)$ screens can be obtained from three of them and a $(16, 14, 3)$ screen can be obtained from the fourth [4].

An important class of two-level designs was derived by Plackett and Burman (PB) [7] (see **Plackett–Burman Designs**). Using the three propositions above, all PB designs can be classified with respect to being of projectivity $P = 2$ and only 2 or of projectivity at least $P = 3$ (Box and Tyssedal [1], see **Plackett–Burman Designs**). Later Samset and Tyssedal [8] discovered by a computer search that PB designs with $n = 68, 72, 80$, and 84 also are of projectivity $P = 4$.

An Example

To show the usefulness of exploiting projective properties in the search for active factors, assume a 12-run PB design (see **Plackett–Burman Designs**) in the 11 factors $x_A, x_B, \dots, x_K$, has been performed and that y is related to the three factors x_A, x_B , and x_C through the functional relationship

$$E(y) = 2x_A + 0.8x_B + \frac{2}{0.4x_B + x_Ax_Bx_C + 2} \quad (1)$$

Noise from a normally distributed error term with mean 0 and standard deviation $\sigma = 0.3$ was added and the projective-based search given in Tyssedal and Samset [9] and Tyssedal *et al.* [10] (see also **Plackett–Burman Designs**) was used to identify the active factors.

The two first columns of Table 1 show the five subspaces of three factors giving the best fit in a projective-based search. The subspace consisting of the three factors x_A, x_B , and x_C is clearly the one that best explains the variation in the data. The two last columns show the result of a search through all possible three-factor models consisting of main effects and two-factor interactions, a rather common

precept for a search based on the heredity principles. We notice that only x_A is included in the five candidate sets having the smallest estimated σ . In fact with zero noise and all three-factor interactions assumed inert, x_B and x_C are not included in the 10 subspaces of three factors giving the best fit.

A Bayesian approach to perform a factor based search is given in Box and Meyer [11].

Some Techniques Used to Construct Designs with Good Projection Properties

The **foldover** of a $n \times k$ two-level design is the design obtained when all mirror images of the original design's runs are added in addition to a column where the first n entries are $+1$ and the last n entries are -1 . Let $\mathbf{1}$ be a $n \times 1$ vector of ones. The foldover, $\tilde{\mathbf{X}}$, of a two-level design $\mathbf{X}$ is given by $\tilde{\mathbf{X}} = \begin{bmatrix} \mathbf{X} & \mathbf{1} \\ -\mathbf{X} & -\mathbf{1} \end{bmatrix}$.

Zeiden and Zemach [12] showed that the foldover of a $n \times k$ orthogonal two-level design of strength 2 is a $2n \times (k+1)$ orthogonal two-level design of strength 3 and therefore according to Cheng [5] also a projectivity $P = 4$ design in $k+1$ factors if n is not a multiple of 8.

Orthogonal two-level arrays have some important hidden projections properties [13, 14]. Assume interactions involving more than two factors can be neglected. The following results are due to Cheng [5] and Cheng [15] respectively.

- If n is not a multiple of 8, then in the projections on to any four factors, all main effects and two-factor interactions are estimable.

Table 1 A comparison between two searches for identifying three active factors in a 12-run PB design. One is a projective-based search. The other is a search where the three-factor interaction is neglected

Projective-based search		Neglecting interactions with three or more factors	
Factors in model	$\hat{\sigma}$	Factors in model	$\hat{\sigma}$
x_A, x_B, x_C	0.23	x_A, x_H, x_I	0.57
x_A, x_F, x_G	0.62	x_A, x_G, x_J	0.73
x_A, x_D, x_J	0.63	x_A, x_E, x_G	0.76
x_A, x_H, x_I	0.64	x_A, x_D, x_F	0.80
x_A, x_G, x_J	0.64	x_A, x_G, x_K	0.85

- If n is not a multiple of 16, then in the projections on to any five factors in a $P = 4$ design, obtained by foldover, all main effects and two-factor interactions are estimable.

Let $\mathbf{X}$ be a $(n, k, 3)$ screen. The design $\begin{bmatrix} \mathbf{X} & \mathbf{X} \\ \mathbf{X} & -\mathbf{X} \end{bmatrix}$ obtained from doubling $\mathbf{X}$ is then a $(2n, 2k, 3)$ screen. Note that $\mathbf{X}$ does not have to be orthogonal. The 68-run PB design is of projectivity $P = 4$. Lin [16] introduced a new class of supersaturated designs (*see Supersaturated Designs*), using half-fractions of Hadamard matrices. Such a construction normally reduces the projectivity by one [8]. A $(34, 66, 3)$ screen can then be constructed using an arbitrary column in the 68-run PB design as a branching column [8]. From this design a $(68, 132, 3)$ screen can be constructed by doubling.

Foldover and doubling have an interesting coupling to split-plot experimentation (*see Split-Plot Designs*). Suppose for each level combination of the whole-plot factors, two runs as mirror image pairs are used at the subplot level. The corresponding design can be written as

$$\begin{bmatrix} \mathbf{W} & \mathbf{S} \\ \mathbf{W} & -\mathbf{S} \end{bmatrix}$$

where $\begin{bmatrix} \mathbf{W} \\ \mathbf{W} \end{bmatrix}$ contains the whole-plot factors and $\begin{bmatrix} \mathbf{S} \\ -\mathbf{S} \end{bmatrix}$ the subplot factors. $\mathbf{W}$ and $\mathbf{S}$ may be different or equal. For instance if $\mathbf{X}$ is a $(n, k, 3)$ screen, $\mathbf{W} = \mathbf{S} = \mathbf{X}$ gives us a split-plot design of projectivity $P = 3$ that can accommodate n whole-plot factors and n subplot factors. For more details on these designs (see Tyssedal and Kulahci [17]).

Concluding Remarks

The reason it is important to focus on projective properties is the need to be able to identify active subspaces of factors without necessarily imposing restrictive assumptions on the underlying model. Nongeometric designs have very favorable projective properties compared to geometric ones, and from that perspective seem to be the preferable choice when many factors need to be investigated

4 Projectivity in Experimental Designs

in few runs and factor sparsity can be assumed. Many of the nongeometric designs have also very favorable hidden projection properties that essentially increase their projectivity by one if the model is adequately approximated by means of main effects and two-factor interactions. Such designs have therefore a great potential for use in screening experimentation and more attention should be paid to their use.

References

- [1] Box, G.E.P. & Tyssedal, J. (1996). Projective properties of certain orthogonal arrays, *Biometrika* **83**(4), 950–955.
- [2] Rao, C.R. (1947). Factorial experiments derivable from combinatorial arrangements of arrays, *Journal of the Royal Statistical Society, Supplement* **9**, 128–139.
- [3] Box, G.E.P., Hunter, J.S. & Hunter, W.G. (1978). *Statistics for Experimenters: An Introduction to Design, Data Analysis, and Model Building*, John Wiley & Sons, New York.
- [4] Box, G.E.P. & Tyssedal, J. (2001). Sixteen run designs of high projectivity for screening, *Communications in Statistics: Simulations and Computation* **30**(2), 217–228.
- [5] Cheng, C.S. (1995). Some projection properties of orthogonal arrays, *The Annals of Statistics* **23**, 1223–1233.
- [6] Hall, M.J. (1961). Hadamard matrices of order 16, *Jet Propulsion Laboratory, Summary* **1**, 21–26.
- [7] Plackett, R.L. & Burman, J.P. (1946). The design of optimum multifactorial experiments, *Biometrika* **33**, 305–325.
- [8] Tyssedal, J. & Samset, O. (1999). *Two-Level Designs with Good Projection Properties*, Preprint, Statistics No. 12/1999, The Norwegian University of Science and Technology.
- [9] Tyssedal, J. & Samset, O. (1997). *Analysis of the 12 Run Plackett-Burman Design*, Preprint, Statistics No. 8/1997, The Norwegian University of Science and Technology.
- [10] Tyssedal, J., Grinde, H. & Røstad, C.C. (2006). The use of a 12 run Plackett-Burman design in the injection moulding of a technical plastic component, *Quality and Reliability Engineering* **22**, 651–657.
- [11] Box, G.E.P. & Meyer, R.D. (1993). Finding the active factors in fractionated screening experiments, *Journal of Quality Technology* **25**, 94–105.
- [12] Seiden, E. & Zemach, R. (1966). On orthogonal arrays, *Annals of Mathematical Statistics* **37**, 1355–1370.
- [13] Lin, D.K.J. & Draper, N.R. (1992). Projection properties of Plackett and Burman designs, *Technometrics* **34**, 423–428.
- [14] Wang, J.C. & Wu, C.F.J. (1995). A hidden projection property of Plackett-Burman and related designs, *Statistica Sinica* **5**, 235–250.
- [15] Cheng, C.S. (1998). Some hidden projection properties of orthogonal arrays with strength three, *Biometrika* **85**(2), 491–495.
- [16] Lin, D.K.J. (1993). A new class of supersaturated designs, *Technometrics* **35**, 28–31.
- [17] Tyssedal, J. & Kulahci, M. (2005). Analysis of split-plot designs with mirror image pairs as sub-plots, *Quality and Reliability Engineering* **21**(5), 539–551.

JOHN TYSSEDAL

Randomization in Experimental Designs

In an experiment to compare several treatments or factor settings, randomization is used as a standard method for deciding which treatment is assigned to each of the experimental units. Randomization is a means to relax the condition of *ceteris paribus* (= all other things being equal). This is quite popular in sports, when the conditions that cannot be made equal for the contestants are allotted by a random process. The argument is that in a large number of cases the different random conditions will cancel out and on the average the better will win.

Therefore, a large number of papers (like e.g. [1]) stress the fact that randomization allows drawing conclusions on **causality** even if there are uncontrollable outer influences, i.e. factors that may influence the experimental outcome but cannot be controlled by the experimenter. The classical example of such a factor is soil quality in agricultural experiments. Further examples of such factors that come to mind in the case of industrial experiments are variations in the raw materials or varying ambient conditions.

To explain the idea, we consider the comparison of two treatments for bending metal plates. The variable of interest is the reduction of material thickness. There are 16 metal plates available for an experiment. Now assume that (unknown to the experimenter) the plates come from two different batches, such that the plates from the first batch are much softer, leading to a higher reduction in thickness. For simplicity, consider the situation where there is no true difference between the two methods. Clearly, if we bend more of the soft plates with method A and more of the hard plates with method B, we will observe that method A on the average produces a higher reduction of thickness. Since we do not even know that the material comes from two batches, it is impossible to allocate them evenly to the two treatments. Here randomization is the solution: If the set of plates that receive method A is a simple random sample of size 8 taken from the set of 16 plates, we observe that the mean reduction of thickness under A *on the average over all possible outcomes of the randomization* is the same as the mean reduction of thickness under B.

However, there is more to randomization than that. In randomization theory, the randomization is used to

introduce a **correlation** structure on the observations, which can be used to justify a model for the analysis of the experiment.

This observation was first pointed out by R. A. Fisher. In his book *The Design of Experiments* [2, 3], he explains how the design of an experiment determines a large part of the properties of the data. He used a sensory experiment (the “lady tasting tea”) to introduce the “principles of experimentation”, which include randomization as a central part.

When describing randomization schemes, he shows that the randomization itself justifies properties of the distribution of the data, which therefore need not be checked by an analysis of the **residuals**, but which are a mathematical fact. Fisher explains that “the purpose of randomisation . . . is to guarantee the validity of the test of significance” [3, p. 64].

To explain this concept, we return to the example with the reduction of thickness of plates. The concept is based on the idea of a hypothetical observation. For each plate i , the number η_i is the hypothetical reduction of thickness that we would measure on the i th plate, if this was bent with a standard procedure. It is clear that η_i is unknown and cannot be observed. We have no information on the distribution of the η_i . It would be rather daring to assume that they are independent identically distributed (and normal). So we make no assumption on them at all. We do make an assumption on the effects of the treatments, however. For each of the two treatments, we assume that there is an effect τ_j , which acts additively on the η_i values. More precisely, consider a plate with the hypothetical observation η . If this plate is treated with method j , then we can write the truly observed material thickness y of this plate as

$$y = \eta + \tau_j. \quad (1)$$

We might of course give each plate a name (for instance “1” to “16”). These names are assigned at random. More precisely, we put the plates on a staple and then randomly select one permutation π among the $16! = 2 \times 10^{13}$ possible permutations of 1, . . . , 16. Then the name “1” is given to the $\pi(1)$ th plate, the name “2” to the $\pi(2)$ th plate, and so on. Since π is randomly selected, $\eta_{\pi(1)}$ is a random variable and the distribution of $\eta_{\pi(1)}$ is the same as that of $\eta_{\pi(2)}$ or of any $\eta_{\pi(i)}$ for $1 \leq i \leq 16$. Also, the joint distribution of $\eta_{\pi(1)}$ and $\eta_{\pi(2)}$ is the same as that of any other pair $\eta_{\pi(i)}$ and $\eta_{\pi(j)}$ for $1 \leq i < j \leq 16$.

2 Randomization in Experimental Designs

We use an experimental design that prescribes for each *name*, what treatment the plate *with that name* will get. Formally, the experimental design d therefore is a mapping from the set of all names to the set of treatments. The plate with name i will receive treatment $d(i)$. For instance, the design might be such that plates with names “1”, “3”, ..., “15” are bent with method A, while plates with names “2”, “4”, ..., “16” are bent with method B. Owing to the randomization used, the 8 plates assigned to each treatment are a random sample from the 16.

We then do the experiment in the order in which the plates lie on the staple. Note, however, that in our notation, y_i refers to the observation on the plate with name “ i ”, not on the i th observation that is made.

For the observation y_i we get from equation (1) that

$$y_i = \tau_{d(i)} + \eta_{\pi(i)}. \quad (2)$$

Then $\eta_{\pi(1)}$ has the same expectation and variance as any other $\eta_{\pi(i)}$, while the correlation between $\eta_{\pi(i)}$ and $\eta_{\pi(j)}$ is the same as that between any other two. Therefore the vector of observations $\mathbf{y} = [y_1, \dots, y_{16}]^T$ follows a linear model: the observations have the same mean vector and the same variance–**covariance** structure as the one-way classification model with a random overall mean. That is, they have the same joint distribution as if they followed the linear model

$$y_i = \tau_{d(i)} + m + e_i \quad (3)$$

where m is a random variable with unknown mean and variance and the e_i are independent identically distributed errors. It follows that the difference between the means of the two treatments is an unbiased estimate of the difference $\tau_1 - \tau_2$ between the two effects, and that the usual estimate for the variance from the one-way classification model (3) is an unbiased estimate of its variance.

The reader should observe that the randomization procedure described above randomizes the order in which the treatments are applied: We do not bend the plate with name “1” as the first one, but the plate that lies on top of the staple. Since the order of the treatments is determined by the names, the treatments are applied in a random order. Hence, a possible time trend in the experiment will not cause **bias**, it will be part of the η . The randomization will also work in cases in which run order is really the only thing that, in principle, distinguishes one observation from

another. For example, suppose we want to run an experiment to test different techniques to throw a ball at some target. Unlike the metal plates, where there is a physical experimental unit associated with each observation, here the same ball and the same thrower is used for each experimental run. Then the randomization process will only permute the order.

Returning to the example with the plates, now assume the experimenter found out about the two batches and he/she knows which plate comes from which batch. Then we will want to assign four plates from each batch to each treatment. This can be achieved by blocking (*see* **Blocking**): the plates will be separated into two blocks of size 8 each. The names now are given by double indices “ (i, j) ”, where $i = 1$ or 2 indicates the block and $j = 1, 2, \dots, 8$ indicates the number of the plate within the block. The experimental design d might then prescribe that plates with names $(1, 1), \dots, (1, 4)$ and $(2, 1), \dots, (2, 4)$ will receive treatment 1, while the others receive treatment 2. The experimenter will separate the two groups and randomly determine which of the groups is called *block 1* and which *block 2*. Furthermore, he/she will randomly order the numbers within the batches, independently for each batch. This procedure implies that η of the plate with number $(1, 1)$ will have the same expected value and the same variance σ^2 as that of any other (i, j) , while the covariance between (i, j) and (r, s) can take one of two values: ρ_1 if $i = r$ and $j \neq s$ or ρ_2 if $i \neq r$. Note that this is the same variance–covariance structure as in a simple block model with random block effects and with a random mean

$$y_{ij} = \tau_{d(i)} + m + b_j + e_{ij} \quad (4)$$

Therefore, the analysis can be done in the simple block model.

It is important to note that the correct model for the bending experiment depends on which randomization was done. A detailed explanation of the connection between the randomization and the model can be found in [4, 5].

It is clear in the example mentioned above that blocking is done to increase precision. The variance of the estimated treatment difference is proportional to $\sigma^2 - \rho_1$, which can become small if the blocks are homogeneous, i.e. the observations from the same block are highly correlated and ρ_1 is large. In industrial experiments, there is a second motivation for the use of blocks. If the experiment is done as a series

of runs that are carried out one after the other, there always is the possibility of a time trend. Therefore, the order of the experimental runs should be randomized. However, often there are factors that are difficult to change, and it might be preferable to change the levels of a certain factor as little as possible. This can be realized by a blocked experimental design, in which the hard-to-change factor is held at the same level within each block. Say, we use a **fractional factorial design** (see **Fractional Factorial Designs**) with five two-level factors in 16 runs. Assume the first factor is hard to change. Then it might make sense to separate the runs into four groups of four each, such that the setting of factor 1 in each group is constant. Then we will randomize the order in which these four groups are done and the order of the runs in each group. This will lead to a block design, where the effect of factor 1 is confounded with the random block effect. Note that this allows **estimation** of the effect of factor 1, but that this factor has another (generally larger) variance than the other factors. Designs of that sort are called *split-plot designs* (see **Split-Plot Designs** and e.g. [6]).

As we try to point out in this article, the great advantage of randomization lies in the justification of a model for the analysis. With proper randomization a model can be made “valid”. Strictly speaking, the word *valid* means only that the estimates from that model are unbiased and that their variances (and the estimates for the variances) are the same as if that model were true (see [5]). However, randomization cannot guarantee normality of the errors and therefore it cannot guarantee normality of the estimates. Fisher himself did not consider this a large problem. In his book on experimental design, he explains the concept of permutation tests, and [3, p. 48] shows for an example that the permutation test leads to the same results as the t -test. He also quotes a paper [7] by Eden and Yates as indication that the F -test is valid, even if the data are nonnormal. Eden and Yates used data from a uniformity trial and a permutation test to compare the empirical distribution of the F statistic to the theoretical F distribution, where the empirical distribution was derived with a permutation test. A uniformity trial is an experiment where all experimental units receive the same treatment. They showed that the data from their uniformity trial were highly skewed, but that the empirical distribution function of the F statistic was very near to the theoretical function. Hence, they concluded that the

nonnormality of their data was not a problem for the validity of the F -test. Fisher [3, p. 57] quotes this experiment and states that Eden and Yates “have demonstrated how closely the theoretical distribution was verified in material that was far from normal”.

However, Fisher’s point of view is slightly too optimistic. As Yates [8, p. xix] points out himself, the Eden and Yates example in [7] is not quite as convincing as it was thought to be at the time: each original data point used for their simulation was the mean of 8 highly skewed observations. Being the mean, it was not as skewed as the single values. Yates wrote: “A random selection of one unit out of each group should have been taken, as taking the means eliminated most of the **skewness** from the data; unfortunately this was not noted at the time.” In spite of the fact that their data were not as nonnormal as it seemed, the method to check the validity of the F -test was an ingenious one, which since that time has been used frequently by other authors. This method also provides an alternative to the parametric test if the empirical distribution function should not fit the theoretical distribution, namely a variant of the permutation test, which is called *randomization test* (see e.g. [9]).

Because the numerator of the t -test or the F -test is generally the mean or a function of the mean, the validity of the normality assumption is largely supported by the central limit theorem, in all but the smallest examples. This explains why in the majority of cases normality is indeed not a problem. If problems arise, they are usually due to **outliers** in the data or due to the fact that the number of different possible permutations may be too small.

For small fractional factorial designs this can indeed be the case. Note that in a factorial design with all factors at two levels, each column of the design matrix (other than that corresponding to the general mean) consists of $n/2$ entries of 1 and $n/2$ entries of -1 . Therefore, there are only $n(n-1)\dots(n-n/2+1)/(n/2)!$ different columns possible and these can be divided into pairs of columns that can be transformed into each other by multiplication with -1 . An important analysis of fractional factorials is based on the half-normal plot (see **Half-Normal Plot** and [10]). In this analysis, the maximum of the absolute values of all estimated contrasts is compared to an estimate of the variance, see e.g. [11] for possible estimates of the variance and for critical values. Note that each fractional factorial design with n runs uses $n-1$ such columns. Therefore the

number of permutations that lead to different maximal contrasts gets relatively small. For instance, if $n = 4$, then there are only three different possible columns and each factorial design uses all of them. Hence the maximum contrast is constant – and the distribution of the maximum is clearly not normal. If $n = 8$, there are 35 different contrasts and each design uses 7. This implies that there are at most five different values of the maximum. So the normality assumption for the maximum contrast in an eight-run factorial design cannot be justified by randomization. For $n = 16$, however, the number of different randomization results gets sufficiently large: there are 6435 different contrasts and each design uses 15, which allows for 429 different values of the maximum. Kunert and Adekeye [12] did a simulation study, where they used a real data set from an experiment with a time trend and the permutation method introduced by Eden and Yates [7]. They demonstrate that the randomization distribution of the maximum contrast of a 16-run 2^{k-1} design fits well to the distribution that it should have under normality. They also show that the randomization of the run order can successfully deal with a time trend.

References

- [1] Rosenbaum, P.R. (2005). Heterogeneity and causality: unit heterogeneity and design sensitivity in observational studies, *The American Statistician* **59**, 147–152.
- [2] Fisher, R.A. (1935). *The Design of Experiments*, First Edition, Oliver & Boyd, Edinburgh.

- [3] Fisher, R.A. (1971). *The Design of Experiments*, Reprint of the 8th Edition, Hafner Publishing Company, New York.
- [4] Bailey, R.A. (1981). A unified approach to design of experiments, *Journal of the Royal Statistical Society A* **144**, 214–223.
- [5] Bailey, R.A. & Rowley, C.A. (1987). Valid randomization, *Proceedings of the Royal Society A* **410**, 105–124.
- [6] Bingham, D. & Sitter, R. (2001). Design issues in fractional factorial split-plot experiments, *Journal of Quality Technology* **33**, 2–15.
- [7] Eden, T. & Yates, F. (1933). On the validity of Fisher's z test when applied to an actual example of non-normal data, *Journal of Agricultural Science* **23**, 6–17.
- [8] Yates, F. (1990). *Foreword to R.A. Fisher: Statistical Methods, Experimental Design and Scientific Inference*, Oxford University Press, Oxford.
- [9] Edgington, E.S. (1995). *Randomization Tests*, 3rd Edition, Dekker, New York.
- [10] Daniel, C. (1959). Use of half normal plots in interpreting factorial two level experiments, *Technometrics* **1**, 311–341.
- [11] Kunert, J. (1997). Use of the factor sparsity assumption to get an estimate of variance in industrial experiments with many factors, *Technometrics* **39**, 81–90.
- [12] Adekeye, K. & Kunert, J. (2006). On the comparison of run orders of unreplicated 2^{k-p} -designs in the presence of a time trend, *Metrika* **63**, 257–229.

Related Articles

Blocking; Fractional Factorial Designs; Half-Normal Plot; Split-Plot Designs.

JOACHIM KUNERT

Factorial Designs, Resolution of

Introduction

The concept of design resolution was introduced by Box and Hunter in 1961 [1] in the context of two-level *fractional Factorial designs* (see **Fractional Factorial Designs**) of the type 2^{k-p} with the intent of defining a general index by which these designs could be easily ranked. The index, the resolution R , is very helpful to users at the stage of design selection. In fact, maximum resolution designs are the most commonly adopted in practice. We briefly recall some notions for an easier understanding of resolution. A 2^{k-p} design allows us to study the effects of k two-level factors on the response by running only a 2^{-p} fraction of the corresponding full factorial design. However the reduction in size (2^{k-p} runs instead of 2^k) is trade-off against a reduction in information. The latter is due to *aliasing of effects* (see **Aliasing in Fractional Designs**), meaning that the 2^k effects (*main effects* (see **Main Effects**) plus *interactions* (see **Interactions**) up to order k plus the constant effect) are split into 2^{k-p} groups containing 2^p fully undistinguishable effects (effects in each group are estimated by the same linear combination of the observations). These groups form the aliasing pattern of the design. Thus an individual effect can be estimated only if the effects with which it is aliased are sufficiently smaller than it. Hence, if we rely on the empirical principle that lower order effects are more likely to be important than higher-order effects (hierarchical ordering principle), a good fraction is one that protects low order effects – typically main-effects and two-factor interactions – by aliasing them with higher-order effects. Roughly speaking, resolution measures the degree to which this desirable aliasing pattern is realized by the design. Following Wu and Chen [2] it is instructive to call a main-effect or a two-factor interaction *clear* if it is not aliased with any main-effect or any two-factor interaction and *strongly clear* if it is also not aliased with any three-factor interaction.

Definition, Rationales, and Use of Resolution

A design is of resolution R if any effect containing q factors is not aliased with any other effect containing fewer than $R - q$ factors. R is denoted by a capital Roman numeral, usually appended as a subscript to the design indication, e.g., 2_{IV}^{4-1} designates a resolution IV half-fraction of the full factorial with four factors. We now focus on designs of resolution III, IV, and V, the most widespread in practice.

Resolution III designs

Main effects are not aliased with one another but some main effects are aliased with two-factor interactions. Thus some main effects and some two-factor interactions are not clear.

Resolution IV designs

Main effects are not aliased with one another or with any two-factor interaction. Some main effects are aliased with three-factor interactions and some two-factor interactions are aliased with other two-factor interactions. Thus all main effects are clear and some two-factor interactions are not clear.

Resolution V designs

Main-effects and two-factor interactions are not aliased with any other main-effects or two-factor interactions. Some main effects are aliased with four-factor interactions and some two-factor interactions are aliased with three-factor interactions. Thus all main effects are strongly clear and all two-factor interactions are clear.

Therefore, while resolution III designs are used in practice primarily for screening purposes (see **Screening Designs**), i.e., to rule out inactive factors from subsequent experiments, resolution IV and V designs may also be used for response prediction. In fact, assuming that main-effects and two-factor interactions are sufficiently larger than higher-order interactions, resolution IV designs capture the linear behavior in the factors of the response while resolution V designs also capture the interactive part of the possible nonlinear behavior in the factors. These are guidelines for users to choose a design resolution consistent with the objective of the experimentation. In practice an efficient approach to knowledge building by designed experiments is

the sequential one [3]. It advises to spend the total experimentation budget on a series of related experiments, each one designed adaptively, i.e., based on information drawn from the previous ones. A typical sequence of experiments would foresee an initial screening phase where highly fractionated factorials rule out inactive factors. This would be followed by higher resolution fractions allowing for the identification of a good operating region in the factor space. Finally, an optimum factor setting would be obtained by fitting a quadratic model to data coming from the efficient three-level designs recommended by the response surface methodology (RSM, *see Response Surface Methodology*). Thus the increase of resolution is a guiding mechanism in *sequential experimentation* (*see Sequential Experimentation*). A useful result in this context is that, after running a resolution III fractional design, main effects can be easily de-aliased from two-factor interactions. In fact, if the same design is run again but switching the signs of levels in all factor columns (*foldover* technique, *see Foldover Designs*), the ensemble of the two designs is a resolution IV design.

It is also worth mentioning a group of resolution III designs that are particularly efficient in terms of run size. If the number of factors k is one less than a power of 2, a resolution III design with only $k + 1$ runs is available. Thus one may study the main effects of seven factors with an 8-run 2^{7-4} , or the main effects of 15 factors with a 16-run 2^{15-11} , and so on. These designs are also a subset of the *Plackett–Burman designs* (*see Plackett–Burman Designs*).

Interestingly, the importance of resolution can be appreciated from another viewpoint. An appealing property of fractional designs is that when a design of resolution R is projected in any $R - 1$ (or less) factors, a full factorial design is generated in those factors (*see Projectivity in Experimental Designs*). (Note that the projection of the design in a subset of the original factors is another design obtained by omitting in the original design the columns of the factors in the complementary subset.) Therefore, if one believes that not more than $R - 1$ of k investigated factors are active (but without knowing which!) a design of resolution R is a good choice: if the conjecture is true, all the effects of the active factors can be estimated without aliasing.

Determination of Resolution

An easy way to determine the resolution of a two-level fractional factorial exploits the fact that it is equal to the number of letters of the shortest word in the *complete defining relation* of the design (sometimes referred to as the *defining contrast subgroup*). The latter is an equation expressing the aliasing in the group of 2^p effects whose contrast column is all 1's. The equation contains only p independent words called *generators of the design* (*see Fractional Factorial Designs*). The constant effect, denoted by I , is referenced as the first in the group but it does not count as a word for determining the resolution. As an example consider the 2^{8-2} design with generators $ABCD$ and $EFHG$. As the complete defining relation $I = ABCD = EFGH = ABCDEFGH$ has two words of length 4 and one word of length 8 the design is of resolution IV. However, the highest possible resolution for a 2^{8-2} design is V. It is obtained by choosing $ABCDG$ and $ABEFH$ as generators. This gives the complete defining relation $I = ABCDG = ABEFH = CDEFGH$ with two words of length 5 and one word of length 6. A notable consequence of the operational definition of resolution is that the maximum resolution of half-fractions (2^{k-1} designs) is equal to the number of factors, k . In fact, because these designs have only one generator, the maximum resolution is simply obtained by choosing the contrast of the highest order interaction as the generator. As an example, the maximum resolution for a 16-run 2^{5-1} design is V and is obtained by choosing the generator $ABCDE$ (and defining relation $I = ABCDE$).

Other Criteria for Design Selection

Resolution is not the only criterion used to select a 2^{k-p} design. A first instance occurs when, for given k and p , there is more than one design of maximum resolution. In such cases the subcriterion of *minimum aberration* (*see Minimum Aberration*) is usually applied. It amounts to choosing the maximum resolution design with the smallest number of words with minimum length. This ensures that the design has the minimum number of main effects aliased with interactions of order $R - 1$, the minimum number of two-factor interactions aliased with interactions of order $R - 2$, and so

on [4]. Consider the three 32-run 2^{7-2} designs defined by the complete relations $I = ABCF = ABDG = CDFG$, $I = ABCF = ADEG = BCDEFG$, and $I = ABCF = ABDEG = CDEFG$ respectively. As the minimum word length in each relation is four, the three designs are of resolution IV, the highest resolution for a 2^{7-2} . Therefore, while main effects are all clear, two-factor interactions are not. However, the third design is also of minimum aberration as its defining relation has only one word of minimum length, while the first design has three and the second has two. This means that the third design has a higher number of clear two-factor interactions than the first and the second design have. In fact, it has 15 clear two-factor interactions out of 21 (all but the three aliased pairs $AB = CF$, $AC = BF$, and $AF = BC$), while the second design has 9 and the first only 6. It may still be that a design with maximum resolution and minimum aberration does not maximize the number of clear effects. For such circumstances Chen *et al.* [5] proposed the number of clear effects as an alternative criterion. A simple illustration refers to the class of 2^{6-2} designs. The maximum resolution design in this class (complete relation $I = ABCE = ABDF = CDEF$) has resolution IV, and only the six main effects are clear. Yet, the 2_{III}^{6-2} design with generators ABE and $ACDF$ (complete relation $I = ABE = ACDF = BCDEF$) is only of resolution III but has nine clear effects, three main effects (C , DF) and six two-factor interactions (BC , BD , BF , CE , DE , EF). Finally, there are cases when we are interested in protecting specific effects which we believe are more important in the application at hand. This usually requires a close scrutiny of possible aliasing patterns since turnkey solutions are not available. One exception is when it is important to estimate clearly the main-effect and all the two-factor interactions involving one particular factor. In this circumstance a simple rule is to choose a design whose generators do not involve that factor. One example is the 16-run 2_{III}^{6-2} design whose generators ABE and ACF do not include factor D . Although the design is of resolution III, one can easily show from the complete defining relation ($I = ABE = ACF = BCEF$) that the **main effect** of factor D is strongly clear and the five two-factor interactions involving D are clear. This is not so for the maximum resolution 2_{IV}^{6-2} design seen before. For more details refer to [6, Section 4.5, p.177]. Tables containing maximum resolution (with minimum aberration) designs can be

found in [7, Section 8–4, pp. 328–330], while [6, Appendix 4A, pp. 193–199] provides more comprehensive tables for $k-p$ in the range 3–6 and k up to 32, where designs maximizing the number of clear effects are also reported.

Resolution may also apply to fractions of three (or higher) level factorial designs. Readers interested in the extension of resolution in those settings may refer to [6, Sections 5.4 and 5.6].

References

- [1] Box, G.E.P. & Hunter, S.J. (1961). The 2^{k-p} fractional factorial designs, *Technometrics* **3**, 311–351 and 449–458.
- [2] Wu, C.F.J. & Chen, Y. (1992). A graph-aided method for planning two-level experiments when certain interactions are important, *Technometrics* **34**, 162–175.
- [3] Box, G.E.P. & Wilson, K.B. (1951). On the experimental attainment of optimum conditions, *Journal of the Royal Statistical Society. Series B* **13**, 1–45.
- [4] Fries, A. & Hunter, W.G. (1980). Minimum aberration 2^{k-p} designs, *Technometrics* **22**, 601–608.
- [5] Chen, J., Sun, D.X. & Wu, C.F.J. (1993). A catalogue of two-level and three-level fractional factorial designs with small runs, *International Statistical Review* **61**, 131–145.
- [6] Wu, C.F.J. & Hamada, M. (2000). *Experiments: Planning, Analysis, and Parameter Design Optimization*, John Wiley & Sons, New York.
- [7] Montgomery, D.C. (2001). *Design and Analysis of Experiments*, 5th Edition, John Wiley & Sons, New York.

Further Reading

NIST/SEMATECH e-Handbook of Statistical Methods, <http://www.itl.nist.gov/div898/handbook/>, 2006. Design resolution is found in a sub-section of the electronic document: 5.3.3.4.4. *Fractional factorial design specifications and design resolution* (<http://www.itl.nist.gov/div898/handbook/pri/section3/pri3344.htm>).

Related Articles

Aliasing in Fractional Designs; Foldover Designs; Fractional Factorial Designs; Interactions; Main Effects; Minimum Aberration; Plackett–Burman Designs; Projectivity in Experimental Designs Response Surface Methodology; Screening Designs; Sequential Experimentation.

DANIELE ROMANO

Sample-Size Determination in Experimental Designs

Introduction

Sample-size problems deal with the general issue of collecting enough data to achieve a stated goal. The importance of **sample size** varies considerably, depending on context. There are several possible criteria for sample size, the two most popular being specifying the width of a **confidence interval** or the **power** of a test with a specified effect size. Those are the main areas of emphasis in this article.

With either approach, it is usually necessary to have prior information concerning one or more variances; thus, one generally needs a pilot study or historical data where measurements are made using the same procedures, instrumentation, and comparable materials or subjects. Conversely, it can be argued that in the absence of such historical pilot data, one should not expect to be in a position to plan a definitive study, and hence sample size is not a very important issue. Indeed, if one's measurements and procedures have not been tested, there is much that can go wrong besides having an inadequate sample size. A pilot study need not be of the same design as the study being planned.

Useful general references for sample size include [1, 2], and [3]. A particularly useful recent reference for the power approach is [4]; while it emphasizes clinical experiments, it covers a great deal of common ground with other areas of application. It also contains a good discussion on study objectives (e.g., testing for significance, noninferiority, equivalence) that deserves the attention of readers in all areas of research.

The treatment in this article emphasizes frequentist methods for **estimation** and **testing**. Good sample-size software is now widely available, whereas many formulas are only rough approximations and apply to very limited – and usually oversimplistic – scenarios. For these reasons, this article emphasizes the use of software. Several standard statistical packages (including Statistical Analysis System (SAS) and *Minitab*) offer good sample-size capabilities; there are also commercial programs specifically for sample

size, such as *nQuery Advisor* and *NCSS-PASS*, and on-line power and sample-size calculators such as *Piface* [5].

Simple Illustration

Suppose that we plan to collect paired data on the percentage of solids at two locations in a pipeline used to transport a liquid. We believe that the paired differences are approximately normal, and we plan to use a paired *t* procedure to analyze the data. We have past data (see [6]) on another similar pipeline on locations about the same distance apart, and those data suggest that the standard deviation σ of the paired differences is about 0.14%.

The percentages of solids are quite small, but somewhat critical. After discussing the matter with the research team, it is decided that we would like to know if there is a 0.10% or greater difference in the means at the two locations. How much data are required?

One traditional approach is that of specifying the width or half-width of a **confidence interval**. In this case, the 95% confidence interval half-width for a sample of n pairs will be $t_{0.025, n-1}s/\sqrt{n}$, where s is the sample standard deviation, and $t_{0.025, n-1}$ is the critical value for a 0.025 tail on the *t* distribution with $n - 1$ **degrees of freedom** (df). If we want the half-width to be 0.10 or smaller, then by solving for n , we obtain $n \geq (t_{0.025, n-1}s/0.10)^2$. We estimate that s will be about 0.14; the *t* critical value depends on n , but only weakly if n is moderate or large. Thus, if we guess that n will be about 30, say, then we obtain $n \geq (2.045 \times 0.14/0.10)^2 \approx 8.2$. Thus, perhaps $n = 9$ is adequate; but this estimate is based on 29 df. Correcting the degrees of freedom to $9 - 1 = 8$, we get $n \geq (2.306 \times 0.14/0.10)^2 \approx 10.4$. Correcting yet again to 10 df, we obtain $n \geq 9.7$. It appears that $n = 10$ will be adequate for obtaining a confidence interval of the desired width.

Note that even in this simple situation, there is a certain amount of handwaving when we substitute 0.14 for a future value of s , plus some iterating on the degrees of freedom. Our previous data may have under- or overestimated σ ; in addition, the future value of s in the planned experiment may be larger or smaller than 0.14. We may want to take such considerations into account; for example, instead of using the point estimate of 0.14 for σ , perhaps an

2 Sample-Size Determination in Experimental Designs

upper confidence limit for σ would be better to ensure that we collect enough data. Such decisions are subjective and based largely on how important it is in this particular case to get the sample size right.

The other traditional approach is to choose n so as to achieve a specified *power* of the t -test when there is a 0.10% difference between the means. The calculations here are more complex, and it is best to use software to solve the problem. Typically, we specify a power of 0.80 or 0.90. The output from SAS (using PROC POWER) is shown in Table 1. For a power of 0.80, we need $n = 18$. To achieve a power of 0.90, we need a lot more data ($n = 23$): once the power exceeds 0.7 or 0.8, there is a large cost for a small increase in power. Note also that the sample size for a power of 0.50 is $n = 10$, the same as we obtained for the confidence interval approach. This is approximately true in general: the sample size for 50% power in a two-sided test is about the same as that for a $(1 - \alpha)$ confidence interval of the same half-width.

SAS also provides for a powerlike concept in confidence intervals – that of specifying the probability that the interval is no wider than a specified value; or the conditional probability thereof, given that the interval contains the actual parameter value (see [7]). Unlike the previous confidence interval approach, this takes into account the variation in s . The output is shown in Table 2. Again, we obtain $n = 10$ for a “power” of 0.50. We need a somewhat larger value of n as the requirements increase; and with $n = 15$, there is a very good (conditional) chance that the half-width of the confidence interval will be 0.10 or

less. When the focus of a planned study is on estimation, this approach deserves more attention than it currently gets.

When is Sample Size Important?

An *observational study* or *experiment* is considered underpowered when there is a low chance of obtaining statistical significance for an effect of such magnitude as to be considered important scientifically; and it is overpowered when a scientifically unimportant effect would still be statistically significant. Power increases as sample size increases, so one can control the power of a study by choosing the sample size.

It is always important to avoid overpowering a study – that is, we should not collect too much data. At best, it is wasteful of resources; and in cases where human or animal subjects are involved, too much data may translate into too many subjects being exposed to a potentially harmful treatment, or being denied a beneficial one.

Having an adequate sample size is most important when a definitive study is planned and there are substantial costs (in money or time or otherwise) involved in collecting more data later. For example, in a situation where it takes a long time to carry out an experiment, it also takes a long time to know whether there are any clear conclusions, and yet more time to run a follow-up experiment to establish results at greater precision.

Table 1 SAS PROC POWER output for the power approach in the paired t -test

Distribution	Normal		
Method	Exact		
Mean	0.1		
Standard deviation	0.14		
Number of sides	2		
Null mean	0		
Alpha	0.05		
Computed N total			
Index	Nominal power	Actual power	N total
1	0.5	0.522	10
2	0.8	0.815	18
3	0.9	0.905	23

Table 2 SAS PROC POWER output based on the “power” of the paired t confidence interval

Distribution	Normal		
Method	Exact		
CI half-width	0.1		
Standard deviation	0.14		
Number of sides	2		
Alpha	0.05		
Probability type	Conditional		
Computed N total			
Index	Nominal probability (width)	Actual probability (width)	N total
1	0.5	0.542	10
2	0.8	0.834	13
3	0.9	0.942	15

Sample size is least important in pilot studies, where the emphasis is less on obtaining definitive scientific results and more on establishing that one's procedures work and estimating one or more error variances for a later, more definitive study. Adequacy of sample size is also relatively unimportant in a situation where *sequential experimentation* is practical – where additional data can be collected readily as needed. However, even then, one must not underestimate administrative barriers – work shifts, funding, availability of personnel, and so on – that might significantly impede the collection of more data.

Almost any course of investigation begins with at least one stage – piloting – where sample size is relatively unimportant, and proceeds to one or more final stages where sample size is relatively important. The early stage(s) establishes the procedures and knowledge needed to assess the sample-size requirements for the later, definitive stage(s). This discipline of piloting and planning is necessary for a successful investigation, and it is a mistake to try to cut it short.

Eliciting Effect Size

One of the toughest hurdles one faces in addressing a sample-size problem is the specification of effect size (or confidence interval half-width). This needs to be based on scientific rather than statistical concerns; but usually the statistician needs to facilitate the conversation on this issue. If the question is asked naively (e.g., “What difference would you consider scientifically important?”) the answer is likely to be equally naive (e.g., “any difference at all” or “you tell me”). A slightly oblique approach is more likely to be successful; for example, ask “How different can these two groups be, and still be considered the same in a practical sense?”

Because effect size is such a daunting question, people tend to seek an easy way out – for example, using a predefined “large”, “medium”, or “small” effect size (see [8]). These “T-shirt” effect sizes are expressed in standard deviation units and are based on surveys of published research in the social sciences. This seems to be a way of following the norms of existing research; however, thinking about this carefully, it becomes clear that choosing a “medium” effect size will always yield the same sample size as a medium-sized study in the social science literature – no calculations are needed. Moreover, this

simpleminded approach yields the same sample size regardless of variation in the instrumentation used to obtain measurements – whereas intuitively, the more precise the measurements, the less data are needed. If the study is important, so is its planning, and for that reason, the focus of the discussion should be on the actual scale of measurement (or some transformation thereof), and the question of effect size should be considered seriously.

An alternative way to specify effect size is through a set of exemplary data—an artificial data set that contains no variation but exhibits a non-null case that would be considered to be of minimal scientific significance. A simple example is that of testing a six-sided die for fairness. A fair die has a probability of $1/6 \approx 0.1667$ for each outcome. An exemplary unfair die might have probabilities 0.15, 0.15, 0.15, 0.15, 0.20, 0.20, so that a corresponding exemplary data set would have proportional frequencies – say 3, 3, 3, 3, 4, 4. The χ^2 goodness-of-fit statistic (see **Statistical Methods for Counts, Rates, and Proportions**) for these 20 exemplary observations is 0.4, and it turns out that the noncentrality parameter λ of the χ^2 statistic based on n observations is proportional: $\lambda = (0.4/20)n$. The Piface user interface is displayed in Figure 1, where we find that 578 die tosses are required to establish the exemplary die to be unfair, with a power of 0.75 and a **significance level** of 0.05.

Using Pilot Data

When the planned study involves measurement data, one must have estimates of one or more error variances in order to determine sample size. Accordingly, one must either have an existing data set for a similar situation or conduct a pilot study to estimate these variances. A pilot study can be beneficial for other reasons, such as checking that the planned experimental procedures actually work as anticipated.

The pilot study need not be of the same design as the planned study. The only requirement is that it should estimate the required variances. However, hidden sources of variation must be carefully considered to make sure that a single variance in one study is not the sum of several variances in the other. For example, suppose a pilot study entails one measurement of hardness on each of several ingots of steel, while the planned study involves several hardness measurements on the same ingot. The pilot study

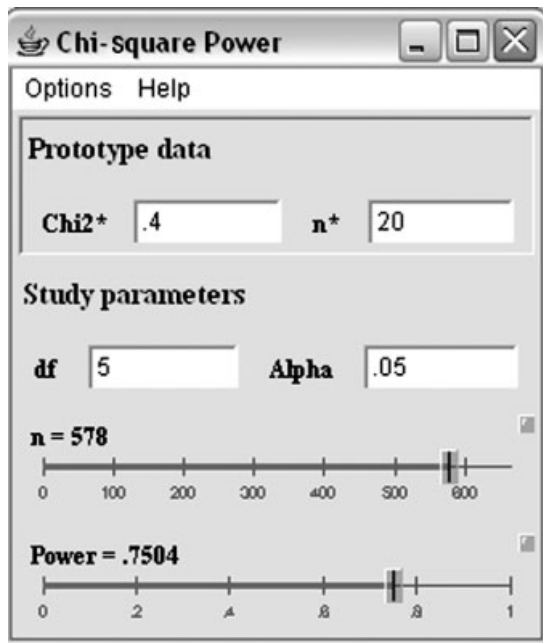


Figure 1 Piface interface for a χ^2 test

estimates the sum of the between- and within-ingot variances, whereas the planned study considers these variances separately.

One question that comes up is to what extent the results of the pilot study should affect the planning of the main study. For example, the pilot will yield an estimate of the effect of interest – should this estimate be used in sample-size planning? Strictly speaking, the answer is “no”, in deference to the effect size we have so carefully elicited based on scientific concerns. However, if the effect size observed in the pilot is much greater than the scientifically important one, it suggests that the main study can be downsized, thus saving cost and reducing exposure of subjects. Accordingly, it is reasonable for the sample-size planning to be based on a lower confidence bound on the effect size, provided that this bound is at least as large as the scientifically important effect size.

Multifactor Example

An experiment is to be conducted on a semiconductor manufacturing process; the response variable of interest is oxide thickness (in Angstroms, denoted Å). The plan for the experiment is to be a *split-plot*

design: silicon wafers will be processed in lots of four wafers each, each wafer being randomly assigned to one of four treatments. Subsequently, oxide thickness will be measured at each of five fixed sites on each wafer (the split-wafer part of the design). The mixed model to be used for the analysis has fixed terms for treatment, site, and treatment $\times$ site, and random effects for lot, lot $\times$ treatment, and residual error.

The sample-size question is simply how many lots should be included in the experiment. But it is a complicated question because three fixed effects came under discussion in the effect-size elicitation process. First, it was decided that we need to detect a difference of ± 10 Å between any two treatments. Second, we want to be able to detect a difference of ± 5 Å between any two sites; and finally, we want to detect a difference of ± 15 Å between any two treatment $\times$ site combinations (see **Interactions**). These effect-size choices are based partly on considerations of the fact that whole-wafer comparisons are less accurate than split-wafer comparisons.

We have past data on the same process, but with a different experimental design – a nested layout with eight lots, three random wafers in each lot, and three random sites in each wafer. The data are given in Littell *et al.* [9], and the analysis provides variance-component estimates of 129.9 for lots, 35.87 for wafers nested in lots, and 12.57 for sites nested in wafers. In the planned experiment, the random effect of lot $\times$ treatment (i.e., the whole-wafer error) corresponds exactly to wafer-to-wafer variations in the pilot data, and so we estimate its standard deviation to be $\sqrt{35.87} \approx 6.0$. The standard deviation of $\sqrt{12.57} \approx 3.5$ for sites within wafers in the pilot data does not correspond to a random effect in the planned study; however, since it includes variations of randomly selected sites, it is reasonable to believe that it overestimates the residual variation for the planned experiment; so using it will provide a conservative assessment for sample-size purposes.

Figure 2 shows the related Piface dialog. The various sample sizes are in the left panel and the standard deviations are in the center panel. In the right panel, we select what effect to study, specify the effect size, and observe the power. In the figure, it is set to study the treatment $\times$ site comparisons, with a restriction that we are comparing two treatments on the same site using the honestly significant difference (HSD)

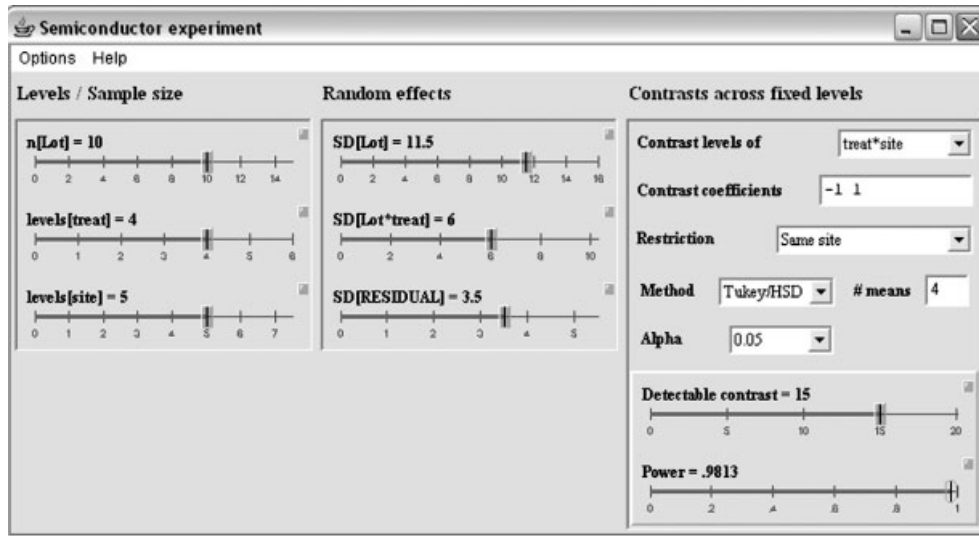


Figure 2 Piface dialog for the split-plot example

procedure. When there are 10 lots in the experiment, the power is quite high for this setting.

There are several other tests that need to be studied, and we need to consider each of them and look for the worst case. Using other settings that are not shown in Figure 2 (all with 10 lots and the HSD method), we find that the power is nearly 1 when comparing two sites with the same treatment (a split-plot comparison has better resolution); the power drops to about 0.84 when no restrictions are placed on the interaction comparisons (because now we have 20 means in the HSD procedure). For the treatment comparison, the power is about 0.80 when the difference is 10 (as per our specifications); and for the site comparisons, the power is nearly 1 when the difference between two site means is 5. Thus, the worst case is the treatment comparison, and we have found that 10 lots is adequate to achieve 80% power for that comparison, using Tukey's HSD.

Allocation

The allocation of sample sizes among several groups is an issue related to sample size. In settings where the variances are equal and all groups are to be compared on an equal basis, a balanced allocation is best. However, if the primary goal is to compare each of k treatment means with a control mean, it

can be shown that for a given total sample size, the **standard error** of the treatment-minus-control comparisons is minimized when the sample size of the control group is $\sqrt{k}$ times the sample size of each treatment group, assuming equal variances. If the variances are unequal and if all means are to be compared symmetrically, the best allocation is in proportion to the standard deviations. These rules can be combined; for example, if we have two treatments to be compared with a control, and the respective standard deviations are σ_{t_1} , σ_{t_2} , and σ_c , the best allocation of sample size is in proportion to σ_{t_1} , σ_{t_2} , and $\sqrt{2}\sigma_c$.

Cost considerations can also drive allocation of sample size. For example, in a nested experiment we may incur a certain cost each time we change to a new batch of raw material, and we have the option of running several tests on each batch at some other cost per test. One can use appropriate optimization techniques to find the number of batches and tests per batch to maximize the power of some test, subject to a constraint on the total cost of the experiment.

Planning a Pilot Study

As explained earlier, it is important in most sample-size problems to have good estimates of one or more standard deviations. That puts us in a circular

situation, in that pilot estimates of standard deviations are themselves subject to sampling error. What then is the required sample size for a pilot study?

One way to address this is to note that the principal risk we face here is that of underestimating the needed sample size. Suppose that if the true error standard deviation σ were known, we could determine that a sample size n is required for our study. The requirements of the pilot study might then be stated as follows: We are willing to take a risk γ of obtaining a sample size of only $\tau \cdot n$, where $\gamma < 1$ and $\tau < 1$ are specified.

In most settings, the required sample size is roughly proportional to σ^2 . Suppose that S^2 is an estimator of σ^2 based on a pilot study. Under a normal-theory model, $\nu S^2/\sigma^2$ has a χ^2 distribution with ν degrees of freedom. We require that $P(S^2/\sigma^2 < \tau) \leq \gamma$, and thus that $P(Y < \nu\tau) \leq \gamma$, where Y has a χ^2_ν distribution. Using this result we can find, for example, that an estimator S^2 having $\nu = 78$ df is sufficient to achieve a risk of $\gamma = 0.10$ of obtaining a sample size only 80% of what is needed ($\tau = 0.8$). Reducing τ to 0.7 reduces the needed degrees of freedom to 33, for the same risk. Piface provides a simple user interface for this method.

We can also use these results to obtain a “fudge factor” to adjust a sample-size calculation. For example, suppose that, based on a variance estimate having 33 df, we calculate a sample size of $n^* = 75$ is needed to achieve certain requirements. If we conduct the study instead with a sample size of $n^{**} = n^*/0.7 \approx 107$, then there is only a 10% chance that this sample size is too small to meet the stated requirements of the study (see the example in the preceding paragraph).

Conclusions

Sample-size determination starts with scientific rather than statistical considerations: we must evaluate the goals of a planned study in scientific terms. In addition, a definitive study can rarely be planned without first collecting pilot data. In many sample-size problems such as the multifactor example in this article, there are multiple targets to consider. New

software makes it increasingly easy to deal with all of these complexities.

The power approach to sample size becomes confounded somewhat by some other uses (and abuses) of power calculations. For instance, some researchers attempt to use power calculations as an aid in data analysis. This is generally a mistake, unless it is done prospectively. The reader is referred to [10] for a discussion of such matters.

References

- [1] Odeh, R. & Fox, M. (1991). *Sample Size Choice: Charts for Experiments with Linear Models*, 2nd Edition, Marcel Dekker, New York.
- [2] Lenth, R. (2001). Practical guidelines for effective sample-size determination, *The American Statistician* **55**, 187–193.
- [3] Parker, R. & Berman, N. (2003). Sample size: more than calculations, *The American Statistician* **57**, 166–170.
- [4] Chow, S.-C., Shao, J. & Wang, H. (2003). *Sample Size Calculations in Clinical Research*, CRC Press, Boca Raton.
- [5] Lenth, R. (2006). Java applets for power and sample size [computer software] <http://www.stat.uiowa.edu/~rlenth/Power/>, retrieved June 25, 2007.
- [6] Mason, R., Gunst, R. & Hess, J. (1989). *Statistical Design & Analysis of Experiments*, John Wiley & Sons, New York, p. 271.
- [7] Beal, S. (1989). Sample size determination for confidence intervals on the population mean and on the difference between two population means, *Biometrics* **45**, 969–977.
- [8] Cohen, J. (1988). *Statistical Power Analysis for the Behavioral Sciences*, 2nd Edition, Academic Press, New York.
- [9] Littell, R., Milliken, G., Stroup, W. & Wolfinger, R. (1996). *SAS System for Mixed Linear Models*, SAS Institute, Cary, pp. 155–157.
- [10] Hoenig, J. & Heisey, D. (2001). The abuse of power: the pervasive fallacy of power calculations for data analysis, *The American Statistician* **55**, 19–24.

Related Article

Confidence Intervals; Hypothesis Testing; Power; Sample-Size Determination.

RUSSELL V. LENTH

Screening Designs

Introduction

In the planning stages of any experimental study, a number of critical questions need to be asked. Among these are questions regarding the choice of experimental factors and factor levels to include in the study. A factor is a variable that is studied in the experiment and is believed to potentially have an effect on the response. The values of the factor at which experimental runs are conducted are called the *levels* of the factor. The actual list of factors to initially include in a study usually arises as the result of discussions with scientists or engineers most familiar with the process being studied. When the list is long, it is often the case that a relatively small subset of the factors actually has an influential effect on the response, and it is important, in the early stages of the experimental study, to identify those influential factors. This is the purpose of screening experiments. Once the influential variables have been identified, a more elaborate experimental study is usually conducted to develop a more detailed mathematical model which can be used to further describe relationships between the influential variables and the response.

A factor to be included in a study may be either quantitative or qualitative. Quantitative factors such as time or temperature can assume all values on a continuous scale. Qualitative factors such as different machine types or different crop varieties have levels which can assume only discrete values and typically cannot be rank ordered. In screening experiments, most factors are tested at no more than three levels and more typically occur at two levels. For a quantitative factor, two levels allow for the detection of a linear effect on the response. When it is suspected that a quantitative factor may affect the response nonlinearly, three or more levels of the factor can be chosen and used to detect curvature in the factor effect. There is usually much less freedom in the selection of levels for a qualitative factor as these levels are determined by the qualitative nature of the particular factor. For example, if four different machine types are being considered in an industrial experiment, then the qualitative factor, machine type, would have four levels.

The most popular screening experiments in use today are factorial experiments (*see* **Factorial Experiments**). A factorial experiment is one whose treatments consist of different combinations of experimental factor levels. In the rest of this article, the use of factorial screening designs to identify influential factors in the early stages of an experimental study is discussed.

Notation and Model

In this article, experimental situations are considered in which m factors are to be studied using a factorial experiment with N runs. We use d to denote a design which can be used in such a situation and interchangeably represent d by $X_d = (x_{d1}, \dots, x_{dm}) = (x_{dij})$, which is an $N \times m$ matrix whose i th row corresponds to run i and j th column corresponds to factor j . If factor j has s_j levels, then the entries in column j of X_d will consist of symbols $1, 2, \dots, s_j$ with $x_{dij} = t \in \{1, 2, \dots, s_j\}$, if factor j occurs at level t in run i . A useful model for analyzing the data from a given design d is

$$Y = \beta_0 \mathbf{1}_N + \sum_{i=1}^m X_{di} \beta_i + \epsilon \quad (1)$$

where Y is an $N \times 1$ vector of observations, β_0 is the overall mean effect, $\mathbf{1}_N$ is an $N \times 1$ vector of 1's, X_{di} is an $N \times (s_i - 1)$ matrix corresponding to factor i having s_i levels, β_i is an associated set of main effect (*see* **Main Effects**) parameters, and ϵ is a vector of random error terms whose entries are assumed to be uncorrelated with mean zero and constant variance σ^2 . The actual entries of X_{di} usually depend upon whether factor i is a quantitative or qualitative factor. If factor i is a quantitative factor with s_i levels, the entries of X_{di} are usually coded (using orthogonal polynomials) so that all polynomial effects up through degree $s_i - 1$ can be estimated orthogonally to one another. For example, if factor i has three equally spaced levels all occurring in $N/3$ runs, then X_{di} has two columns with entries $-1, 0, 1$ and $1, -2, 1$ depending on whether factor i occurs at its low, middle, or high levels. If factor i is a qualitative factor with s_i levels, then the entries of X_{di} can be coded in several different ways. We shall assume for convenience that the average effect among levels of factor i is zero which implies one parameterization,

where column j of X_{di} corresponding to level j , $j = 1, \dots, s_i - 1$, has entries 1 if factor i occurs at level j , -1 if factor i occurs at level s_i , and 0 otherwise. The actual parameters in model (1) are typically estimated using the method of **least squares**.

We note that the model given in equation (1) only allows a main-effects analysis for a given set of observations. However, additional columns can be added to model (1) corresponding to interaction parameters (see **Interactions**) by taking component-wise (Hadamard) products between the appropriate columns in model (1) involved in the interaction. We shall say that a design d is saturated if $\sum_{i=1}^m (s_i - 1) + 1 = N$. For saturated designs, only main-effect parameters can be estimated. In many screening experiments, the designs used are either saturated or nearly saturated and leave few **degrees of freedom** for estimating error. In these situations, analysis of the data to find significant main effects in model (1) usually depends upon the form of the matrices X_{di} in model (1). If the columns of all the X_{di} are orthogonal to one another and the corresponding parameter estimates can be assumed to be normal random variables with equal variances, then graphical methods such as half-normal plots (see **Half-Normal Plot**) can be used to identify significant effects. However, because the use of graphical methods is somewhat subjective in nature, other more objective statistical methods have been developed, for example, see Hamada and Balakrishnan [1] for a good discussion of methods available such as Lenth's method (see **Lenth's Method for the Analysis of Unreplicated Experiments**). When the columns of the X_{di} in model (1) are nonorthogonal or are orthogonal but the corresponding parameter estimates do not satisfy the assumptions appropriate for the use of graphical or other available methods, regression analysis and variable selection techniques can be used to identify significant effects.

Basic Principles for Screening Experiments

With regard to the usage of factorial designs as screening experiments, there are several principles which are often of use in both the choice of an experimental design and in the analysis of data obtained from a given experiment. These are the principles of hierarchical ordering, effect sparsity, and effect heredity. The general validity of these

principles has been established over time and a great deal of empirical observation.

The principle of hierarchical ordering says that lower-order interactions are more likely to be important than higher-order interactions, and it implies that main effects are more likely to be important than two-factor interactions, which are more likely to be important than three-factor interactions, and so on. This principle often leads to the additional assumption used in practice that in models such as equation (1) which might include interactions, three-factor and higher-order interactions are negligible. This latter assumption concerning three-factor and higher-order interactions is further justified through empirical evidence and the fact that three-factor and higher-order interactions are difficult to interpret.

The principle of effect sparsity says that the number of influential effects in a factorial experiment is likely to be small and has been observed to be the case in many factorial experiments. This principle is one reason why factorial screening experiments often prove successful in the identification of significant factors on a response.

The last principle, that of effect heredity, says that in order for an interaction to be significant, at least one of the factors involved in the interaction should be significant. This latter principle is of use in the determination of models which make physical sense, since in most experiments it is difficult to envision a significant interaction without at least one of the involved factors being significant. It is also of use in the development of follow-up design strategies.

Two-Level Orthogonal Designs

In this section, we consider two-level designs which allow for the orthogonal **estimation** of all main effects as these are the types of designs which are most often used as screening experiments. For these designs, the matrices X_{di} , $i = 1, \dots, m$ in model (1) consist of a single column with entries -1 or 1 depending on whether factor i occurs at a low or a high level in a given experimental run. The actual levels of a quantitative factor usually occur at different values but are coded so as to assume values -1 or $+1$ in equation (1). For qualitative factors, the same coding is used, with one of the factor levels being declared to be the "high" level.

Complete 2^m Designs

A complete 2^m design d is a design containing runs at all possible 2^m treatments that can be constructed from all possible level combinations of the m factors. The statistical analysis of such designs under a model such as equation (1) (with interactions included) typically involves the use of half-normal plots or standard analysis of variance (see **Analysis of Variance**) techniques.

Advantages

The primary advantage of complete 2^m designs is that they allow for the orthogonal estimation of all main effects and interactions between experimental factors (when such terms are included in the model). This lets an experimenter gain a greater understanding of how the response variable is related to the experimental factors being studied.

Disadvantages

Because 2^m grows exponentially as m increases, the number of runs required for a complete 2^m design becomes large quickly. For this reason, complete 2^m designs are not used much as screening designs unless m is relatively small ($m \leq 5$).

Regular 2^{m-k} Fractional Factorials (FF)

To handle experimental situations where m is relatively large and complete 2^m designs are not practical, designs called *regular 2^{m-k} fractional factorial (FF) designs* (see **Fractional Factorial Designs**) are often used. Regular 2^{m-k} FF designs are constructed using defining relations groups. In constructing a 2^{m-k} FF design d , the m factors consist of two types. There are $m - k$ “basic” factors, assumed to be those factors labeled $1, \dots, m - k$, such that the 2^{m-k} design is a complete factorial when considered just in terms of the basic factors. The other factors, labeled $m - k + 1, \dots, m$, are called *added factors* and are obtained by associating each added factor with some interaction in the basic factors. Each matrix X_{di} , $i = m - k + 1, \dots, m$, corresponding to an added factor and consisting of a single column in model (1) is obtained by taking the component-wise product of the single column matrices corresponding to the basic factors in its associated interaction. By combining each added factor label with those of the basic factors in their associated interaction, we obtain

the treatment defining effect words of the design. All possible products among the treatment defining words (according to the rule, if a factor label appears an even number of times in the product, it is deleted, whereas if it appears an odd number of times, it is kept) form an algebraic group called the *treatment defining relations group* which we denote by $G_t(d)$. We note that in any regular 2^{m-k} FF design d , the columns corresponding to any two distinct factorial effects (when included in model (1)) have euclidean inner product $\pm N$ or 0. When the columns corresponding to two factorial effects have inner product of $\pm N$, we say those effects are totally aliased (see **Aliasing in Fractional Designs**) or confounded. To find the **aliasing** set of a given factorial effect, that is, those factor effects that are totally confounded with it in d , simply multiply all words in $G_t(d)$ by the factorial effect using the multiplication rule for $G_t(d)$ given above.

Example 1: To study $m = 5$ factors in eight runs, one would use a regular 2^{5-2} design d . In such a design, factors 1, 2, and 3 are the basic factors and the levels of the two added factors 4 and 5 can be defined in terms of the basic two-factor interactions by $4 = 12$ and $5 = 13$. In d , the defining effect words are 124 and 135 and $G_t(d) = \{I, 124, 135, 124 \cdot 135 = 2345\}$, where I is the identity element of the group. The aliasing set for an interaction such as 134 is obtained by multiplying all words in $G_t(d)$ by 134 to get $\{134, 23, 45, 125\}$.

For more information on construction and a listing of good regular 2^{m-k} FF designs, the reader is referred to Chen *et al.* [2].

With regard to the usage of regular 2^{m-k} FF designs as screening experiments, they have certain advantages and disadvantages that should be kept in mind.

Advantages over 2^m designs

1. Economy of run size. In a 2^{m-k} regular FF design, there are $r = m - k$ basic factors and $N = 2^{m-k}$ runs. Such designs can accommodate up to $N - 1$ factors including the r basic factors and $2^r - r - 1$ added factors, one added factor for each interaction in the basic factors.
2. The analysis of 2^{m-k} designs is simple and well understood. All main effects can be estimated orthogonally and hence efficiently. The actual identification of significant effects in a

4 Screening Designs

given regular 2^{m-k} FF design is often accomplished through the use of half-normal plots or through the use of one of the more objective statistical methods discussed in Hamada and Balakrishnan [1].

3. The aliasing structure for such designs is simple, that is, factor effects are either totally confounded or orthogonal. This makes it easier to construct follow-up designs when aliased effects need to be broken out.
4. Many 2^{m-k} designs have desirable projection properties (see **Projectivity in Experimental Designs**). In particular, if a 2^{m-k} design has resolution R (see **Factorial Designs, Resolution of**), that is, the shortest word in $G_I(d)$ has exactly R letters in it, then any projection of the design onto $R - 1$ factors is a full factorial in the $R - 1$ factors. This fact can be used when analyzing the data from a 2^{m-k} design. For example, if only $R - 1$ of the factors from the original 2^{m-k} resolution R design are found to be significant, then all the interactions among the $R - 1$ factors can also be estimated without taking any additional observations because the projected design onto the $R - 1$ factors is a complete factorial.

Disadvantages

1. The number of runs for a 2^{m-k} design must be a power of 2. This limits their flexibility and availability as a viable option in many experimental settings.
2. The fact that effect estimates are either completely confounded or orthogonal that was cited in 3 above as an advantage can also be a disadvantage. Complete confounding of main effects and low-order interactions can sometimes cause difficulty in the identification of truly significant factor effects and require additional observations to dealias confounded effects. For example, suppose the estimating contrast for factor i , which is aliased with interaction $i_1 i_2$, is found to be significant. If neither factors i_1 nor i_2 are significant, then one can use the effect heredity principle to possibly conclude that the significant effect was due to the main effect of i . On the other hand, if one or both of factors i_1 or i_2 are significant, then it may be difficult to conclude whether the effect was due to the main effect of i or the interaction

$i_1 i_2$. In the latter case, additional observations may be needed to resolve the issue.

Comment Complete s^m and regular s^{m-k} , $s > 2$ FF designs can be constructed using defining relations along the lines given in this section for 2^m and regular 2^{m-k} designs. However, because regular s^{m-k} FF designs are not used as often for screening experiments, we shall not consider them here.

Nonregular Orthogonal FF Designs

The regular 2^{m-k} designs discussed in the section titled “Regular 2^{m-k} Fractional Factorials (FF)” are constructed using defining relations among factors. In these designs, any two factorial effects are either completely confounded or estimated orthogonally with all main effects being orthogonally estimable. FF designs that allow for the orthogonal estimation of all main effects but have some factorial effects that are neither orthogonal nor fully aliased are called *nonregular orthogonal FF designs*. In this section, we discuss a collection of nonregular orthogonal two-level FF designs which are sometimes useful in screening experiments.

Plackett and Burman (PB) [3] introduced a series of two-level orthogonal FF designs, where N is a multiple of 4 and $1 \leq m \leq N - 1$ (see **Plackett–Burman Designs**). If N is a power of 2, these designs contain the 2^m and the regular 2^{m-k} designs given earlier as special cases. A catalog of PB designs for $N = 4$ to $N = 100$ can be found online at www.research.att.com/~njas/oadir/index.html.

The constructions of PB designs are not based on well-defined defining relations groups, hence they have what is often referred to as a *complex aliasing structure*. Two-factorial effects are partially aliased in an N -run two-level design if the euclidean product between their corresponding columns in model (1) has absolute value between 0 and N . PB designs involve a great deal of partial aliasing between main effects and two-factor interactions as well as between pairs of two-factor interactions. For example, in the 12-run PB design involving 11 factors, every main effect is partially aliased with every two-factor interaction not involving itself. Thus every main effect in the 12-run PB design is aliased with 45 two-factor interactions. The complex alias structure associated with PB designs needs to be carefully

considered when such designs are used for screening as it can serve as both an advantage and disadvantage.

Advantages over regular 2^{m-k} FF designs

1. Economy of run size. For example, if between 8 and 11 two-level factors are to be included in a screening experiment, the smallest 2^{m-k} design that can be used contains 16 runs. On the other hand, a 12-run PB design can be used to accommodate between 8 and 11 factors and requires four fewer runs. Similar situations occur for number of factors between 16 and 27.
2. Estimates for all main effects are orthogonal and will be unbiased as long as all two-factor and higher-order interactions are negligible.
3. PB designs have a useful projection property which occurs as a result of their complex aliasing structures. For example, consider a 12-run PB design. It has been shown that any projection onto the columns associated with four or fewer factors has the property that all main effects and two-factor interactions among the four factors are estimable when higher-order interactions are negligible. Additional results for other numbers of factors and other PB designs have also been established [4]. Even though the estimates for main effects and interactions yielded in the associated projection models are not necessarily orthogonal, the projection models do allow for additional analysis of main effects and interactions without requiring additional observations.

Disadvantages

1. An assumption typically made when PB designs are used is that two-factor and higher-order interactions are negligible. When this assumption is valid, PB designs can be of great use as screening experiments. Even when this assumption is not completely valid but only a few two-factor interactions are nonnegligible, methods of analysis such as those described in the next section can successfully identify influential effects. However, when the assumption is violated and more than a few two-factor interactions are nonnegligible, problems of interpretation can arise. For example, because the complex aliasing structure of PB designs involves main effects being partially aliased with so many two-factor interactions, whenever a main effect contrast is determined to

be significant, it may be difficult to tell if the significant contrast is due to the associated factor or some combination of the partially aliased two-factor interactions. In such situations, complex aliasing can lead to the identification of experimental factors that are not truly influential as being influential and completely miss other influential factors. While the first situation is not as serious, failure to identify influential factors in a screening process can have serious consequences. Because of the ambiguity associated with main-effects estimates in PB designs, they are typically used only for screening.

Analysis of Designs with Complex Aliasing

In the section titled “Nonregular Orthogonal FF Designs”, screening designs were introduced that had many main effects partially aliased with two-factor interactions. When this partial aliasing is severe, the identification of influential main effects can be difficult. However, Hamada and Wu [5] suggested two analysis strategies for designs with complex aliasing, which can be successful when the principles of effect sparsity and effect heredity hold, for example, see the section titled “Notation and Model”, and the correlations between partially aliased effects are small to moderate. Their method 1 proceeds as follows:

- Step 1. For each factor, X_i , consider the variable X_i and all its two-factor interactions $X_i X_j$ with other factors. Use a stepwise regression procedure to identify significant effects from the candidate variables and denote the selected model by M_{X_i} . Repeat this for each of the factors and then choose the best model. Go to step 2.
- Step 2. Use a stepwise regression procedure to identify significant effects among the effects identified in the previous step as well as all the main effects. Go to step 3.
- Step 3. Using effect heredity, consider all variables consisting of the effects identified in step 2 and (ii) the two-factor interactions that have at least one member appearing among the main effects in (i). Also, consider (iii) interactions suggested by the experimenter. Use stepwise regression to identify significant effects among the effects in (i)–(iii). Go to step 4.

Step 4. Iterate between steps 2 and 3 until the model stops changing.

Advantages of method 1

1. The method of analysis suggested above provides a systematic method for analyzing data from designs with complex aliasing. When the conditions mentioned in the first paragraph hold, simulation studies and empirical evidence suggest that the method can be used successfully.

Disadvantages of method 1

1. When the method is used and the conditions cited do not hold, then the method can miss influential effects.
2. If more than a few two-factor interactions are partially aliased with main effects, the method may identify some incompatible models, that is, it may identify several models that explain the data equally well but contain different variables.

Method 2

Assume effect sparsity holds and that no meaningful model can have more than h effects. Search over all models that have no more than h effects (plus an intercept term) and satisfy the effect heredity principle. Choose the best model according to a sensible model selection criterion.

Advantages of method 2

1. It is computationally simple and easy to understand.
2. It provides a search over a greater number of models than does method 1.

Disadvantages of method 2

1. If h is large, the implementation of the method may not be computationally feasible.
2. The disadvantages cited for method 1 also hold for method 2.

Nonorthogonal Two-Level FF Designs

All of the designs given in the section titled “Two-Level Orthogonal Designs” allow for the orthogonal estimation of all main effects. However, other FF designs are available which can be used in screening experiments but where main effects are not all

orthogonal. In this section, we introduce three such sets of designs.

Model Robust FF (MRFF) Designs

Li and Nachtsheim [6] introduced a series of nonorthogonal two-level model robust fractional factorial (MRFF) designs having $N = 8, 12$, and 16 runs and m factors with m varying between 4 and 10. These designs are model robust in the sense that each design given allows for the estimation of all main effects and any combination or a large proportion of the combinations of up to g (or fewer) two-factor interactions, where g is specified by the user. These designs were computer generated and the reader is referred to Li and Nachtsheim [6] for their actual structure. The analysis of these designs and identification of significant effects is usually done through the use of regression and variable selection methods.

Advantages

1. The primary advantage of the MRFF designs described in this section is that they can be used to investigate both main effects and a small number of two-factor interactions. This provides an experimenter the opportunity to investigate model parameters in addition to main effects without collecting additional data in a follow-up experiment.

Disadvantages

1. The primary disadvantage of MRFF designs is that main effects are partially aliased with each other and with two-factor interactions, hence these designs have a complex aliasing structure. All problems associated with designs having complex aliasing thus hold for these designs also.

Resolution IV Nonorthogonal Two-Level FF Designs

In this section, we consider two-level designs which are similar in some respects to the resolution IV regular 2^{m-k} FF designs considered in the section titled “Regular 2^{m-k} Fractional Factorials (FF)”. An arbitrary two-level design d is said to have resolution III if it allows for the estimation of all main effects assuming that all interactions involving two or more factors are negligible, and it is said to have resolution IV if it allows for the estimation of all

main effects assuming all interactions involving three or more factors are negligible. One simple method of constructing resolution IV two-level designs is the “foldover method” (see **Foldover Designs**). In this section, the term *foldover design* is used to indicate a design $\tilde{d}$ which is obtained from an original design d by combining the runs in d with those mirror image runs that can be obtained by reversing the signs of all the factors in d . Thus if the original design d has m factors and N runs, the foldover design $\tilde{d}$ obtained from d has m factors and $2N$ runs. Margolin [7] showed that if d is a two-level FF design having resolution III, then the foldover design $\tilde{d}$ derived from d has resolution IV. Margolin discussed folding over minimal and efficient, nonorthogonal, two-level, resolution III designs to create resolution IV designs having m factors and $2m$ runs. These designs have been further discussed by Miller and Sitter [8, 9] who also recommended a method of analysis based on the techniques of Hamada and Wu [5]. The reader is referred to Miller and Sitter [8, 9] for details.

Advantages

1. The estimates for main effects in a resolution IV design are not aliased with any two-factor interactions.
2. The Margolin nonorthogonal resolution IV two-level designs having m factors and $2m$ runs require fewer runs than their regular resolution IV 2^{m-k} counterparts.
3. A limited number of two-factor interactions can typically be estimated in some designs, though there is usually a great deal of partial aliasing among the two-factor interactions.

Disadvantages

1. Because of the nonorthogonality involved in many of the initial resolution III designs, main effects in the foldover design are not orthogonal. Hence there is some loss in estimation efficiency of main effects (though in many cases it is small).
2. The method of analysis is not as straightforward as orthogonal designs, but Miller and Sitter [8, 9] have suggested ways to analyze these foldover designs.

Supersaturated Designs

Supersaturated designs (see **Supersaturated Designs**) are most often used as screening experiments

to study factors occurring at two levels where the number of factors m being studied is at least as large as the number of runs N in the experiment. Because $N \leq m$ in a supersaturated design d , any such design must necessarily have main effects which are nonorthogonal to one another and have a complex aliasing structure. A number of methods for constructing supersaturated designs have been suggested, for example, see Lin [10, 11], Wu [12], and Nguyen [13, 14], and while these designs have been used successfully as screening experiments, they should be used with caution. The actual analysis of supersaturated designs and identification of significant main effects usually involves the use of some sort of regression variable selection technique.

Advantages

1. Run size economy. Supersaturated designs can be used to study a large number of factors using relatively few observations.

Disadvantages

1. The analysis of supersaturated designs is not as straightforward as that of orthogonal designs since parameter main effect estimates between factors are not necessarily orthogonal. The lack of orthogonality also leads to some inefficiency in terms of estimating main effects.
2. The aliasing structure of supersaturated designs is generally more complex than that of orthogonal and other nonorthogonal designs because $N < m$. Hence the same problems cited earlier for designs with complex aliasing structures hold here but to an even greater extent.

Mixed-Level Designs

While two-level designs are predominantly used as screening experiments, there are some screening situations where some factors must have more than two levels or where different factors might require different number of levels. For example, if a quantitative factor is believed to have a quadratic effect on the response, then it would be necessary to include at least three levels for the factor to detect the quadratic effect. In other situations, a qualitative factor might have more than two levels or an experiment might require the use of both qualitative and quantitative factors requiring different number of levels. To handle such experimental situations, there are a number

of designs available such as regular s^{m-k} FF designs, $s > 2$ or mixed-level designs, that is, designs containing factors with differing number of levels. However, because such designs are not used as often in screening experiments as two-level designs, we do not consider them for further discussion here.

References

- [1] Hamada, M. & Balakrishnan, N. (1998). Analyzing unreplicated factorial experiments: a review with some new proposals (with discussion), *Statistica Sinica* **8**, 1–41.
- [2] Chen, J., Sun, D.X. & Wu, C.F.J. (1993). A catalogue of two-level and three-level fractional factorial designs with small runs, *International Statistical Review* **61**, 131–145.
- [3] Plackett, R.L. & Burman, J.P. (1946). The design of optimum multifactorial experiments, *Biometrika* **33**, 305–325.
- [4] Lin, D.K.G. & Draper, N.R. (1992). Projection properties of Plackett and Burman designs, *Technometrics* **34**, 423–428.
- [5] Hamada, M. & Wu, G.F.J. (1992). Analysis of experiments with complex aliasing, *Journal of Quality and Technology* **24**, 130–137.
- [6] Li, W. & Nachtsheim, C.J. (2000). Model robust factorial designs, *Technometrics* **42**, 345–352.
- [7] Margolin, B.H. (1969). Results on factorial designs of resolution IV for 2^n and $2^n 3^m$ series, *Technometrics* **11**, 431–444.
- [8] Miller, A. & Sitter, R.R. (2001). Using the folded-over 12 run Plackett-Burman design to consider interactions, *Technometrics* **43**, 44–55.
- [9] Miller, A. & Sitter, R.R. (2005). Using the folded-over non-orthogonal designs, *Technometrics* **47**, 502–513.
- [10] Lin, D.K.G. (1993). A new class of supersaturated designs, *Technometrics* **35**, 28–31.
- [11] Lin, D.K.G. (1995). Generating systematics supersaturated designs, *Technometrics* **37**, 213–225.
- [12] Wu, C.F.J. (1993). Construction of supersaturated designs through partially aliased interactions, *Biometrika* **80**, 661–669.
- [13] Nguyen, N.K. (1996a). A note on the construction of near orthogonal arrays and mixed levels and economic run size, *Technometrics* **38**, 279–283.
- [14] Nguyen, N.K. (1996b). An algorithmic approach to constructing supersaturated designs, *Technometrics* **38**, 69–73.

Further Reading

- Daniel, C. (1959). Use of half-normal plots in interpreting factorial experiments, *Technometrics* **1**, 311–341.
- Wang, J.C. & Wu, C.F.J. (1995). A hidden projection property of Plackett-Burman and related designs, *Statistica Sinica* **5**, 235–250.

Related Articles

Aliasing in Fractional Designs; Factorial Experiments; Fractional Factorial Designs; Half-Normal Plot; Interactions; Lenth's Method for the Analysis of Unreplicated Experiments; Main Effects; Main Effect Designs; Orthogonal Arrays; Plackett-Burman Designs; Projectivity in Experimental Designs; Supersaturated Designs; Factorial Designs, Resolution of.

MIKE JACROUX

Split-Plot Designs

Classical Agricultural Split-Plot Designs

The terminology used in the context of split-plot designs derives from its initial agricultural applications, where experiments were carried out on different plots of land. In these experiments, levels of one or more factors were assigned to large plots of land referred to as *whole plots*, while levels of other factors were assigned to smaller plots of land referred to as *split-plots* or *subplots*. This means that split-plot designs have two types of experimental units, whole plots and subplots. The smaller experimental units, the subplots, are nested within the larger ones, the whole plots. The importance of the split-plot structure for the design and analysis of experiments was first recognized by Fisher [1].

For example, in experiments for investigating the effect of different fertilizers and varieties on the yield of crops, the fertilizers are often sprayed from planes and large plots of land must then be treated with the same fertilizer. Crop varieties can be planted on smaller plots of land that are obtained by dividing up the large plots that were sprayed together. The larger plots of land are called *whole plots*, whereas the smaller plots are named *split-plots* or *subplots*. Because the levels of the factor fertilizer are applied to whole plots, it is called the *whole-plot factor* of the experiment. The second factor, crop variety, is applied to the subplots, and is referred to as the *subplot factor* of the experiment. A schematic representation of a split-plot design involving fertilizer as the whole-plot factor and variety as the subplot factor is given in Figure 1. In the figure, the levels of the whole-plot factor are represented by $F_1 - F_3$, and those of the subplot factor by $V_1 - V_3$. Evidently, split-plot designs with more than one whole-plot factor and more than one subplot factor can be constructed as well. In classical split-plot designs, all combinations of levels of the subplot factors occur within each combination of levels of the whole-plot factors.

The layout of the design in Figure 1 resembles that of a randomized block design (see **Blocking; Block Designs, Incomplete**). In a randomized block design however, the factor level combinations can be randomly assigned to the experimental units in the

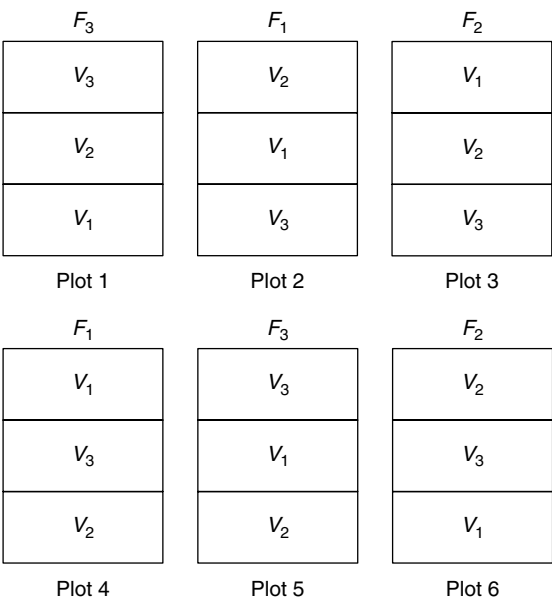


Figure 1 Classical agricultural split-plot design where $F_1 - F_3$ and $V_1 - V_3$ represent the levels of the factors fertilizer and variety, respectively. The design has six whole plots (labeled Plot 1–6) each of which is divided into three subplots

blocks. In the split-plot design shown in Figure 1, such a complete random assignment of the factor level combinations to the subplots is impossible because all the subplots within one whole plot must be treated with the same fertilizer. Therefore, only a *restricted randomization* can be used. Typically, this restricted randomization is performed in two stages. First, the combinations of levels of the whole-plot factors are randomly assigned to the whole plots. Next, the combinations of levels of the subplot factor are randomly assigned to the subplots within each whole plot.

Classical Analysis of Data from Split-Plot Experiments

In the analysis of data from split-plot designs, the whole plots act as levels of a blocking factor (see **Blocking; Block Designs, Incomplete**). Therefore, *whole-plot effects* are included in the statistical model along with the whole-plot and subplot factors and advantages similar to the ones for block designs can

2 Split-Plot Designs

be obtained using split-plot designs. For example, the effect of the varieties in Figure 1 is calculated from the differences between the yields within each whole plot. As these differences eliminate the whole-plot effect, the effect of the varieties is measured with greater precision than the effect of the fertilizers. The calculation of the latter effect is not as precise because it involves calculating and comparing averages over all the subplots within a whole plot, and errors between subplots in different whole plots are usually larger than errors between subplots in a single whole plot. The effect of the whole-plot factor fertilizer is therefore tested against the *whole-plot error*, whereas the effect of the subplot factor variety is tested against the *subplot error*.

Traditionally, the whole-plot and subplot factors in a classical split-plot design are treated as categorical. The whole-plot effects are modeled using random effects (unlike block effects, which are typically modeled using fixed effects). The analysis of variance (ANOVA) model used to analyze the split-plot data can be written as

$$Y_{ijk} = \mu_{...} + \gamma_{i(j)} + \alpha_j + \beta_k + (\alpha\beta)_{jk} + \varepsilon_{ijk},$$

$$i = 1, \dots, m; j = 1, \dots, a; k = 1, \dots, b \quad (1)$$

with Y_{ijk} the response for the j th fertilizer and the k th variety in whole plot i , $\mu_{...}$ the overall mean, $\gamma_{i(j)}$ the random effect of the i th whole plot nested within the j th level of the whole-plot factor fertilizer, α_j the **main effect** of the j th fertilizer, β_k the main effect of the k th variety, $(\alpha\beta)_{jk}$ the interaction effect of the j th fertilizer and the k th variety, ε_{ijk} the random error, a the number of levels of the whole-plot factor, b the number of levels of the subplot factor, and m the number of whole plots at each whole-plot factor level. The constants α_j , β_k , and $(\alpha\beta)_{jk}$ are subject to the constraints $\sum_{j=1}^a \alpha_j = 0$, $\sum_{k=1}^b \beta_k = 0$, $\sum_{j=1}^a (\alpha\beta)_{jk} = 0$ for all k and $\sum_{k=1}^b (\alpha\beta)_{jk} = 0$ for all j to make the model estimable. The α_j , β_k , and $(\alpha\beta)_{jk}$ are the whole-plot effects, the subplot effects, and the whole-plot-by-subplot interaction effects, respectively. The random effects $\gamma_{i(j)}$ and ε_{ijk} are often referred to as *whole-plot error* and *subplot error*, respectively. The whole-plot errors $\gamma_{i(j)}$ are assumed to be independently normally distributed with zero mean and variance σ_γ^2 (the *whole-plot error variance*), and the subplot

errors ε_{ijk} are assumed to be independently normally distributed with zero mean and variance σ_ε^2 (the *subplot error variance*). It is also assumed that $\gamma_{i(j)}$ and ε_{ijk} are independent. These assumptions are identical to those made for randomized block designs with random block effects and imply that observations from different whole plots are independent, whereas observations within a whole plot are correlated. The resulting **correlation** structure of the responses is called *compound symmetric*. This structure is special in the sense that, while two observations in a given whole plot are correlated in advance of the experiment, they are independent once a whole plot has been selected.

A general ANOVA table for a classical two-factor split-plot design is given in Table 1. In the table, the whole-plot and subplot factors are denoted by A and B , respectively, and a , b , and m denote the number of levels of factor A , the number of levels of factor B , and the number of whole plots at each whole-plot factor level, respectively. The number of observations in the split-plot design equals $N = abm$. For the design in Figure 1, a , b , m , and N equal 3, 3, 2, and 18, respectively. The total number of whole plots in a classical split-plot design equals am and is denoted by c . For the example, that number is six.

An important aspect of split-plot designs is that the effects of the whole-plot factors are estimated less precisely than those of the subplot factors. This is taken into account in the ANOVA analysis. It implies that split-plot designs result in a low **power** for detecting whole-plot effects and a high power for detecting subplot effects. Because interactions between a whole-plot factor and a subplot factor are estimated precisely too, split-plot designs have a high power for detecting these as well. Ignoring the split-plot nature of the experiment and analyzing the data from a split-plot design as if they were data from a completely randomized experiment can lead to erroneous conclusions. Such an analysis typically underestimates the variance of the estimates of the whole-plot effects and overestimates the variance of the estimates of the subplot effects and the whole-plot-by-subplot interaction effects. Because of this, ignoring the split-plot nature of the design results in labeling the whole-plot effects as significant too often and in labeling the subplot effects and the whole-plot-by-subplot interaction effects as significant not often enough.

Table 1 ANOVA table for a classical two-factor split-plot design

Source of variation	SS	df	MS	E(MS)
<i>Whole plots</i>				
Factor A	$SSA = bm \sum_{j=1}^a (\bar{Y}_{.j} - \bar{Y}_{...})^2$	$a - 1$	MSA	$\sigma^2 + b\sigma_b^2 + bm \frac{\sum_{j=1}^a \alpha_j^2}{a - 1}$
Whole-plot error	$SSW(A) = b \sum_{i=1}^m \sum_{j=1}^a (\bar{Y}_{ij.} - \bar{Y}_{.j})^2$	$a(m - 1)$	MSW(A)	$\sigma^2 + b\sigma_b^2$
<i>Subplots</i>				
Factor B	$SSB = am \sum_{k=1}^b (\bar{Y}_{..k} - \bar{Y}_{...})^2$	$b - 1$	MSB	$\sigma^2 + am \frac{\sum_{k=1}^b \beta_k^2}{b - 1}$
AB interaction	$SSAB = m \sum_{j=1}^a \sum_{k=1}^b (Y_{.jk} - \bar{Y}_{.j} - \bar{Y}_{..k} + \bar{Y}_{...})^2$	$(a - 1)(b - 1)$	MSAB	$\sigma^2 + m \frac{\sum_{j=1}^a \sum_{k=1}^b (\alpha\beta)_{jk}}{(a - 1)(b - 1)}$
Subplot error	$SSB.W(A) = \sum_{i=1}^m \sum_{j=1}^a \sum_{k=1}^b (Y_{ijk} - \bar{Y}_{ij.} - \bar{Y}_{.jk} + \bar{Y}_{.j})^2$	$a(m - 1)(b - 1)$	MSB.W(A)	σ^2
Total	$SSTO = \sum_{i=1}^m \sum_{j=1}^a \sum_{k=1}^b (Y_{ijk} - \bar{Y}_{...})^2$	$abm - 1$		

Industrial Split-Plot Experiments

Classical Industrial Split-Plot Experiments

Classical split-plot designs are not only applied in agriculture, but also in industrial and other settings. Although a better name for split-plot designs would therefore be *split-unit* designs, the original name of the design is still used in recent literature. An example of an industrial split-plot experiment is given in Box *et al.* [2]. The experiment was intended to investigate the corrosion resistance of steel bars treated with four different coatings C_1 , C_2 , C_3 , and C_4 at three duplicated furnace temperatures T_1 , T_2 , and T_3 (360, 370, and 380 °C, respectively). For each temperature, four differently coated bars were randomly arranged in the furnace. A schematic representation of this industrial split-plot design is shown in Figure 2, where the levels of the whole-plot factor, the furnace temperature, are represented by T_1 – T_3 , and those of the subplot factor, the coating, by C_1 – C_4 . In this experiment, the six runs of the furnace serve as the whole plots of the experiment and the four positions in the furnace can be interpreted as subplots. The furnace temperature and the coating are the whole-plot and subplot factors, respectively. When constructing the design, the temperatures are assigned to the runs of the furnace first. Next, the coatings are randomly assigned to the positions in the furnace. The split-plot design involving six operations of the furnace is more convenient and cheaper to conduct than a completely randomized design which would require 24 heatings each of which would be used for testing one single coating only.

The data for the corrosion resistance study depicted in Figure 2 are given in Table 2. The

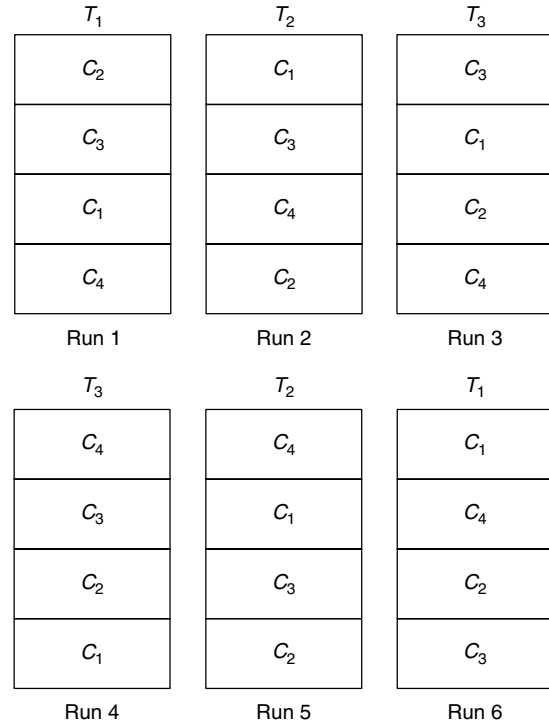


Figure 2 Classical industrial split-plot design where T_1 – T_3 and C_1 – C_4 represent the levels of the factors temperature and coating, respectively. The design has six whole plots each of which contains four observations and corresponds to a run of the furnace

corresponding ANOVA table is displayed in Table 3. In the ANOVA table, the effect of the whole-plot factor temperature is tested against the whole-plot error. The effect of the subplot factor coating is tested against the subplot error, just like the interaction

Table 2 Data for the corrosion resistance study. The whole plots of the study are indicated in the column labeled WP. The columns labeled T , C , and y contain the furnace temperatures, coatings, and observed responses, respectively

WP	T	C	y	WP	T	C	y	WP	T	C	y
1	360	2	73	3	380	3	147	5	370	4	150
	360	3	83		380	1	155		370	1	140
	360	1	67		380	2	127		370	3	121
	360	4	89		380	4	212		370	2	142
2	370	1	65	4	380	4	153	6	360	1	33
	370	3	87		380	3	90		360	4	54
	370	4	86		380	2	100		360	2	8
	370	2	91		380	1	108		360	3	46

Table 3 ANOVA table for the corrosion resistance study

Source of variation	<i>SS</i>	<i>df</i>	<i>MS</i>	<i>F</i>	<i>p</i>
<i>Whole plots</i>					
Temperature	26 519.25000	2	13 259.62500	10.23	0.0026
Whole-plot error	14 439.62500	3	4813.20833		
<i>Subplots</i>					
Coating	4289.12500	3	1429.70833	11.48	0.0020
Temperature $\times$ coating	3269.75000	6	544.95833	4.38	0.0241
Subplot error	1120.87500	9	124.54167		
Total	49 638.62500	23			

effect between the two experimental factors. The main effects of the two factors are significant at the 1% level, whereas the interaction effect is significant at the 5% level. Note that the mean squared subplot error, which equals 124.54, is much smaller than the mean squared whole-plot error, which amounts to 4813.21. This shows that the errors between subplots in a single whole plot are much smaller than errors between subplots in different whole plots.

Other Industrial Split-Plot Experiments

As split-plot designs in industry often involve several quantitative whole-plot and subplot factors, and as, for reasons of cost and time, industrial experiments are frequently limited in size, classical split-plot designs are usually not satisfactory. As a matter of fact, classical split-plot designs in which all combinations of levels of one or more factors occur within levels of one or more other factors to allow for an ANOVA analysis (which involves the **estimation** of a large number of model parameters) are typically quite large. The split-plot designs used in industry and described in the literature on response-surface methodology (see **Response Surface Methodology**) therefore often look different from the classical ones. The main difference is that only a fraction of all combinations of subplot factor levels appear in every whole plot. In general, the analysis of data from these response-surface split-plot designs can be done using a regression approach rather than an ANOVA approach. For the analysis of two-level **factorial** split-plot designs and some specific two-level factorial split-plot designs however, a simpler ANOVA type of approach can still be used.

An example of a response-surface split-plot experiment in the food industry, where a regression approach is needed for analysis, was provided by

Trinca and Gilmour [3]. In this experiment, the impact of five factors on the yield of a protein extraction process was investigated: the feed position for the inflow of a mixture, the feed flow rate, the gas flow rate, and the concentrations of a protein A and a protein B. Twenty-one days were available for experimentation and three levels were used for each factor because the interest was in estimating a second-order response-surface model. The problem with this experiment was that setting the feed position could be done only once per day because it involved taking apart and reassembling the equipment. By keeping the position fixed during one day, two experimental runs could be performed per day instead of one, allowing 42 experimental runs in total. The split-plot design that was actually used is given in Table 4. In the table, the whole-plot factor, feed position of the inflow, is denoted by w and the subplot factors are denoted by s_1-s_4 . Notice that the level of the whole-plot factor does not change within each whole plot.

It is not always as obvious that experiments are run in a split-plot format as in the protein extraction experiment. In the following scenarios, a split-plot design is often used.

1. When different factor level combinations are tested while some other factor levels are not reset, then the resulting design is a split-plot design. This happens, for example, when experimenters try out different factor level combinations during one run of an oven. The oven temperature is not reset for all these observations. It is therefore the whole-plot factor of the experiment. The other factors are the subplot factors.
2. In two-stage experiments, it sometimes happens that some of the experimental factors are applied in the first stage, whereas others are applied in the

6 Split-Plot Designs

Table 4 Response-surface split-plot design for a protein extraction experiment with 21 whole plots containing two observations each

WP	w	s_1	s_2	s_3	s_4	WP	w	s_1	s_2	s_3	s_4	WP	w	s_1	s_2	s_3	s_4
1	-1	-1	-1	-1	-1	8	0	-1	-1	-1	1	15	1	-1	-1	-1	-1
	-1	1	-1	-1	1		0	1	-1	1	-1		1	1	0	0	0
2	-1	-1	1	-1	1	9	0	-1	-1	1	-1	16	1	-1	-1	1	1
	-1	0	0	0	0		0	1	-1	1	1		1	0	1	0	0
3	-1	-1	0	0	0	10	0	-1	-1	1	1	17	1	-1	1	1	-1
	-1	1	1	-1	-1		0	1	-1	-1	-1		1	0	0	-1	0
4	-1	-1	1	1	-1	11	0	-1	1	-1	-1	18	1	-1	1	-1	1
	-1	1	1	-1	1		0	1	1	1	1		1	1	1	-1	-1
5	-1	0	0	0	-1	12	0	-1	1	1	1	19	1	0	0	-1	0
	-1	0	1	0	0		0	1	1	1	-1		1	0	-1	0	0
6	-1	0	0	0	0	13	0	-1	0	0	0	20	1	1	-1	-1	1
	-1	1	-1	1	-1		0	0	0	0	1		1	0	0	0	-1
7	-1	0	-1	0	0	14	0	0	0	1	0	21	1	0	0	0	1
	-1	1	1	1	1		0	1	0	0	0		1	0	0	1	0

second stage. In those cases, batches of experimental material are often split in subbatches after the first stage. The subbatches then undergo different treatments dictated by the levels of the factors applied in the second stage of the experiment. The factors applied in the first stage are the whole-plot factors, whereas those applied in the second stage are the subplot factors. In experiments involving more than two stages, more complicated designs than split-plot designs are used. Examples of such designs for multistage experiments in the semiconductor industry can be found in Mee and Bates [4].

- Experiments for simultaneously investigating the impact of ingredients of a mixture and of process factors, i.e. mixture-process variable experiments, are frequently conducted as split-plot experiments [5]. This can happen in two ways. First, several batches can be prepared that undergo different process conditions. In that case, the ingredients of the mixture are the whole-plot factors and the factors determining the tested process conditions act as subplot factors. Second, processing conditions can be held fixed while consecutively trying out different mixtures. In that case, the ingredients of the mixture are the subplot factors of the design.
- When experimental designs need to be run sequentially, experimenters are often reluctant to change the levels of one or more factors because this is impractical, time consuming, or expensive. All the observations that are obtained by not resetting the levels of the *hard-to-change factors*

then play the role of a whole plot. The hard-to-change factors are the whole-plot factors of the design, whereas the remaining factors are subplot factors. A useful reference in this context is Webb *et al.* [6].

- A special case of an experiment involving hard-to-change factors is a *prototype experiment*, in which the levels of some of the factors define a prototype, whereas the levels of other factors determine operating conditions under which the prototype is tested. As changing the levels of a prototype factor involves assembling a new prototype, these levels are held constant for several observations under different operating conditions. In that type of experiment, the factors related to the prototype are the whole-plot factors, whereas the ones connected to the operating conditions serve as subplot factors. Prototype experiments run using split-plot designs are discussed in Bisgaard and Steinberg [7].
- As indicated in Box and Jones [8], robust product experiments (*see Robust Design*) involving crossed arrays are often run as split-plot designs. Often, this is because the noise factors are hard to control and changing their levels is cumbersome. An interesting feature of split-plot designs for robust product experimentation is that control-by-noise interaction effects can be estimated more precisely than by using a completely randomized experiment.

A common denominator of these industrial split-plot experiments is that the levels of some of

the factors are not independently reset for several observations. The groups of observations for which these levels were not reset constitute the whole plots of the experimental design. The number of observations in each of these groups is the number of subplots. In many of the scenarios described above, though not always, the number of subplots within each whole plot is dictated by practical considerations such as the number of observations that can be made in one day, using one batch or in one oven run. In those cases, the numbers of subplots within each whole plot are equal. However, the regression or response-surface analysis of data from split-plot designs can handle situations in which the numbers of subplots are unequal across whole plots. Goos [9] demonstrates it can be advantageous in terms of design efficiency to use unequal numbers of subplots within each whole plot.

Response-Surface Split-Plot Model and Analysis

Response-Surface Model

The data from industrial response-surface split-plot designs with c whole plots and $n_1, \dots, n_c$ subplots in each of them can in general be analyzed by means of a regression approach. If the levels of the k whole-plot factors and the l subplot factors in the design are denoted by $\mathbf{w} = [w_1, \dots, w_k]'$ and $\mathbf{s} = [s_1, \dots, s_l]'$ respectively, then the regression model can be written as

$$Y_{ij} = \mathbf{f}'(w_i, s_{ij})\boldsymbol{\beta} + \gamma_i + \varepsilon_{ij} \quad (2)$$

where Y_{ij} represents the response of the j th observation in the i th whole plot, $\mathbf{w}_i$ contains the levels of the k whole-plot factors in that whole plot, $\mathbf{s}_{ij}$ contains the levels of the l subplot factors at

the j th observation in the i th whole plot, $\boldsymbol{\beta}$ is the p -dimensional vector of unknown factor effects, and γ_i and ε_{ij} are the i th whole-plot error and the subplot error associated with the j th observation in the i th whole plot, respectively. Finally, $\mathbf{f}(\mathbf{w}_i, \mathbf{s}_{ij})$ contains the p -dimensional model expansion of the experimental factors. The model expansion typically contains the levels of the factors (corresponding to the factors' linear effects), their cross-products (corresponding to interactions) and, for second-order models, their squares (corresponding to the quadratic effects). In matrix notation, and denoting $n = \sum_{i=1}^c n_i$, the model can be written as

$$\mathbf{Y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{Z}\boldsymbol{\gamma} + \boldsymbol{\varepsilon} \quad (3)$$

where $\mathbf{Y} = [Y_{11}, \dots, Y_{cn_c}]'$, $\mathbf{X} = [\mathbf{f}(\mathbf{w}_1, \mathbf{s}_{11}), \dots, \mathbf{f}(\mathbf{w}_c, \mathbf{s}_{cn_c})]'$ represents the $n \times p$ model matrix, $\mathbf{Z}$ is an $n \times c$ matrix of zeros and ones assigning the n runs to the c whole plots, $\boldsymbol{\gamma}$ is the c -dimensional vector containing the random effects of the c whole plots, and $\boldsymbol{\varepsilon}$ is the n -dimensional vector containing the random errors. It is assumed that

$$E(\boldsymbol{\varepsilon}) = \mathbf{0}_n \text{ and } \text{cov}(\boldsymbol{\varepsilon}) = \sigma_\varepsilon^2 \mathbf{I}_n \quad (4)$$

$$E(\boldsymbol{\gamma}) = \mathbf{0}_c \text{ and } \text{cov}(\boldsymbol{\gamma}) = \sigma_\gamma^2 \mathbf{I}_c \quad (5)$$

and

$$\text{cov}(\boldsymbol{\gamma}, \boldsymbol{\varepsilon}) = \mathbf{0}_{c \times n} \quad (6)$$

Under these assumptions, the **covariance** matrix of the responses, $\text{var}(\mathbf{Y})$, is

$$\mathbf{V} = \sigma_\varepsilon^2 \mathbf{I}_n + \sigma_\gamma^2 \mathbf{Z}\mathbf{Z}' \quad (7)$$

which is a block diagonal matrix if the entries of $\mathbf{Y}$ are grouped per whole plot. For example, for a design with two whole plots of three observations and one whole plot of two observations, the covariance matrix equals

$$\mathbf{V} = \begin{bmatrix} \sigma_\varepsilon^2 + \sigma_\gamma^2 & \sigma_\gamma^2 & \sigma_\gamma^2 & 0 & 0 & 0 & 0 & 0 \\ \sigma_\gamma^2 & \sigma_\varepsilon^2 + \sigma_\gamma^2 & \sigma_\gamma^2 & 0 & 0 & 0 & 0 & 0 \\ \sigma_\gamma^2 & \sigma_\gamma^2 & \sigma_\varepsilon^2 + \sigma_\gamma^2 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & \sigma_\varepsilon^2 + \sigma_\gamma^2 & \sigma_\gamma^2 & \sigma_\gamma^2 & 0 & 0 \\ 0 & 0 & 0 & \sigma_\gamma^2 & \sigma_\varepsilon^2 + \sigma_\gamma^2 & \sigma_\gamma^2 & 0 & 0 \\ 0 & 0 & 0 & \sigma_\gamma^2 & \sigma_\gamma^2 & \sigma_\varepsilon^2 + \sigma_\gamma^2 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & \sigma_\varepsilon^2 + \sigma_\gamma^2 & \sigma_\gamma^2 \\ 0 & 0 & 0 & 0 & 0 & 0 & \sigma_\gamma^2 & \sigma_\varepsilon^2 + \sigma_\gamma^2 \end{bmatrix} \quad (8)$$

Model Estimation and Inference

Because of the correlation between the observations from a split-plot design, the best linear unbiased estimator of the model parameters contained within β is the generalized least-squares estimator

$$\hat{\beta} = (\mathbf{X}'\mathbf{V}^{-1}\mathbf{X})^{-1}\mathbf{X}'\mathbf{V}^{-1}\mathbf{y} \quad (9)$$

with covariance matrix

$$\text{var}(\hat{\beta}) = (\mathbf{X}'\mathbf{V}^{-1}\mathbf{X})^{-1} \quad (10)$$

To use these expressions, the whole-plot and subplot error variances σ_γ^2 and σ_ε^2 contained within $\mathbf{V}$ have to be replaced by estimates. In the response-surface split-plot design literature, two approaches are advocated to estimate the variances. The first approach, restricted (or residual) maximum-likelihood estimation, is generally applicable to any split-plot design and was recommended by Letsinger *et al.* [10]. The second approach uses pure error estimates for the whole-plot and subplot error variances and can be applied only if the design has replicated points. This approach offers the advantage that the estimates of the **variance components** σ_ε^2 and σ_γ^2 do not depend on the specification of the model, $\mathbf{f}'(\mathbf{w}_i, \mathbf{s}_{ij})\beta$. It is outlined in Vining *et al.* [11]. A disadvantage is that replications and thus a large experiment are required for estimating the variance components, which are not of primary interest to the experimenter.

The statistical inference for split-plot response-surface designs is not straightforward. First, it is not clear what denominator **degrees of freedom** have to be used for the hypothesis tests. Goos *et al.* [12] show that the approach suggested in Kenward and Roger [13] is best in terms of controlling the type I error if generalized least-squares estimation is combined with restricted maximum likelihood. Second, industrial split-plot designs sometimes have a small number of whole plots so that the whole-plot error variance is not estimable or cannot be estimated accurately. In such situations, the Bayesian approach of Gilmour and Goos [14] to analyze data from split-plot designs is recommended. Also, it is not clear what degrees of freedom to use when the pure error estimation approach of Vining *et al.* [11] for the variance components is used. The good performance of restricted maximum likelihood with the Kenward and Roger option for determining the denominator degrees of freedom for hypothesis tests, together

with the growing availability of software, makes this combination the current state-of-the-art method for analyzing data from split-plot designs.

Simplified Analysis for Two-Level Designs

For two-level factorial split-plot designs and for the two-level **fractional factorial** designs described in Huang *et al.* [15], Bingham and Sitter [16], and Bingham *et al.* [17, 18], a simplified analysis is possible because the generalized least-squares estimator produces the exact same estimates as the ordinary least-squares estimator. As a result, estimating the model from such a design can be done in the same fashion as for completely randomized two-level factorial and fractional factorial designs (*see* **Factorial Experiments; Fractional Factorial Designs**). An appropriate ANOVA then consists of two separate parts just like the classical analysis in Table 1. Bisgaard [19] however prefers using half-normal plots to identify significant effects just like for completely randomized two-level designs (*see* **Half-Normal Plot**). It is important however that separate half-normal plots are prepared for the whole-plot effects and for the subplot effects because they are estimated with a different precision. Alternatively, Lenth's method (*see* **Lenth's Method for the Analysis of Unreplicated Experiments**) can be applied separately to each set of effects.

Design of Response-Surface Split-Plot Experiments

Several construction methods for split-plot designs can be found in the literature. A distinction between methods for first- and second-order response-surface designs can be made.

First-Order Split-Plot Response-Surface Designs

The design of two-level factorial and fractional factorial split-plot experiments has been discussed by Huang *et al.* [15], Bingham and Sitter [16], and Bingham *et al.* [17, 18] (*see also* **Fractional Factorial Designs**). The design of fractional factorial split-plot experiments for robust parameter design (*see* **Product Array Designs**) received attention in Bingham and Sitter [20]. The design criterion used by these authors is an extension of the

minimum aberration criterion (see **Minimum Aberration**).

Standard Second-Order Split-Plot Designs

Vining *et al.* [11] suggest using modified central composite designs and Box–Behnken designs for conducting second-order split-plot experiments. The **central composite** and Box–Behnken designs are modified and arranged in whole plots so that the generalized least-squares estimator (which takes into account the split-plot error structure of the data) produces the same estimates as the ordinary least-squares estimator (which ignores the split-plot nature of the experiment). These so-called *equivalent-estimation designs* have a lot of observations at the **center point** to allow for a pure error estimate of the variance components. A drawback of the designs is that, despite their sizes, they lead to imprecise estimates of the model parameters. This is because a large part of the designs consists of replicated center points for the pure error estimation of the variance components. These replicates provide no information about the factor effects. Also, it is not clear how the standard designs can be adapted for different numbers and sizes of whole plots.

Tailor-Made Split-Plot Response-Surface Designs

Two approaches have been presented in the literature to construct tailor-made second-order response-surface designs. Trinca and Gilmour [3] suggest a sequential approach for constructing *multistratum designs*, a special case of which are split-plot designs. For split-plot designs, the sequential approach starts by selecting a design for the whole-plot factors. Next, a design for the subplot factors is chosen, and an algorithmic approach is used to assign the factor level combinations in that design to the whole plots. The algorithm searches for a split-plot design that is as orthogonal as possible. A second approach that can be used for constructing tailor-made split-plot designs is the optimal design approach (see also **Optimal Design**). This approach is advocated by Goos [21] and in general leads to precise parameter estimates and accurate predictions. In theory (and unlike the sequential approach of Trinca and Gilmour [3]), the optimal designs depend on the (unknown) ratio of the whole-plot error variance σ_γ^2

to the subplot error variance σ_ε^2 , which has to be specified by the experimenter. Fortunately however, specifying different ratios of the variance components within the range of practical values usually leads to the same design when the optimal design approach is used. The practical range of the ratio of the whole-plot error variance to the subplot error variance is 0.5–10. Algorithms for computing D-optimal split-plot designs have been proposed by Goos and Vandebroek [22] and Jones and Goos [23]. Goos and Donev [24] show how to adapt the optimal design approach to include replicated points in the designs and to allow lack-of-fit tests to be conducted.

The two approaches for constructing tailor-made split-plot designs are flexible in that they can be used for any number and size of whole plots. They can also handle constraints on the factor levels and situations involving different types of factors, i.e., they can be used for problems involving quantitative, categorical, and/or mixture factors.

Extensions and Related Designs

Experiments involving more than two stages will sometimes have factors that are applied at each of these stages. In such cases, the resulting design will often be a *split-split-plot design*. A similar design might be useful when there are several hard-to-change factors. In such cases, the factor that is hardest to change will be reset the smallest number of times, and factors that are easier to change will be reset more often (but not as often as easy-to-change factors). More complicated designs than split-plot and split-split-plot designs for multistage experiments are discussed in Mee and Bates [4].

Split-plot and split-split-plot type of designs are often used inadvertently when experimental designs are run sequentially. This is because it is not sufficient to run a design in a random order to have a completely randomized design with observations that are statistically independent. For this to be the case, it is important that the levels of all the factors are reset for every single observation even when the level of one or more factors is the same in two successive runs. If this is not done, then it is unrealistic to assume that the successive runs are independent from each other and an analysis that takes into account

the correlation pattern is recommended. Details about this type of experimental approach can be found in Webb *et al.* [6]. In case resetting the factors levels for every run is impossible for reasons of time or cost, it is however better to construct a proper split-plot design using one of the methods outlined above rather than to use a random run order and not reset the factors levels. This is because, for properly designed split-plot experiments, many of the parameter estimates are often independent from each other, while this is not true for *randomized-not-reset experiments* in general. Also, the parameter estimates from a designed split-plot experiment will usually be more precise.

References

- [1] Fisher, R.A. (1925). *Statistical Methods for Research Workers*, Oliver & Boyd, Edinburgh.
- [2] Box, G.E.P., Hunter, J.S. & Hunter, W.G. (2005). *Statistics for Experimenters: Design, Innovation, and Discovery*, Wiley-Interscience, New York.
- [3] Trinca, L.A. & Gilmour, S.G. (2001). Multi-stratum response surface designs, *Technometrics* **43**, 25–33.
- [4] Mee, R.W. & Bates, R.L. (1998). Split-plot designs: experiments for multistage batch processes, *Technometrics* **40**, 127–140.
- [5] Cornell, J.A. (1988). Analyzing data from mixture experiments containing process variables: a split-plot approach, *Journal of Quality Technology* **20**, 2–23.
- [6] Webb, D., Lucas, J.M. & Borkowski, J.J. (2004). Factorial experiments when factor levels are not necessarily reset, *Journal of Quality Technology* **36**, 1–11.
- [7] Bisgaard, S. & Steinberg, D.M. (1997). The design and analysis of $2^{k-p} \times s$ prototype experiments, *Technometrics* **39**, 52–62.
- [8] Box, G.E.P. & Jones, S.P. (1992). Split-plot designs for robust product experimentation, *Journal of Applied Statistics* **19**, 3–26.
- [9] Goos, P. (2006). Optimal versus orthogonal and equivalent-estimation design of blocked and split-plot experiments, *Statistica Neerlandica* **60**, 1–18.
- [10] Letsinger, J.D., Myers, R.H. & Lentner, M. (1996). Response surface methods for bi-randomization structures, *Journal of Quality Technology* **28**, 381–397.
- [11] Vining, G.G., Kowalski, S.M. & Montgomery, D.C. (2005). Response surface designs within a split-plot structure, *Journal of Quality Technology* **37**, 115–129.
- [12] Goos, P., Langhans, I. & Vandebroek, M. (2006). Practical inference from industrial split-plot designs, *Journal of Quality Technology* **38**, 162–179.
- [13] Kenward, M.G. & Roger, J.H. (1997). Small sample inference for fixed effects from restricted maximum likelihood, *Biometrics* **53**, 983–997.
- [14] Gilmour, S.G. & Goos, P. (2006). *Bayesian Analysis of Data from Multi-Stratum Response Surface Designs*, Working Paper 2006/005, Universiteit Antwerpen, Faculty of Applied Economics.
- [15] Huang, P., Chen, D. & Voelkel, J. (1998). Minimum-aberration two-level split-plot designs, *Technometrics* **40**, 314–326.
- [16] Bingham, D. & Sitter, R.R. (1999). Minimum-aberration two-level fractional factorial split-plot designs, *Technometrics* **41**, 62–70.
- [17] Bingham, D.R., Schoen, E.D. & Sitter, R.R. (2004). Designing fractional factorial split-plot experiments with few whole-plot factors, *Journal of the Royal Statistical Society, Series C* **53**, 325–339.
- [18] Bingham, D.R., Schoen, E.D. & Sitter, R.R. (2005). Designing fractional factorial split-plot experiments with few whole-plot factors, *Journal of the Royal Statistical Society, Series C* **54**, 955–958.
- [19] Bisgaard, S. (2000). The design and analysis of $2^{k-p} \times 2^{q-r}$ split plot experiments, *Journal of Quality Technology* **32**, 39–56.
- [20] Bingham, D.R. & Sitter, R.R. (2003). Fractional factorial split-plot designs for robust parameter experiments, *Technometrics* **45**, 80–89.
- [21] Goos, P. (2002). *The Optimal Design of Blocked and Split-plot Experiments*, Springer, New York.
- [22] Goos, P. & Vandebroek, M. (2003). D-optimal split-plot designs with given numbers and sizes of whole plots, *Technometrics* **45**, 235–245.
- [23] Jones, B. & Goos, P. (2007). A candidate-set-free algorithm for generating D-optimal split-plot designs, *Journal of the Royal Statistical Society: series C* **56**, 347–364.
- [24] Goos, P. & Donev, A.N. (2007). Tailor-made split-plot designs for mixture and process variables, *Journal of Quality Technology* **39**, 326–339.

PETER GOOS

Supersaturated Designs

Introduction

A supersaturated design is essentially a **fractional factorial** design in which the number of experimental runs (observations) is less than the number of parameters to be estimated. These designs are useful as screening tools, as the paring of candidate factors is performed in a cost-efficient manner.

Consider the simple fact that where there is an effect, there is a cause. Quality engineers are constantly faced with distinguishing between the factors which have a real effect and those estimated effects that are due to random error. Those “null” factors are then adjusted to lower the cost; those “nonnull” (effective) factors are used to yield better quality. Often, a large number of factors can be listed as possible sources of effect and among those factors only a small portion are actually active. This is sometimes called *effect sparsity*. A problem frequently encountered in this area is how to reduce the total number of experiments.

In order to obtain an unbiased estimate of the **main effect** of each factor, the number of experiments must exceed (or at least be equal to) the number of factors plus one (for estimating the overall grand average). When these two are equal, the design is called a *saturated design* and is the minimum effort required to estimate all main effects. The standard advice given to users in such a screening process is to use the saturated design, which is “optimal” according to certain theoretical optimality criteria. However, the nonsignificant effects are not of interest. Estimating all main effects may be wasteful if the goal is simply to detect those few active factors. If the number of active factors is indeed small, the use of a slightly biased estimate will still allow one to accomplish the identification of the active factors but significantly reduce the amount of experimental work. This is particularly important in situations where the cost of an individual run is high. Developing such **screening designs** has long been a well-recognized problem, certainly since Satterthwaite [1].

When all factors can be reasonably arranged into several groups, the so-called group screening designs can be used (see [2, 3]). Only those factors in groups that are found to have large effects are studied further. The grouping scheme seems to be crucial but has

seldom been discussed. While such methods may be appropriate in certain situations (e.g., blood tests), we are interested in systematic supersaturated designs that can examine k factors in $N < k + 1$ experiments in which no grouping scheme needs to be made.

Optimality Criteria

Optimality criteria that are useful for characterizing the quality of supersaturated designs are discussed in this section. When a supersaturated design is employed, the abandonment of orthogonality is inevitable. It is well known that lack of orthogonality results in lower efficiency; therefore we seek a design that is as “near orthogonal” as possible. One way to measure the degree of nonorthogonality between two columns, c_i and c_j , is to consider their cross-product $s_{ij} = c'_i c_j$. A design is orthogonal when all $s_{ij} = 0$. A larger $|s_{ij}|$ implies less orthogonality. Intuitively, a design should be efficient if $E(s^2)$, the average value of s_{ij}^2 over all pairs is small. This criterion, first proposed by Booth and Cox [4], is given by:

$$E(s^2) = \sum_{1 \leq i < j \leq k} \frac{s_{ij}^2}{\binom{k}{2}} \quad (1)$$

Alternatively, one could consider the maximal value of either $|s_{ij}|$ or s_{ij}^2 over all pairs as a criterion.

For three-level supersaturated designs, Yamada and Lin [5] defined a measure of dependency between two design columns $\chi^2(\mathbf{x}_i, \mathbf{x}_j) = \sum_{u,v=1}^3 \frac{(n_{uv}^{(ij)} - n/9)^2}{n/9}$, where $n_{uv}^{(ij)}$ is the number of (u, v) pairs in $(\mathbf{x}_i, \mathbf{x}_j)$. Then they defined a criterion for the whole design $\mathbf{X}$ by

$$\text{ave} \chi^2 = \sum_{1 \leq i < j \leq k} \frac{\chi^2(\mathbf{x}_i, \mathbf{x}_j)}{\binom{k}{2}} \quad (2)$$

and also showed a lower bound of $\text{ave} \chi^2 \geq \frac{2n(2k - n + 1)}{(n - 1)(k - 1)}$.

Recently, Fang *et al.* [6] proposed a general criterion for supersaturated design, suitable for any mixed-level design. For any two design columns $\mathbf{x}_i$ and $\mathbf{x}_j$, define $f_{\text{NOD}}^{ij} = \sum_{u=1}^{q_i} \sum_{v=1}^{q_j} \left(n_{uv}^{(ij)} - \frac{n}{q_i q_j} \right)^2$, where

2 Supersaturated Designs

$n_{uv}^{(ij)}$ is the number of (u, v) pairs in $(\mathbf{x}_i, \mathbf{x}_j)$, and $n/(q_i q_j)$ stands for the average frequency of level combinations in each pair of columns $\mathbf{x}_i$ and $\mathbf{x}_j$. A new criterion $E(f_{\text{NOD}})$ is defined as the following,

$$E(f_{\text{NOD}}) = \sum_{1 \leq i < j \leq k} \frac{f_{\text{NOD}}^{ij}}{\binom{k}{2}} \quad (3)$$

It is obvious that $E(f_{\text{NOD}}) = 0$ for an **orthogonal array**. Let $\lambda_{ml} = \sum_{j=1}^k 1_{\{x_{mj}=x_{lj}\}}$, where 1_A is the indicator function of A , that is, λ_{ml} is the number of coincidences between the two columns $\mathbf{x}_m$ and $\mathbf{x}_l$. It is obvious that $\lambda_{mm} = k$. It can be shown [6] that

$$\begin{aligned} E(f_{\text{NOD}}) &= \frac{\sum_{m,l=1, m \neq l}^n \lambda_{ml}^2}{k(k-1)} + C(n, q_1, \dots, q_k) \\ &\geq \frac{n \left(\sum_{j=1}^k n/q_j - k \right)^2}{k(k-1)(n-1)} \\ &\quad + C(n, q_1, \dots, q_k) \end{aligned} \quad (4)$$

where $C(n, q_1, \dots, q_k) = \frac{nk}{k-1} - \frac{1}{k(k-1)} \left(\sum_{i=1}^k \frac{n^2}{q_i} + \sum_{i,j=1, j \neq i}^k \frac{n^2}{q_i q_j} \right)$ depends on $\mathbf{X}$ only through $n, q_1, \dots, q_k$, and the lower bound of $E(f_{\text{NOD}})$ can be achieved if and only if $\lambda = \left(\sum_{j=1}^k n/q_j - k \right) / (n-1)$ is a positive integer and all λ_{ml} 's ($m \neq l$) are equal to λ .

For any two factors $\mathbf{x}_i$ and $\mathbf{x}_j$ of the design $\mathbf{X}$, the equal occurrence of -1 and 1 implies that $f_{\text{NOD}}^{ij} = \sum_{u,v=-1,1} \left(n_{uv}^{(ij)} - n/4 \right)^2 = s_{ij}^2/4$. This, in turn, implies that f_{NOD}^{ij} is essentially the same measure as s_{ij}^2 for two-level designs, and as $\chi^2(\mathbf{x}_i, \mathbf{x}_j) = f_{\text{NOD}}^{ij} q_i q_j / n$ for three-level columns. Namely, $E(f_{\text{NOD}}) = \frac{1}{4} E(s^2)$ and $\frac{n}{9} \text{ave} \chi^2$ when all $q_i = 2$ and 3 , respectively.

All the above criteria, however, are only for pairwise relationships. More general criteria have been obtained in Wu [7] and Deng *et al.* [8, 9]. One important feature of the goodness of a supersaturated screening design is its projection property (see [10, 11, 12]). Deng *et al.* [13, 14] thus proposes the r-rank

criterion as *the resolution rank (or r-rank, for short) of $\mathbf{X}$ is defined as $f = d - 1$, where d is the minimum number of subset columns that are linearly dependent.* Detailed properties of r-rank can be found in Deng *et al.* [14].

Supersaturated Designs Using Hadamard Matrices

In this section we discuss some basic methods for constructing supersaturated designs. Lin [15] proposed a class of special supersaturated designs which can be easily constructed via half fractions of the Hadamard matrices. These designs can examine $k = N - 2$ factors with $n = N/2$ runs, where N is the order of the Hadamard matrix used. The Plackett and Burman [16] designs, which can be viewed as a special class of Hadamard matrices, are used to illustrate the basic construction method.

Table 1 shows the original 12-run Plackett and Burman design. We will take column 11 as a *branching* column, meaning that we divide the 12 runs into two groups of 6 runs each, in accord with the sign in column 11. Then the runs (rows) can be split into two groups: group I with the sign of $+1$ in column 11 (rows 2, 3, 5, 6, 7, and 11), and group II with the sign of -1 in column 11 (rows 1, 4, 8, 9, 10, and 12). Deleting column 11 from group I causes columns 1–10 to form a supersaturated design to examine $N - 2 = 10$ factors in $N/2 = 6$ runs (runs 1–6, as indicated in Table 2). It can be shown that if group II is used, the resulting supersaturated design is an equivalent one. In general, a Plackett and Burman [16] design matrix can be split into two half fractions according to a specific branching column whose signs equal $+1$ or -1 . Specifically, take only the rows that have $+1$ in the branching column. Then, the $N - 2$ columns other than the branching column will form a supersaturated design for $N - 2$ factors in $N/2$ runs. Judged by $E(s^2)$, these designs have been shown to be optimal among all other existing supersaturated designs [17].

Moreover, the half-fraction Hadamard matrix of order $n = N/2 = 4t$ is closely related to a *balanced incomplete block design* with $(v, b, r, k) = (2t - 1, 4t - 2, 2t - 2, t - 1)$ and $\lambda = t - 1$. Consequently, the $E(s^2)$ value for a supersaturated design from a half-fraction Hadamard matrix is $n^2(n - 3)/[(2n - 3)(n - 1)]$, which can be shown to

Table 1 A supersaturated design derived from the Hadamard matrix of order 12

R	#	1	2	3	4	5	6	7	8	9	10	11
1	1	+	+	−	+	+	+	−	−	−	+	−
	2	+	−	+	+	+	−	−	−	+	−	+
2	3	−	+	+	+	−	−	−	+	−	+	+
	4	+	+	+	−	−	−	+	−	+	+	−
3	5	+	+	−	−	−	+	−	+	+	−	+
	6	+	−	−	−	+	−	+	+	−	+	+
4	7	−	−	−	+	−	+	+	−	+	+	+
	8	−	−	+	−	+	+	−	+	+	+	−
5	9	−	+	−	+	+	−	+	+	+	−	−
	10	+	−	+	+	−	+	+	+	−	−	−
6	11	−	+	+	−	+	+	+	−	−	−	+
	12	−	−	−	−	−	−	−	−	−	−	−

Table 2 The resulting supersaturated design for $(n, k) = (6, 10)$

Number	$\mathbf{x}_1$	$\mathbf{x}_2$	$\mathbf{x}_3$	$\mathbf{x}_4$	$\mathbf{x}_5$	$\mathbf{x}_6$	$\mathbf{x}_7$	$\mathbf{x}_8$	$\mathbf{x}_9$	$\mathbf{x}_{10}$
1	+	−	+	+	+	−	−	−	+	−
2	−	+	+	+	−	−	−	+	−	+
3	+	+	−	−	−	+	−	+	+	−
4	+	−	−	−	+	−	+	+	−	+
5	−	−	−	+	−	+	+	−	+	+
6	−	+	+	−	+	+	+	−	−	−

be the minimum within the class of designs with same size. Potentially promising theoretical results seem possible for the construction of a half-fraction Hadamard matrix. Theoretical implications deserve detailed scrutiny and are discussed below. For more details regarding this issue, see Cheng [17] and Nguyen [18]. The half-fraction idea can be generalized to fractions of higher-level orthogonal arrays for construction of higher-level supersaturated designs [6].

Note that the interaction columns of Hadamard matrices are only partially confounded with other main-effect columns [10, 11]. Wu [7] makes use of such a property and proposes a supersaturated design that consists of all main-effect and two-factor interaction columns from any given Hadamard matrix of order N . The resulting design has N runs and can accommodate up to $N(N-1)/2$ factors, namely, $N-1$ columns from main effects and $(N-1)(N-2)/2$ columns from interaction effects. When there are $k < N(N-1)/2$ factors to be studied, choosing columns becomes an important issue to be addressed. Provided the interactions are not fully confounded with the main effects, Lin [19] extended

this idea to analyze interaction effects for first-order designs.

Data Analysis Methods

Several methods have been proposed in the literature to analyze the experimental data, given that the number of potential effects exceeds the number of observations. The stepwise regression procedure is probably the most popular one. Chen and Lin [20] show that the largest effect can always be identified, under some very mild conditions. Recent advances in this area, including adjusted p value, Bayesian, and penalized least-squares methods, are discussed below.

Adjusted p -Value Method

When studying so many effects in only a few runs, the probability of a false-positive reading (type I error) is a major risk. Recall that we have null and alternative pairs $H_j : \beta_j = 0$ and $H_j^c : \beta_j \neq 0$. Forward selection proceeds by identifying the maximum F statistics at successive stages. Let $\mathbf{S}$ denote the

4 Supersaturated Designs

complete set of potential effects and $j_1, \dots, j_{s-1}$ denote the effects selected at the first $s - 1$ stages. The F statistic for testing H_j at stage s is then $F_j^{(s)} = \text{SSR}(j|j_1, \dots, j_{s-1})/\text{MSE}(j, j_1, \dots, j_{s-1})$ (SSR: Sum Square Regression; MSE: Mean Square Error). Choose the new effect by

$$j_s = \arg \max_{j \in S - \{j_1, \dots, j_{s-1}\}} F_j^{(s)} \quad (5)$$

Letting $\max F_j^{(s)} = F^{(s)}$ and $p^{(j)} = \text{Prob}\{F^{(j)} > f^{(j)} | \text{all previous } H_i \text{ are true}\}$, where $f^{(j)}$ is the observed value of $F^{(s)}$. The forward selection procedure is defined by selecting variables $j_1, \dots, j_f$, where $p^{(f)} \leq \alpha$ and $p^{(f+1)} > \alpha$. If $p^{(1)} > \alpha$, then no variables are selected.

Algorithmically, at stage j , if $p^{(j)} > \alpha$, then stop; otherwise, enter X_j and continue. This procedure controls the type I error rate exactly at level α under the complete null hypothesis since P (rejects at least one H_i | all H_i true) $= P(F^{(1)} \leq f_\alpha^{(1)}) = \alpha$. In addition, if the first s variables are *forced* and the test is used to evaluate the significance of the next entering variable (of the remaining $k - s$), the procedure is again exact under the complete null hypothesis of no effects among the $k - s$ remaining variables. The exactness disappears with simulated p values, but the errors can be made very small, particularly with control variates. The analysis of data from supersaturated designs along this direction can be found in Westfall *et al.* [21].

Bayesian Method

Beattie *et al.* [22] detail a two-stage Bayesian model selection strategy combining recent methodologies: the stochastic search variable selection (SSVS) method and the intrinsic Bayes factor (IBF) method. The two-stage procedure is able to keep all possible models under consideration while providing a level of robustness akin to Bayesian analyses incorporating noninformative priors. Note that Bayesian methods are able to supplement observational information using prior information on the parameters. This allows straightforward computation of posterior probabilities, a more intuitive concept than the p value. Recent development of the Gibbs sampling algorithm and other Markov chain Monte Carlo techniques, in tandem with computational advances, have brought many Bayesian ideas into the mainstream.

However, the use of these methods in analyzing supersaturated designs may require a departure from noninformative priors to guarantee the existence of a proper posterior distribution, thus leaving room for controversy on the objectivity of the results.

Penalized Least-Squares Method

Li and Lin [23] proposed a variable selection procedure for identifying active factors in supersaturated designs, via nonconvex penalized **least squares**. Suppose that observations $y_1, \dots, y_n$ are an independent sample from the linear regression model

$$y_i = x_i^T \beta + \epsilon_i \quad (6)$$

with $E(\epsilon_i) = 0$, and $\text{Var}(\epsilon_i) = \sigma^2$, where x_i corresponds to a fixed experiment design. Define a penalized least-squares function as

$$Q(\beta) \equiv \frac{1}{2n} \sum_{i=1}^n (Y_i - x_i^T \beta)^2 + \sum_{j=1}^d p_\lambda(|\beta_j|) \quad (7)$$

where $p_\lambda(\cdot)$ is a penalty function, and λ may depend on the **sample size** n and is a tuning parameter, which controls model complexity and can be selected by some data-driven methods. Fan and Li [24] and Li and Lin [23, 25] applied the generalized cross validation to choose λ .

Fan and Li [24] studied the issue of selection of penalty function in depth and advocated the use of the smoothly clipped absolute deviation (SCAD), whose first-order derivative is defined by

$$p'_\lambda(\beta) = \lambda \left\{ I(\beta \leq \lambda) + \frac{(a\lambda - \beta)_+}{(a-1)\lambda} I(\beta > \lambda) \right\} \quad (8)$$

for some $a > 2$ and $\beta > 0$

with $p_\lambda(0) = 0$. In our analysis, we use $a = 3.7$ as advocated in Fan and Li [24]. Given an initial value $\beta^{(0)}$ that is close to the true value of β , when $\beta_j^{(0)}$ is not very close to 0, the penalty $p_\lambda(|\beta_j|)$ can be locally approximated by a quadratic function as $[p_\lambda(|\beta_j|)]' = p'_\lambda(|\beta_j|) \text{sgn}(\beta_j) \approx \{p'_\lambda(|\beta_j^{(0)}|)/|\beta_j^{(0)}|\} \beta_j$, otherwise, set $\hat{\beta}_j = 0$. With the local quadratic approximation, the solution for the penalized least squares can be found by iteratively computing the following **ridge** regression with an initial value $\beta^{(0)}$: $\beta^{(1)} = \{X^T X + n \Sigma_\lambda(\beta^{(0)})\}^{-1} X^T y$, where $\Sigma_\lambda(\beta^{(0)}) = \text{diag}\{p'_\lambda(|\beta_1^{(0)}|)/|\beta_1^{(0)}|, \dots, p'_\lambda(|\beta_d^{(0)}|)/|\beta_d^{(0)}|\}$. For

analysis of supersaturated design, Li and Lin [23, 25] set the level of significance to be 0.1 for both F enter and F remove in the stepwise variable selection procedure for choosing an initial value for the penalized least-squares approaches.

Example: Williams Rubber Experiment

The original problem given in Williams [26] concerned the effects of 24 experimental factors, and 28 runs of a Plackett and Burman design were used. The full data display was also given in Box and Draper [27, p. 175]. Note that columns 13 and 16 are fully confounded in the original setup. Taking the unused orthogonal column (+ - + + - + - + + + - - + - - - + - - - + + + + - - +) as the branching column results in a supersaturated design with 14 runs, as displayed in Table 3 (with corresponding observations). Using the entire data set of 28 observations, Williams [26] identified factors 4,

10, 15, and 20 as the most influential. This result will be viewed as the **benchmark** for all comparisons.

Lin [15] utilized a forward selection procedure on this half fraction for identifying the most important factors. His analysis also identified these four factors in addition to factor 12. Note that sequential selection procedures are prone to overfitting. It is worth mentioning that a stepwise procedure utilizing the AIC criterion (implemented in S-Plus) selected factor 15 only.

The **adjusted p -value** method [21] used **resampling** to estimate the distribution of the maximal F statistic at each step of the procedure. They used those estimates to adjust the p -value cutoffs for inclusion in forward selection. Using their methodology on the Williams half-fraction data set, they concluded that only factor 15 was active.

For the Bayesian method [22], {4, 12, 15, 20} was chosen by the selections of the SSVS algorithm as the model for input into the IBF scheme. The IBF follow-up procedure to SSVS results in a

Table 3 Supersaturated design: half fraction of Williams [26] data

Factors	Experimental runs													
	1	2	3	4	5	6	7	8	9	10	11	12	13	14
1	1	1	1	1	-1	-1	-1	-1	-1	1	-1	1	1	-1
2	1	-1	1	1	-1	-1	-1	1	-1	1	1	-1	1	-1
3	1	-1	-1	-1	1	1	-1	1	-1	1	-1	-1	1	1
4	-1	-1	1	1	1	1	-1	-1	-1	1	1	-1	1	-1
5	-1	-1	1	-1	1	1	1	-1	-1	-1	1	1	1	-1
6	-1	-1	-1	1	1	1	-1	1	1	1	-1	1	-1	-1
7	1	1	-1	-1	-1	1	-1	-1	1	1	-1	1	1	-1
8	1	1	-1	-1	1	-1	1	1	-1	1	1	-1	-1	-1
9	1	1	-1	-1	1	1	-1	-1	-1	-1	1	1	1	-1
10	1	-1	1	1	-1	1	1	1	-1	-1	-1	1	-1	-1
11	1	-1	-1	1	-1	1	-1	-1	1	-1	1	1	-1	1
12	-1	-1	1	-1	-1	-1	1	-1	1	1	-1	1	1	1
13	-1	1	1	1	1	1	-1	-1	1	1	-1	-1	-1	-1
14	1	1	1	1	1	-1	1	-1	-1	-1	-1	1	-1	-1
15	-1	1	1	-1	-1	1	1	-1	-1	1	1	-1	-1	1
16	-1	1	1	1	1	1	-1	-1	1	1	-1	-1	-1	-1
17	1	-1	1	-1	1	1	1	-1	1	1	-1	-1	-1	-1
18	-1	1	-1	1	-1	1	1	-1	1	-1	1	-1	1	-1
19	-1	-1	-1	1	-1	1	1	1	-1	1	1	1	-1	-1
20	1	-1	-1	1	1	1	1	-1	-1	-1	-1	-1	1	1
21	-1	1	-1	-1	-1	1	1	1	-1	1	-1	1	1	-1
22	-1	1	1	-1	1	1	-1	1	-1	-1	-1	1	-1	1
23	-1	-1	1	-1	1	-1	-1	1	1	-1	1	1	1	-1
24	1	-1	-1	1	-1	1	-1	1	1	1	1	1	-1	-1
Responses	133	62	45	52	56	47	88	193	32	53	276	145	130	127

Table 4 Comparative results for analyses of the Williams half-fraction experiment

Model selection method	Factors identified as important
Williams final model [26]	4, 10, 15, 20
Stepwise regression [15]	4,10,12,15,20
Stepwise regression using AIC	15
Adjusted p -value method [21]	15
SSVS	4, 12, 15, 20
IBF using 4, 12, 15, 20 from SSVS [22]	4, 12, 15, 20
Penalized least-squares method with SCAD [23]	4, 12, 15, 20

more moderate, stable final model, with four factors selected from the factors identified by SSVS. Note that the result from the two-stage procedure is *insensitive* to the estimate of σ^2 . Using the all-subsets regression procedure recommended by Abraham *et al.* [28], assuming five or fewer active factors, yields the same four-factor model as the two-stage procedure. However, the Bayesian approach does not need to specify the number of active factors in advance.

The penalized least-squares method [23] identified {4, 12, 15, 20} as the final model. Table 4 summarizes the results of the various analytical techniques applied to the Williams half-fraction experiment. The classical stepwise techniques tend to be either very conservative (using $p = 0.01$), selecting one factor, or very liberal (with $p = 0.10$), selecting 12 factors (not shown in the table). The SSVS/IBF agrees closely with the conservative stepwise methods but while using a different estimate is extremely liberal in factor inclusion.

Theoretical Construction Methods

Deng *et al.* [14] proposed a supersaturated design of the form $\mathbf{X}_c = [\mathbf{H}, \mathbf{RHC}]$, where $\mathbf{H}$ is a normalized Hadamard matrix, $\mathbf{R}$ is an orthogonal matrix, and $\mathbf{C}$ is an $n \times (n - c)$ matrix representing the operation of column selection. The $\mathbf{X}_c$ matrix covers many existing supersaturated designs. This includes the supersaturated designs proposed by Lin [15], Wu [7], Yamada and Lin [29], and Tang and Wu [30]. It can be shown that

$$\mathbf{X}_c' \mathbf{X}_c = \begin{pmatrix} n\mathbf{I}_n & \mathbf{H}'\mathbf{RHC} \\ \mathbf{C}'\mathbf{H}'\mathbf{R}'\mathbf{H} & n\mathbf{I}_{n-c} \end{pmatrix} = \begin{pmatrix} n\mathbf{I}_n & \mathbf{WC} \\ \mathbf{C}'\mathbf{W}' & n\mathbf{I}_{n-c} \end{pmatrix} \quad (9)$$

where $\mathbf{W} = \mathbf{H}'\mathbf{RH} = (w_{ij}) = (\mathbf{h}_i'\mathbf{R}\mathbf{h}_j)$ and $\mathbf{h}_j$ is the j th column of $\mathbf{H}$. Following the argument in Tang and Wu [30] and the theorem below, these designs will always have the minimum $E(s^2)$ values within the supersaturated design of this type.

Theorem 1 Let $\mathbf{H}$ be a Hadamard matrix of order n and $\mathbf{B} = (b_1, \dots, b_r)$ be a $n \times r$ matrix with all entries ± 1 and $\mathbf{V} = \mathbf{H}'\mathbf{B} = (v_{ij}) = \mathbf{h}_i'b_j$, then for any fixed $1 \leq j \leq r$, $n^2 = \sum_{i=1}^n v_{ij}^2$. Furthermore, let $\mathbf{B} = \mathbf{RH}$ and $\mathbf{W} = \mathbf{H}'\mathbf{RH} = (w_{ij})$. We have

1. $\frac{1}{n}\mathbf{W}$ is an $n \times n$ orthogonal matrix,
2. $n^2 = \sum_{i=1}^n w_{ij}^2 = \sum_{j=1}^n w_{ij}^2$,
3. w_{ij} is always a multiple of 4, and
4. if $\mathbf{H}'$ is column balanced, then $\pm n = \sum_{i=1}^n w_{ij} = \sum_{j=1}^n w_{ij}$.

Extension of these results to a more general class of supersaturated designs in the form $\mathbf{S}_K = (\mathbf{R}_1\mathbf{H}\mathbf{C}_1, \dots, \mathbf{R}_K\mathbf{H}\mathbf{C}_K)$ are promising (see [31, 29]). Combinatorial construction methods for supersaturated designs have also been studied, some results can be found in Fang *et al.* [32]. Other recent advances for construction methods are discussed below.

Marginally Oversaturated Designs

A special class of supersaturated designs, called *marginally oversaturated designs* (MOSDs), in which the number of variables under investigation (k) is only slightly larger than the number of experimental runs (n), is presented in Deng *et al.* [13]. The construction method builds on the following theorem.

Theorem 2

1. Let $\mathbf{X} = [\mathbf{H}, \mathbf{v}]$, $\mathbf{w} = \mathbf{H}'\mathbf{v}$, where $\mathbf{H}$ is a Hadamard matrix of order n . Let R_1 be the number of

nonzero entries in $\mathbf{w}$. Then $r = R_1$, where r is the resolution rank of $\mathbf{X}$.

2. Let $\mathbf{X} = [\mathbf{H}, \mathbf{v}_1, \mathbf{v}_2]$ and $\mathbf{w}_1 = \mathbf{H}'\mathbf{v}_1$, $\mathbf{w}_2 = \mathbf{H}'\mathbf{v}_2$, where $\mathbf{H}$ is a Hadamard matrix of dimension n . Let $R_1 = \min[S(\mathbf{w}_1), S(\mathbf{w}_2)]$ and $R_2 = \min[S(b_1\mathbf{w}_1 + b_2\mathbf{w}_2)] + 1$, where $S(\mathbf{u})$ represents the number of nonzero elements in the vector $\mathbf{u}$ and b_1, b_2 can take on all possible values. Then $r = \min[R_1, R_2]$, where r is the resolution rank of $\mathbf{X}$.

For example, with $n = 12$, given $(1, \mathbf{x}_1, \dots, \mathbf{x}_{11})$ form a 12-run Hadamard matrix (as in Table 1), the optimal added columns $\mathbf{v}_1$ and $\mathbf{v}_2$ for MOSD from the theorem are $v_1 = (+ + + + + - - - - -)$ and $v_2 = (- + + + + - - - - -)$.

Three-Level Supersaturated Designs

Given a two-level orthogonal array $\mathbf{C}$ of size (n, m) , Yamada and Lin [5] provided a series of three-level supersaturated designs of the form $(N, K) = (3n, 4m)$ as follows.

$$\mathbf{D} = \begin{bmatrix} \phi^{12}(\mathbf{C}) & \phi^{12}(\mathbf{C}) & \phi^{13}(\mathbf{C}) & \phi^{23}(\mathbf{C}) \\ \phi^{23}(\mathbf{C}) & \phi^{13}(\mathbf{C}) & \phi^{23}(\mathbf{C}) & \phi^{12}(\mathbf{C}) \\ \phi^{31}(\mathbf{C}) & \phi^{23}(\mathbf{C}) & \phi^{12}(\mathbf{C}) & \phi^{13}(\mathbf{C}) \end{bmatrix} \quad (10)$$

where $\phi^{ab}(\cdot)$ is an operator which transforms the elements from -1 to a and from 1 to b on the matrix/vector in $(\cdot)$. Such a design can study $K = 4m$ factors in $N = 3n$ runs. It is shown in Yamada and Lin [5] that such designs have a relatively small value of $\text{ave } \chi^2$ as previously defined.

Multilevel Supersaturated Designs

Fang *et al.* [33] obtained a new class of multilevel supersaturated design by collapsing a U-type **uniform design** to an orthogonal array. This can be illustrated by the following example. Suppose that we extend the orthogonal design $\mathbf{L} = \mathbf{L}_9(3^4)$ and a generating U-type design $\mathbf{U} = \mathbf{U}(9, 9^2)$ to a S-design $S_9(3^8)$ by the collapsing method.

$$\mathbf{U} \oplus \mathbf{L} = \begin{bmatrix} 1 & 1 \\ 2 & 7 \\ 3 & 3 \\ 4 & 9 \\ 5 & 5 \\ 6 & 6 \\ 7 & 2 \\ 8 & 8 \\ 9 & 4 \end{bmatrix} \oplus \begin{bmatrix} 0 & 0 & 0 & 0 \\ 0 & 1 & 1 & 1 \\ 0 & 2 & 2 & 2 \\ 1 & 0 & 1 & 2 \\ 1 & 1 & 2 & 0 \\ 1 & 2 & 0 & 1 \\ 2 & 0 & 2 & 1 \\ 2 & 1 & 0 & 2 \\ 2 & 2 & 1 & 0 \end{bmatrix}$$

$$= \mathbf{X} = \begin{bmatrix} 0 & 0 & 0 & 0 & | & 0 & 0 & 0 & 0 \\ 0 & 1 & 1 & 1 & | & 2 & 0 & 2 & 1 \\ 0 & 2 & 2 & 2 & | & 0 & 2 & 2 & 2 \\ 1 & 0 & 1 & 2 & | & 2 & 2 & 1 & 0 \\ 1 & 1 & 2 & 0 & | & 1 & 1 & 2 & 0 \\ 1 & 2 & 0 & 1 & | & 1 & 2 & 0 & 1 \\ 2 & 0 & 2 & 1 & | & 0 & 1 & 1 & 1 \\ 2 & 1 & 0 & 2 & | & 2 & 1 & 0 & 2 \\ 2 & 2 & 1 & 0 & | & 1 & 0 & 1 & 2 \end{bmatrix}$$

where, the operation $\oplus$ is defined such that the i th block of $\mathbf{X}$ is essentially the matrix $\mathbf{L}$ by taking the row order of column i in matrix $\mathbf{U}$. Namely, the first block (the first four columns) of $\mathbf{X}$ is the original $\mathbf{L}$ and the second block (the last four columns) is formed by rearranging row vectors of $\mathbf{L}$ according to the order defined by the second column of $\mathbf{U}$. The crucial step of this collapsing method is how to choose the best U-type design $U(n, n')$ in the sense of minimizing any specific criterion (see [33]).

Conclusion

Since the study of *random balanced design* in late 1950s [1], the work of supersaturated design has been dormant until the appearance of Lin [15, 34]. The related work about supersaturated design is now much more mature, from better design construction to clearer understanding on design capability, to more reliable data analysis methods. In today's informatic era, as we face more and more variables/parameters, the use of supersaturated design is considered to be a relatively novel endeavor.

Supersaturated designs are very useful in early stages of the experimental investigation of complicated systems and processes involving many factors. They are not used for a terminal experiment. Knowledge of the confounding patterns makes possible the interpretation of the results and provides the understanding of how to plan the follow-up experiments.

Using supersaturated designs involves more risk than using designs with more runs. However, their use is far superior to other experimentation approaches such as subjective selection of factors or changing factors one at a time. The latter can be shown to have unresolvable confounding patterns, though such confounding patterns are important for data analysis and follow-up experiments. The success of a supersaturated design depends heavily on the "effect sparsity" assumption. Consequently, the projection properties

play an important role in designing a supersaturated experiment.

More and more researchers are benefiting from using computer power to construct designs for specific needs. Lin [34, 35] introduced the first computer algorithm to construct supersaturated designs. Nguyen [18] uses an exchange procedure to construct supersaturated designs and near-orthogonal arrays. Li and Wu [36] develop a so-called columnwise – pairwise exchange algorithm. Such an algorithm seems to perform well for constructing supersaturated designs by various criteria. Commercial statistical software for constructing supersaturated designs is also available, see for example, JMP.

References

- [1] Satterthwaite, F. (1959). Random balance experimentation, *Technometrics* **1**, 111–137 (with discussion).
- [2] Lewis, S.M. & Dean, A.M. (2001). Detection of interactions in experiments with large numbers of factors (with discussion), *Journal of the Royal Statistical Society. Series B* **63**, 633–672.
- [3] Watson, G.S. (1961). A study of the group screening methods, *Technometrics* **3**, 371–388.
- [4] Booth, K.H.V. & Cox, D.R. (1962). Some systematic supersaturated designs, *Technometrics* **4**, 489–495.
- [5] Yamada, S. & Lin, D.K.J. (1999). Three-level supersaturated designs, *Statistics and Probability Letters* **45**, 31–39.
- [6] Fang, K.T., Lin, D.K.J. & Liu, M.Q. (2003). Optimal mixed-level supersaturated designs and computer experiment designs, *Metrika* **58**, 279–291.
- [7] Wu, C.F.J. (1993). Construction of supersaturated designs through partially aliased interactions, *Biometrika* **80**, 661–669.
- [8] Deng, L.Y., Lin, D.K.J. & Wang, J.N. (1994b). Supersaturated design using Hadamard matrix, IBM Research Report, RC19470, IBM Watson Research Center.
- [9] Deng, L.Y., Lin, D.K.J. & Wang, J.N. (1996b). A measurement of multifactor orthogonality, *Statistics and Probability Letters* **28**, 203–209.
- [10] Lin, D.K.J. & Draper, N.R. (1992). Projection properties of Plackett and Burman designs, *Technometrics* **34**, 423–428.
- [11] Lin, D.K.J. & Draper, N.R. (1993). Generating alias relationships for two-level Plackett and Burman designs, *Computational Statistics and Data Analysis* **15**, 147–157.
- [12] Lin, D.K.J. (1993b). Another look at first-order saturated designs: the p -efficient designs, *Technometrics* **35**, 284–292.
- [13] Deng, L.Y., Lin, D.K.J. & Wang, J.N. (1996a). Marginally oversaturated designs, *Communications in Statistics* **25**(11), 2557–2573.
- [14] Deng, L.Y., Lin, D.K.J. & Wang, J.N. (1999). On resolution rank criterion for supersaturated design, *Statistica Sinica* **9**, 605–610.
- [15] Lin, D.K.J. (1993a). A new class of supersaturated designs, *Technometrics* **35**, 28–31.
- [16] Plackett, R.L. & Burman, J.P. (1946). The design of optimum multifactorial experiments, *Biometrika* **33**, 303–325.
- [17] Cheng, C.S. (1997). $E(s^2)$ -optimal supersaturated designs, *Statistica Sinica* **7**, 929–939.
- [18] Nguyen, N.K. (1996). An algorithmic approach to constructing supersaturated designs, *Technometrics* **38**, 69–73.
- [19] Lin, D.K.J. (1998). Spotlight interaction effects in main-effect designs, *Quality Engineering* **11**, 133–139.
- [20] Chen, J.H. & Lin, D.K.J. (1998). On identifiability of supersaturated designs, *Journal of Statistical Planning and Inference* **72**, 99–107.
- [21] Westfall, P.H., Young, S.S. & Lin, D.K.J. (1998). Forward selection error control in the analysis of supersaturated designs, *Statistica Sinica* **8**, 101–117.
- [22] Beattie, S.D., Fong, D.H. & Lin, D.K.J. (2002). Two-stage model selection strategy for supersaturated designs, *Technometrics* **44**, 55–63.
- [23] Li, R. & Lin, D.K.J. (2002). Data analysis in supersaturated designs, *Statistics and Probability Letters* **59**, 135–144.
- [24] Fan, J. & Li, R. (2001). Variable selection via non-concave penalized likelihood and its oracle properties, *Journal of the American Statistical Association* **96**, 1348–1360.
- [25] Li, R. & Lin, D.K.J. (2003). Comparisons of variable selection approaches for analysis of supersaturated design, *Journal of Data Science* **1**, 249–260.
- [26] Williams, K.R. (1968). Designed experiments, *Rubber Age* **100**, 65–71.
- [27] Box, G.E.P. & Draper, N.R. (1987). *Empirical Model Building and Response Surfaces*, John Wiley & Sons, New York.
- [28] Abraham, B., Chipman, H. & Vijayan, K. (1999). Some risks in the construction and analysis of supersaturated designs, *Technometrics* **41**, 135–141.
- [29] Yamada, S. & Lin, D.K.J. (1997). Supersaturated designs including an orthogonal base, *The Canadian Journal of Statistics* **25**, 203–213.
- [30] Tang, B. & Wu, C.F.J. (1997). A method for construction of supersaturated designs and its $E(s^2)$ optimality, *The Canadian Journal of Statistics* **25**, 191–201.
- [31] Liu, Y.F. & Dean, A.M. (2004). k -circulant supersaturated designs, *Technometrics* **46**, 32–43.
- [32] Fang, K.T., Ge, G.N., Liu, M.Q. & Qin, H. (2004). Combinatorial construction for optimal supersaturated designs, *Discrete Mathematics* **279**, 191–202.
- [33] Fang, K.T., Lin, D.K.J. & Ma, C.X. (2000). On the construction of multi-level supersaturated designs, *Journal of Statistical Planning and Inference* **86**, 239–252.

-
- [34] Lin, D.K.J. (1991). *Systematic Supersaturated Designs*, Working Paper, No. 264, College of Business Administration, University of Tennessee.
- [35] Lin, D.K.J. (1995). Generating systematic supersaturated designs, *Technometrics* **37**, 213–225.
- [36] Li, W.W. & Wu, C.F.J. (1997). Columnwise-pairwise algorithms with applications to the construction of supersaturated designs, *Technometrics* **39**, 171–179.

Further Reading

Deng, L.Y. & Lin, D.K.J. (1994a). Criteria for supersaturated designs, *Proceedings of the Section on Physical and Engineering Sciences*, American Statistical Association, pp. 124–128.

DENNIS K.J. LIN

Uniform Experimental Designs

Introduction

Because laboratory experiments and computer simulations cost time and money, researchers want to plan their experiments or computer simulations in a way that provides the most insight from a limited number of *runs*. This involves carefully choosing the values of the *factors*, also known as *explanatory variables*, *independent variables*, or *controllable parameters*, in order to understand how they affect the *response*, also known as the *dependent variable*. A given choice of factor values for each of the runs is called the *experimental design*. This article describes *uniform designs*, that is, experimental designs that sample the space of possible factor values, called the *experimental domain*, as evenly as possible.

Consider the problem of estimating the response (y) as a function of several controlled factors ($x_1, \dots, x_s$). We may have a model

$$y = f(x_1, \dots, x_s) + \varepsilon \quad (1)$$

The noise term ε is usually present in laboratory experiments, but may vanish in computer simulations. When the function f is unknown in physical experiments or f is too complicated in **computer experiments**, one would like to find a good approximation to the function f , denoted by $\hat{f}$. Let us consider two examples: one is a laboratory experiment from the chemical industry and another is a computer experiment arising in robotics.

Example 1: Chemical yield. A chemical experiment is conducted in order to find the manufacturing conditions that maximize the yield. The experimenter chooses four factors considered crucial to the production process: x_1 , the amount of formaldehyde (mol); x_2 , the reaction temperature ($^{\circ}\text{C}$); x_3 , the reaction time (h); and x_4 , the amount of potassium carbolic. The response y is the yield of the process. The experimenter does not know the model f relating the response and the four factors in equation (1) or even its form. The experimenter wants to find a good approximation to the model f via an experiment and the “best” four-factor-value combination in the sense of maximizing y .

Example 2: Robot arm. The movement trajectory of a robot arm is frequently used as an illustrative example in the **neural network** literature as well as in computer experiments. Consider a robot arm with m segments. The shoulder of the arm is fixed at the origin in the (u, v) plane. The segments of this arm have lengths L_j , $j = 1, \dots, m$. The first segment is at angle θ_1 with respect to the horizontal (u) axis. For $k = 2, \dots, m$, segment k makes an angle θ_k with segment $k - 1$. The end of the robot arm is at (u_m, v_m) , where

$$\begin{cases} u_m = \sum_{j=1}^m L_j \cos(\theta_1 + \dots + \theta_j) \\ v_m = \sum_{j=1}^m L_j \sin(\theta_1 + \dots + \theta_j) \end{cases} \quad (2)$$

and the response y is the distance $y = \sqrt{u_m^2 + v_m^2}$ from the end of the arm to the origin expressed as a function of $2m$ variables $\theta_j \in [0, 2\pi]$ and $L_j \in [0, 1]$. One wants to find a simple approximate model $y \approx \hat{f}(\theta_1, \dots, \theta_m, L_1, \dots, L_m)$ for the robot arm with m -segments as a surrogate for the true model f .

The above examples require the following:

1. **Design.** A design can help the experimenter discover valuable information about the underlying model. Because the form of the underlying model is typically unknown, a natural idea is to spread the experimental points as evenly as possible on the experimental domain.
2. **Modeling.** By various parametric or nonparametric model selection and **estimation** techniques, one wants to find a good approximation model $\hat{f}$ – also called a *metamodel* in the computer experiment literature.

The uniform design and other designs (e.g., the Latin hypercube design has been popular, *see Latin Hypercube Designs* for the details) can be used for the above purpose. The uniform design has been widely used in computer experiments (*see Computer Experiments; Latin Hypercube Designs*), physical experiments under model uncertainty and experiments with mixtures (*see Mixture Experiments*). It was proposed by Fang *et al.* [1–3]. Comprehensive reviews are given by Fang and Hickernell [4], Fang and Lin [5], Fang *et al.* [6], and Fang and Chan [7].

Experimental Design

The overall mean model has been used for developing theory and methodology of the uniform design: $y = \mu + h(x_1, \dots, x_s) + \varepsilon$. Let $\mathcal{P} = \{x_1, \dots, x_n\} \subset \mathcal{X}$ be an n run, s factor design. Suppose that the experimenter wants to estimate μ , the overall mean of the response y on the experimental domain $\mathcal{X}$. *The best design for this purpose is one whose empirical distribution approximates the uniform distribution.* This idea arose first in numerical integration methods for high-dimensional problems, called *quasi-Monte Carlo methods*, that were proposed in the early 1960s. The Koksma–Hlawka inequality (see Niederreiter [8] and Hickernell [9]) defines the discrepancy, $D(\cdot)$, or measure of uniformity, that quantifies the difference between the uniform distribution and the empirical distribution of the design. Designs with minimum discrepancy are called *uniform designs*.

The discrepancy is a Kolmogorov–Smirnov type goodness-of-fit statistic. Hickernell [10] showed that for estimating μ in the overall mean model the uniform design has optimal *average mean-square error* assuming random h and optimal *maximum mean-square error* assuming deterministic h . These facts mean that the uniform design is a kind of robust design (see **Robust Design**).

In fact, there are different forms of discrepancy, depending on the norm used to measure the difference between the uniform distribution and the empirical distribution of the design. The centered L_2 discrepancy (CD) proposed by Hickernell [11]

is a popular choice due to several nice properties, including invariance when the points are reflected about a plane through the middle of the domain.

Many authors have proposed constructions of uniform designs with various n and s . The notation $U_n(q^s)$ is for symmetric uniform designs and $U_n(q_1 \times \dots \times q_s)$ is for nonsymmetric uniform designs, very similar to the case for orthogonal design tables (see **Orthogonal Arrays**). For example, the first five columns of Table 1 represent a design $U_{15}(15^4)$. Constructions of uniform designs are made via the good-lattice-point method or integration lattices, digital nets, combinatorial methods, cutting methods as well as numerical search. A number of uniform designs can be found in the web site “<http://uniformdesign.org>”.

Even though the original theory of the uniform design is based on the overall mean model, such designs are useful for more general situations where the true model is unknown. Let us illustrate the method of the uniform design via Example 1. On the basis of familiarity with the process, the experimenter chooses the experimental region to be $[2, 6.2] \times [6, 55] \times [0.5, 7.5] \times [15, 85]$ and 15 equally spaced levels for each factor. Suppose that the amount of formaldehyde, the reaction temperature, the reaction time, and the amount of potassium carbolic are arranged according to the columns **1, 2, 3, 4** of $U_{15}(15^4)$ and use the real levels to replace levels 1, 2, $\dots$, 15 for each factor that forms four columns x_1, x_2, x_3, x_4 in Table 1. This results in a design with 15 level combinations. Their corresponding response,

Table 1 $U_{15}(15^4)$, related design and responses

Number	1	2	3	4	x_1	x_2	x_3	x_4	y
1	9	1	4	10	4.4	6	2	60	0.0835
2	11	14	5	5	5	51.5	2.5	35	0.1847
3	10	7	10	1	4.7	27	5	15	0.1121
4	3	12	3	2	2.6	44.5	1.5	20	0.1121
5	13	10	2	13	5.6	37.5	1	75	0.1663
6	5	4	8	15	3.2	16.5	4	85	0.1522
7	2	2	12	6	2.3	9.5	6	40	0.0046
8	12	5	14	11	5.3	20	7	65	0.2445
9	1	8	6	12	2	30.5	3	70	0.1436
10	7	6	1	7	3.8	23.5	0.5	45	0.1138
11	4	15	9	9	2.9	55	4.5	55	0.1357
12	14	3	7	3	5.9	13	3.5	25	0.1618
13	6	9	15	4	3.5	34	7.5	30	0.0640
14	15	11	11	8	6.2	41	5.5	50	0.2014
15	8	13	13	14	4.1	48	6.5	80	0.2618

that is, chemical yield, is given in the last column of the table. The larger the y value, the higher the output.

Modeling

On the basis of the experimental data one wishes to find a metamodel to fit the data. This metamodel should approximate the true model well, where the true model is known in computer experiments and may be unknown in some physical experiments. There are many modeling techniques in the literature, such as linear, quadratic, polynomial or nonlinear regression, **spline**, kriging, the Bayesian approach, and so on. The reader may refer to Fang *et al.* [6] and references therein for details.

For linear regression models

$$y = \beta_1 + \beta_2 g_2(x) + \cdots + \beta_p g_p(x) + \varepsilon \quad (3)$$

one usually applies model selection techniques to decide the factors, x_j , to include in the model, and the forms of the terms, $g_l(x)$, involving those factors. To estimate the regression coefficient, $\beta = (\beta_1, \dots, \beta_p)^T$, the regression matrix $\mathbf{G}(\mathbf{X}) = (g_l(x_i))_{i,l=1}^{n,p}$ must have full rank, that is, linearly independent columns, and this depends on the design $\mathbf{X} = \{x_1, \dots, x_n\}$.

For the data in Table 1 we might try many modeling methods to find metamodels and then choose the best one among these candidate models. If we employ a centered quadratic regression model with model selection to choose the significant terms, we obtain

$$\begin{aligned} \hat{y} = & -0.103 + 0.0287x_1 + 0.00136x_2 + 0.00295x_3 \\ & + 0.00146x_4 + 0.000429(x_1 - 4.1)^2 \\ & + 0.00509(x_1 - 4.1)(x_3 - 4) \\ & + 0.000499(x_3 - 4)(x_4 - 50) \end{aligned} \quad (4)$$

with $R^2 = 0.91$ and error standard deviation $\hat{\sigma} = 0.0283$. The maximum of $\hat{y}$ on the experimental domain according to this model occurs at $x_1 = 6.2$, $x_2 = 55$, $x_3 = 7.5$, $x_4 = 85$ and the maximum $\hat{y}$ value is 0.40779, which is much larger than the maximum value among 15 runs, $y_{15} = 0.2618$. The uniform design with linear regression modeling suggests the optimal level-combination, which should then be verified by some additional experiments.

Applications and Discussion

The uniform design has a number of good statistical properties. In addition to its optimality for the overall mean model mentioned above, Xie and Fang [12] pointed out that the uniform measure is admissible and minimax under a nonparametric regression model. Hickernell and Liu [13] considered the efficiency and robustness of uniform experimental designs for linear regression models. They found that “Although it is rare for a single design to be both maximally efficient and robust, it is shown here that uniform designs limit the effects of **aliasing** to yield reasonable efficiency and robustness together.”

The uniform design provides flexibility in design and modeling for the user and can be used for complicated nonlinear models. The uniform design is robust against model specification (*see Robust Design*). Although the number of levels in a uniform design is typically equal to the number of runs, a smaller number of levels can be used by collapsing adjacent levels into one level. The uniform design can be applied to experiments with qualitative and quantitative factors and it can also be used for experiments with mixtures. Many products are formed by mixtures of several ingredients, for example, building construction concrete consists of sand, water, and one or more types of cement. Designs for deciding how to mix the ingredients have played an important role in various fields such as chemical engineering, rubber industry, and material and pharmaceutical engineering. Cornell [14] gave a comprehensive treatment of the design and analysis of experiments with mixture (*see Mixture Experiments*). The uniform design can be also applied to such experiments even with some constraints (*see Fang and Yang [15]*).

There are close relationships among the **fractional factorial** design, the supersaturated design (*see Supersaturated Designs*) and the uniform design. For example, various discrepancies (centered L_2 discrepancy, wraparound L_2 discrepancy, and discrete discrepancy) can be used for comparing orthogonal/supersaturated designs and detecting nonisomorphic orthogonal designs.

References

- [1] Fang, K.T. (1980). The uniform design: application of number-theoretic methods in experimental design, *Acta Mathematicae Applicatae Sinica* **3**, 363–372.

- [2] Wang, Y. & Fang, K.T. (1981). A note on uniform distribution and experimental design, *Kexue Tongbao* **26**, 485–489.
- [3] Fang, K.T. & Wang, Y. (1994). *Number-Theoretic Methods in Statistics*, Chapman & Hall, London.
- [4] Fang, K.T. & Hickernell, F.J. (1995). *The Uniform Design and Its Applications*, Bulletin of The International Statistical Institute, Beijing, 50th Session, Book 1, pp. 339–349.
- [5] Fang, K.T. & Lin, D.K.J. (2003). Uniform designs and their application in industry, in *Handbook on Statistics 22: Statistics in Industry*, Khattree, R. & Rao, C.R., eds, Elsevier, North-Holland, pp. 131–170.
- [6] Fang, K.T., Li, R. & Sudjianto, A. (2005). *Design and Modeling for Computer Experiments*, Chapman & Hall/CRC Press, London.
- [7] Fang, K.T. & Chan, L.Y. (2006). Uniform design and its industrial applications, in *Springer Handbook of Engineering Statistics*, Pham, H., ed, Springer, New York, pp. 229–247.
- [8] Niederreiter, H. (1992). Random number generation and quasi-Monte Carlo methods, in *SIAM CBMS-NSF Regional Conference Series in Applied Mathematics*, Philadelphia.
- [9] Hickernell, F.J. (2006). Koksma-Hlawka Inequality, in *Encyclopedia of Statistical Sciences*, 2nd Edition, S. Kotz, N.L. Johnson, C.B. Read, N. & Balakrishnan, B. Vidakovic, eds, John Wiley & Sons, Hoboken, pp. 3862–3867.
- [10] Hickernell, F.J. (1999). Goodness-of-fit statistics, discrepancies and robust designs, *Statistics and Probability Letters* **44**, 73–78.
- [11] Hickernell, F.J. (1998). A generalized discrepancy and quadrature error bound, *Mathematics of Computation* **67**, 299–322.
- [12] Xie, M.Y. & Fang, K.T. (2000). Admissibility and minimaxity of the uniform design in nonparametric regression model, *Journal of Statistical Planning and Inference* **83**, 101–111.
- [13] Hickernell, F.J. & Liu, M.Q. (2002). Uniform designs limit aliasing, *Biometrika* **89**, 893–904.
- [14] Cornell, J.A. (1990). *Experiments with Mixtures, Designs, Models and the Analysis of Mixture Data*, John Wiley & Sons, New York.
- [15] Fang, K.T. & Yang, Z.H. (1999). On uniform design of experiments with restricted mixtures and generation of uniform distribution on some domains, *Statistics and Probability Letters* **46**, 113–120.

Further Reading

- Fang, K.T., Lin, D.K.J., Winker, P. & Zhang, Y. (2000). Uniform design: theory and applications, *Technometrics* **42**, 237–248.
- Fang, K.T. & Mukerjee, R. (2000). A connection between uniformity and aberration in regular fractions of two-level factorials, *Biometrika* **87**, 193–198.
- McKay, M.D., Beckman, R.J. & Conover, W.J. (1979). A comparison of three methods for selecting values of input variables in the analysis of output from a computer code, *Technometrics* **21**, 239–245.

KAI-TAI FANG AND FRED J. HICKERNELL

Box–Behnken Designs

Historical Background

The topic of experimental designs for second-order models is one that has received a great deal of attention in the literature since Box and Wilson's [1] inception of response surface methodology (*see Response Surface Methodology*). In order to estimate a second-order term for a given factor, there must be at least three levels of that factor in the experimental design. A natural design choice for estimating a quadratic model is the 3^k factorial design (*see Factorial Experiments*) in which k denotes the number of experimental factors. However, when k becomes moderately large, the 3^k factorial design may involve an impractical number of design points. In order to reduce the number of design points, Finney [2] considered fractional factorial (*see Fractional Factorial Designs*) subsets of the complete 3^k design. Box and Behnken [3] defined a *redundancy factor* equal to the number of design points divided by the number of model parameters. They pointed out that for situations in which the experimental error variance is not so large as to require large numbers of observations to obtain necessary precision of model parameters, designs having small redundancy factors are desirable. Applying the redundancy factor to a one-third replicate of a 3^5 factorial design when the interest is in estimating 21 model parameters (intercept, 5 linear terms, 10 two-factor interactions, and 5 pure quadratic effects) the redundancy factor is $81/21 \approx 3.9$. In order to reduce the redundancy factor, Box and Behnken [3] proposed the class of designs which are now commonly referred to as *Box–Behnken designs*. Henceforth, we refer to the Box–Behnken design as a *BBD*.

Box–Behnken Design Construction

BBDs represent a family of three-level designs that are highly efficient for estimating second-order response surfaces. The formulation of these designs is accomplished by combining two-level factorial designs (*see Factorial Experiments*) with incomplete block designs (*see Block Designs, Incomplete*) in a specific manner. This formulation is best illustrated *via* an example. Consider the balanced incomplete block design with three treatments and three

blocks, each of size two, given in Table 1. The BBD in three factors is obtained by combining this incomplete block design with 2^2 factorial designs which involve all combinations of the three factors. The two asterisks in every row of the balanced incomplete block design of Table 1 are replaced by the two columns of the 2^2 factorial design in the corresponding factors. Wherever an asterisk does not appear, a column of 0's is inserted. The resulting three-factor BBD is given in Table 2. Note that the last row in the design is a vector of n_c center runs.

Although many BBDs can be constructed by combining factorial designs with balanced incomplete block designs, sometimes the construction must involve a partially balanced incomplete block design instead. In the construction this implies that not every factor will occur in a two-level factorial structure the same number of times with every other factor. As an example, consider $k = 6$ factors and the $N = 48 + n_c$ run design provided in Table 3. Here, each row of the design matrix involves eight design points and the ± 1 notation indicates a 2^3 factorial structure among the corresponding pairs of variables. Note the partial

Table 1 Balanced incomplete block design with three treatments and three blocks

	Factor		
	X_1	X_2	X_3
Block 1	*	*	
Block 2	*		*
Block 3		*	*

Table 2 Box–Behnken design for $k = 3$ factors

X_1	X_2	X_3
−1	−1	0
−1	1	0
1	−1	0
1	1	0
−1	0	−1
−1	0	1
1	0	−1
1	0	1
0	−1	−1
0	−1	1
0	1	−1
0	1	1
0	0	0

2 Box–Behnken Designs

balance by observing that X_1 occurs with X_4 twice (rows 1 and 4) whereas X_1 occurs with X_2, X_3, X_5 , and X_6 only one time. See Myers and Montgomery [4, pp. 345–347] for more details regarding the construction of BBDs.

Characteristics of the Box–Behnken Design

Upon observing the design in Table 2, it is evident that the BBD is comparable to the central composite design (CCD; *see Central Composite Designs*) in terms of design size. The CCD for $k = 3$ factors contains $14 + n_c$ runs and the BBD contains $12 + n_c$ runs. The CCD and the BBD are also quite comparable in size for larger numbers of factors. Borkowski [5, 6] develops prediction variance properties of the CCD and BBD and then uses variance dispersion plots (see Giovannitti-Jensen and Myers [7]) to compare the two designs in terms of minimum, maximum, and average spherical prediction variances. He notes that the two designs possess good prediction variance properties, especially when a sufficient number of center runs are utilized. In terms of other alphabetic optimality (*see Optimal Design; Assessment of Experimental Designs*) efficiencies for the BBD, Lucas [8] reports D-, G-, A-, and E-efficiencies for varying numbers of center runs. Borror *et al.* [9] compare the BBD to the CCD for estimating the slope variance when one is in the robust design (*see Robust Design*) setting. Davison [10, pp. 91–102] compares the CCD and BBD in terms of prediction variance properties and D-efficiency within the context of a split-plot **randomization** (*see Split-Plot Designs*) scheme. Parker *et al.* [11, 12] provide a generalized derivation of equivalence conditions for the BBD and CCD when split-plot randomization is present.

Other than offering appealing prediction variance properties as well as small design sizes, the BBD also

displays other desirable characteristics. Whittinghill [13] discusses the robustness of BBDs to the unavailability of data. The BBD generally allows sufficient information for testing **lack of fit**, it is often close to being rotatable (*see Rotatable Designs and Rotatability*), and the design is often suitable for orthogonal blocking (*see Blocking*). To illustrate the degrees of freedom (df) for lack of fit as well as suitability for orthogonal blocking, consider the BBD for $k = 4$ factors with $n_c = 6$ center runs (*see Center Points*), which is provided in Table 4. The ± 1 notation here indicates the 2^2 factorial structure among the corresponding pairs of variables. Note that the vectors of 0's in the rows which contain factorial points are 4×1 vectors whereas the vectors of 0's in the last row of each block are 2×1 vectors. The design in Table 4 is composed of $N = 24$ factorial points + 6 center runs. Ignoring blocks and assuming a quadratic model (intercept, four linear main effects (*see Main Effects*), six two-factor interactions (*see Interactions*), and four pure quadratics), there are 10 df for lack of fit and 5 df for pure error. For a catalog of BBDs with up to $k = 10$ factors and possible orthogonal blocking schemes, see Box and Behnken [3].

A final important characteristic of BBDs is that their design region is spherical. Figure 1 displays the design points of the $k = 3$ BBD from Table 2. Note that all of the points lie on the edges of the cube and are equidistant (distance = $\sqrt{2}$ units) from the design center. Although the BBD provides good coverage of the cube, there are no points on the corners and the BBD is not a good choice in situations when the experimenter is interested in predicting at the corners. When the experimenter desires more complete coverage of the cube, a face-centered CCD is preferred to the BBD.

Table 3 Box–Behnken design for $k = 6$ factors

X_1	X_2	X_3	X_4	X_5	X_6
± 1	± 1	0	± 1	0	0
0	± 1	± 1	0	± 1	0
0	0	± 1	± 1	0	± 1
± 1	0	0	± 1	± 1	0
0	± 1	0	0	± 1	± 1
± 1	0	± 1	0	0	± 1
0	0	0	0	0	0

Table 4 Box–Behnken design for $k = 4$ factors in three blocks

	X_1	X_2	X_3	X_4
Block 1	± 1	± 1	0	0
	0	0	± 1	± 1
	0	0	0	0
Block 2	± 1	0	0	± 1
	0	± 1	± 1	0
	0	0	0	0
Block 3	± 1	0	± 1	0
	0	± 1	0	± 1
	0	0	0	0

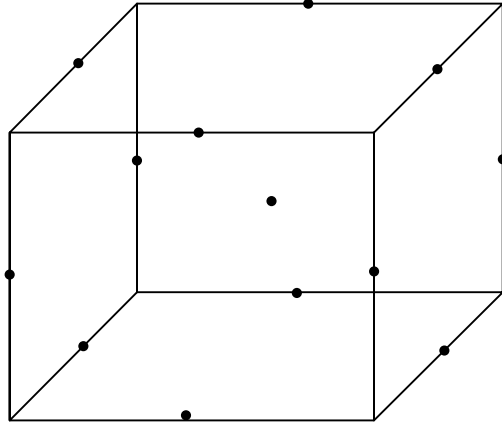


Figure 1 The $k = 3$ Box–Behnken design with a single center point

Table 5 Box–Behnken design for viscosity example

X_1 (temperature)	X_2 (agitation)	X_3 (addition rate)	Y (viscosity)
0	1	1	65
1	−1	0	58
1	0	−1	45
0	0	0	65
0	1	−1	64
0	0	0	59
1	0	1	60
−1	1	0	59
0	−1	1	53
−1	0	1	35
−1	0	−1	64
1	1	0	56
0	0	0	62

Table 6 Estimated regression coefficients for full quadratic model in viscosity example

Term	Coefficient estimate	SE	T	P
Constant	62	1.848	33.542	0.000
X_1	1	1.132	0.883	0.417
X_2	2.625	1.132	2.319	0.068
X_3	−2.375	1.132	−2.098	0.090
X_1^2	−7.375	1.666	−4.426	0.007
X_2^2	1.875	1.666	1.125	0.312
X_3^2	−3.625	1.666	−2.176	0.082
$X_1 \times X_2$	−2	1.601	−1.249	0.267
$X_1 \times X_3$	11	1.601	6.872	0.001
$X_2 \times X_3$	1.75	1.601	1.093	0.324

Example

Consider the following example from Myers and Montgomery [4, p. 347] concerning the production of a particular polyamide resin. One step in the production involves the addition of amines. It was felt that the manner of addition has an impact on the molecular weight distribution of the resin. Three factors were under consideration: temperature at the time of addition (X_1 , °C), agitation (X_2 , rpm), and rate of addition (X_3 , min^{−1}). The viscosity of the resin (y) was recorded as an indirect measure of molecular weight. The BBD used, along with the coded levels of the factors and corresponding resin values, is provided in Table 5.

Notice that this design corresponds to the three-factor BBD ($n_c = 3$) provided in Table 2, but with the run order randomized. A full second-order model (intercept, three linear main effects, three pure quadratics, and three two-factor interactions) was fit and the results are found in Tables 6 and 7.

Note that with only 15 experimental runs, the BBD

Table 7 Summary analysis of variance table for full quadratic model in viscosity example

Source	df	Seq SS ^(a)	Adj SS ^(a)	Adj MS ^(a)	F	P
Regression	9	882.48	882.48	98.054	9.57	0.011
Linear	3	108.25	108.25	36.083	3.52	0.105
Square	3	261.98	261.98	87.328	8.52	0.021
Interaction	3	512.25	512.25	170.75	16.66	0.005
Residual error	5	51.25	51.25	10.25	—	—
Lack of fit	3	33.25	33.25	11.083	1.23	0.477
Pure error	2	18	18	9	—	—
Total	14	933.73	—	—	—	—

^(a)Seq SS: sequential **sum of squares**; Adj SS: adjusted sum of squares; Adj MS: adjusted mean square

allows for **estimation** of all model coefficients, 3 df for lack of fit and 2 df for pure error. The lack-of-fit test suggests that a quadratic model is sufficient. Further details of the analysis of this example can be found in Myers and Montgomery [4, pp. 348–350].

Conclusions

The BBD is a popular design for experimenters who wish to fit a second-order response surface model. The design region is spherical and should only be used in situations where the experimenter is not interested in predicting at the corners of a cuboidal region. The design is not only economical in terms of design runs, but it also possesses good prediction variance properties when conducted in either the completely randomized setting or a split-plot setting. Box and Behnken [3] provide a catalog of BBDs for up to $k = 10$ factors along with possible orthogonal blocking schemes. The popularity of the BBD has been recognized by major statistical software packages such as SAS JMP, MINITAB, and Design Expert and these designs can be easily constructed using these software packages.

References

- [1] Box, G.E.P. & Jones, S. (1951). On the experimental attainment of optimum conditions, *Journal of the Royal Statistical Society. Series B* **13**, 1–45.
- [2] Finney, D.J. (1945). The fractional replication of experiments, *Annals of Eugenics* **12**, 291–301.
- [3] Box, G.E.P. & Behnken, D.W. (1960). Some new three-level designs for the study of quantitative variables, *Technometrics* **2**, 455–475.
- [4] Myers, R.H. & Montgomery, D.C. (2002). *Response Surface Methodology: Process and Product Optimization using Designed Experiments*, John Wiley & Sons, New York, pp. 345–350.
- [5] Borkowski, J.J. (1995a). Minimum, maximum, and average spherical prediction variances for central composite and Box–Behnken designs, *Communications in Statistics: Theory and Methods* **24**, 2581–2600.
- [6] Borkowski, J.J. (1995b). Spherical prediction-variance properties of central composite and Box–Behnken designs, *Technometrics* **37**, 399–410.
- [7] Giovannitti-Jensen, A. & Myers, R.H. (1989). Graphical assessment of the prediction capability of response surface designs, *Technometrics* **31**, 159–171.
- [8] Lucas, J.M. (1977). Design efficiencies for varying numbers of centre points, *Biometrika* **64**, 145–147.
- [9] Borror, C.M., Montgomery, D.C. & Myers, R.H. (2002). Evaluation of statistical designs for experiments involving noise variables, *Journal of Quality Technology* **34**, 54–70.
- [10] Davison, J.J. (1995). *Response surface designs and analysis for bi-randomization error structures*. Unpublished Ph.D. dissertation, Virginia Tech Department of Statistics.
- [11] Parker, P.A., Kowalski, S.M. & Vining, G.G. (2006). Classes of Split-Plot Response Surface Designs for Equivalent Estimation, *Quality and Reliability Engineering International* **22**, 291–305.
- [12] Parker, P.A., Kowalski, S.M. & Vining, G.G. (2007). Construction of Balanced Equivalent Estimation Second-Order Split-Plot Designs, *Technometrics* **49**, 56–65.
- [13] Whittinghill, D.C. (1998). A note on the robustness of Box–Behnken designs to the unavailability of data, *Metrika* **48**, 49–52.

Related Articles

Analysis of Variance; Assessment of Experimental Designs; Block Designs, Incomplete; Blocking; Factorial Experiments; Fractional Factorial Designs; Interactions; Main Effects; Optimal Design; Response Surface Methodology; Robust Design; Rotatable Designs and Rotatability; Split-Plot Designs.

TIMOTHY J. ROBINSON

Center Points

Introduction

In an experimental design with k quantitative factors, it is customary to code the levels of each factor so that the minimum level is -1 and the maximum level is 1 . A *center point* is an observation taken at the design point $(x_1, x_2, \dots, x_k) = (0, 0, \dots, 0)$ that is located in the center of the design region. The primary reasons that center points have been included in experimental designs are to provide an estimate of pure experimental error, to allow for tests of **lack of fit** for first-order **response surface** design and for tests of curvature in 2^k designs, and to ensure orthogonal **blocking** and uniform precision for second-order response surface design. Descriptions of these statistical applications of center points will be reviewed. First, several comments related to testing for model lack of fit will be given.

Decisions and predicted responses are based on a fitted model. Therefore, it is important to assess the adequacy of the fitted model because an incorrect model can lead to conclusions that are erroneous, and, in particular, lead to unreliable predicted responses. A statistical test for checking the adequacy of a fitted response surface model is called a *lack-of-fit test* of the fitted model. When we say that a proposed response surface model is “lacking in fit”, we are implying that the model is inadequate because it does not contain enough terms because either (a) there exist one or more factors that are omitted from the proposed model but affect the response of interest or (b) one or more higher-order terms involving factors already in the proposed model are needed to adequately explain the relationship between the response and these factors [1].

The following two design conditions must be met to perform a test for the adequacy of the fitted response surface model (i.e., testing for no lack of fit): (a) the number of distinct design points must exceed the number of terms in the fitted model and (b) the design must contain replicated points that will provide an estimate of the experimental error variance that does not depend on the form of the fitted model. When these two conditions are satisfied, the error **sum of squares** (SSE) can be partitioned into two component sources of variation: the variation associated with the replicated points (referred to as the

pure error sum of squares and denoted SS_{PE}) and the variation attributable to model lack of fit (denoted SS_{LOF}). Thus, $SSE = SS_{PE} + SS_{LOF}$. Although condition (b) can be satisfied by having at least two replicates for one or more design points, it is common to use several center points as the replicates [1].

Center Points and Two-Level Response Surface Designs

Two-level response surface designs are often used when the researcher is interested in fitting either a first-order [2] or interaction [3] response surface model, where

$$y = \beta_0 + \sum_{i=1}^k \beta_i x_i + \epsilon \quad (\text{First-order model}) \quad (1)$$

$$y = \beta_0 + \sum_{i=1}^k \beta_i x_i + \sum_{i < j}^k \beta_{ij} x_i x_j + \epsilon \quad (\text{Interaction model}) \quad (2)$$

Popular two-level designs are the Plackett–Burman designs [4] and the low-resolution 2^{k-p} **fractional factorial** designs [5] for the first-order model and the 2^k full factorial designs [6] and high-resolution 2^{k-p} fractional factorial designs for the second-order model. Note that these designs cannot detect lack of fit due to quadratic x_i^2 terms.

For a 2^k design, let n_f be the number of factorial design points. To test for curvature (i.e. the presence of x_i^2 terms in the true model) n_c center points are included in the design. Let $\bar{y}_f$ and $\bar{y}_c$ be the response means for the factorial and center points, respectively. Thus, if curvature is not present, then a first-order or interaction model may adequately model the response surface, and we would expect that $\bar{y}_f \approx \bar{y}_c$. Thus, if the difference $\bar{y}_f - \bar{y}_c$ is large, we conclude curvature is present.

If the second-order response model is considered as an alternative model, then the null and alternative hypotheses for the test for curvature are $H_0 : \sum_{i=1}^k \beta_{ii} = 0$ and $H_a : \sum_{i=1}^k \beta_{ii} \neq 0$. The sum of squares for curvature based on the contrast $\bar{y}_f - \bar{y}_c$ is $SS_{\text{curv}} = [n_f n_c (\bar{y}_f - \bar{y}_c)^2] / (n_f + n_c)$. Because it has a single degree of freedom, $MS_{\text{curv}} = SS_{\text{curv}}$, and it may be compared to MS_{PE}

2 Center Points

Table 1 A test for curvature

Uncoded		Coded		Yield	$n_f = 4 \quad \bar{y}_f = 44.3$
ξ_1	ξ_2	x_1	x_2	y	
35	150	-1	-1	43.8	
35	160	-1	1	44.6	
45	150	1	-1	43.9	
45	160	1	1	44.9	$n_c = 5 \quad \bar{y}_c = 44.9$
40	155	0	0	44.6, 45.0, 44.8, 45.2, 44.9	

to test for curvature (where the pure error $MS_{PE} = SS_{PE}/(n_c - 1)$ is calculated from the replicated center points). Note that although the test may indicate that curvature is present, we do not have estimates for $\beta_{11}, \dots, \beta_{kk}$. One should augment the design with additional points to allow for **estimation** of these model parameters [7].

Numerical Example suppose the goal is to determine the operating conditions for two process variables (mixing time ξ_1 and oven temperature ξ_2) that maximize the process yield y . The process is currently operating with a mixing time of 40 min and a temperature of 155 °F. It was determined that the experimental region should be from 35 to 45 min for mixing and from 150 to 160 °F for oven temperature. A completely randomized 2^2 design with five center points with process variables mixing time ξ_1 and oven temperature ξ_2 was run. A test of curvature is performed assuming an interaction model. The values of the process yield (y) response are given in Table 1.

Then, $SS_{\text{curv}} = MS_{\text{curv}} = [(4)(5)(44.3 - 44.9)^2]/9 = 0.8$ and $MS_{PE} = 0.05$. Thus, $F = 0.8/0.05 = 16$, and the p value based on a $F(1, 4)$ distribution is 0.0161 suggesting the presence of curvature.

Another important issue regarding two-level designs arises when no design points are replicated. Specifically, the unreplicated design will not produce data that yields an independent estimate of the experimental error variance. For example, when a first-order or interaction model is fitted, then the SSE contains variation attributable to both experimental error and to the lack of fit of that model. Significant lack of fit of a first-order model will occur when the model is fitted over a range of factor levels where terms of higher order are needed to adequately model the response [1]. In the case of a saturated first-order design, it not

possible to test for lack of fit of the first-order model.

The properties of a two-level design, however, may change once it is augmented with replicated design points. For an orthogonal two-level design, an equal number of replications for each design point is needed to preserve the orthogonality property. Unfortunately, the cost of such replication may be prohibitive. A cost-efficient alternative is to augment an unreplicated design with center point replications. This enables the experimenter to obtain an independent estimate of the experimental error while maintaining the orthogonality of the design. It also allows for traditional model lack-of-fit tests to be performed [1]. Estimates resulting from a 2^k design also remain unaffected by the addition of replicates at the center of the design region [7].

Center Points in Second-Order Response Surface Designs

Two of the designs most often used to fit second-order response surface models are the **central composite** designs (CCDs) [8] and the Box–Behnken designs (BBDs) [9]. CCD points are composed of factorial points from a 2^k or 2^{k-p} design, axial points, and center points, while BBD points are constructed from **incomplete block** designs [10] to which center points are added. Information regarding the use of center points in second-order response surface designs can be found in [11–14].

The CCD is a very efficient design for fitting the second-order response surface model. The factorial points form a design that minimizes the variances of the estimated coefficients in a first-order or interaction model. The center points provide information about the existence of curvature, and if curvature is

present they also contribute toward the estimation of quadratic terms.

For a k -factor CCD, the experimenter has the flexibility to control both the axial point distance α and the number of center points n_c . While choice of α depends to a great extent on the regions of operability and interest, the choice of n_c affects prediction variances in the region of interest. If $\alpha = \sqrt{k}$ and $n_c = 0$, then the CCD would be a spherical design (i.e., all design points are equidistant from the center of the design region). The experimental design matrix would be singular, meaning the second-order model parameters cannot be estimated. Thus, n_c must be ≥ 1 when $\alpha = \sqrt{k}$ to avoid the singularity and allow estimation of all model parameters. In addition, the use of center points also provides reasonable stability for the prediction variance function in the design region. Hence, inclusion of center points in a CCD is very beneficial. Options regarding n_c and α were studied in [11], and it was concluded that the best CCDs should have $\alpha \approx \sqrt{k}$ and $2 \leq n_c \leq 5$ [7].

The spherical nature of the BBD, combined with the fact that the designs are rotatable or near rotatable [15], suggests that center points should be included in the design. When $k = 4$ or 7, BBDs are spherical designs making center points necessary to avoid singularity of the design matrix. Spherical or nearly spherical designs require three to five center points to avoid a severe imbalance in the prediction variance through the design region.

In general, in an unblocked response surface design, the number of center points can be used to control various properties of the design matrix X . Under certain conditions, the number of center points can make X have uniform precision (or have near-uniform precision). A design has *uniform precision* if the prediction variance at the center of the experimental region and at a unit distance away from the center are equal.

The number of center points can also be chosen to cause a rotatable CCD to have the uniform (or near-uniform) precision property. Setting n_c to be the integer closest to $\lambda_4(\sqrt{F} + 2)^2 - F - 2k$ yields a uniform precision rotatable design. For an N -point CCD, λ_4 is the constant mixed fourth-order design moment $(1/N) \sum_{u=1}^N x_{ui}^2 x_{uj}^2$ for any pair of design factors $i \neq j$. Values of λ_4 can be found in [1]. If it is desired for a CCD to be orthogonal as well as rotatable, it is possible to choose α and n_c to

achieve both properties. To achieve the rotatability, we only need to choose α such that $\alpha^4 = F$ where F is the number of factorial points in the CCD. To achieve orthogonality, n_c is the integer closest to $4\sqrt{F} + 4 - 2k$ [1].

There are many practical situations in which the experimenter specifies strict minimum and maximum levels for the design factors. That is, the region of interest and the region of operability are the same, and the resulting design region is a cube. By setting $\alpha = 1$, we have a CCD in the cube which is referred to as a face-center cube design (FCCD). The recommendations for the number of center points differ between FCCDs and CCDs in spherical design regions. For a FCCD, only one or two center points are sufficient to provide reasonable stability to the prediction variance. Thus, the stability of the prediction variance is much more sensitive for a CCD in a spherical region than for an FCCD in a cuboidal region [7].

References

- [1] Khuri, A.I. & Cornell, J.A. (1996). *Response Surfaces: Designs and Analyses*, Marcel Dekker, New York, pp. 38–42, 54–56, 96–97, 113, 163–164.
- [2] Jacroux, M. (2007). Main effect plans, *Encyclopedia of Quality and Reliability*, John Wiley & Sons, Chichester, pp. xxx–xxx.
- [3] Reese, S. (2007). Interactions, *Encyclopedia of Quality and Reliability*, John Wiley & Sons, Chichester, pp. xxx–xxx.
- [4] Tyssedal, J. (2007). Plackett-Burman designs, *Encyclopedia of Quality and Reliability*, John Wiley & Sons, Chichester, pp. xxx–xxx.
- [5] Kenett, R. (2007). Factorial designs, *Encyclopedia of Quality and Reliability*, John Wiley & Sons, Chichester, pp. xxx–xxx.
- [6] Voelkel, J. (2007). Fractional-factorial designs, *Encyclopedia of Quality and Reliability*, John Wiley & Sons, Chichester, pp. xxx–xxx.
- [7] Myers, R.H. & Montgomery, D.A. (2002). *Response Surface Methodology: Process and Product Optimization Using Designed Experiments*, John Wiley & Sons, New York, pp. 128–130, 321–323, 332–337.
- [8] Draper, N. (2007). Central composite designs, *Encyclopedia of Quality and Reliability*, John Wiley & Sons, Chichester, pp. xxx–xxx.
- [9] Robinson, T.J. (2007). Box-Behnken designs, *Encyclopedia of Quality and Reliability*, John Wiley & Sons, Chichester, pp. xxx–xxx.
- [10] Prescott, P. (2007). Incomplete block designs, *Encyclopedia of Quality and Reliability*, John Wiley & Sons, Chichester, pp. xxx–xxx.

4 Center Points

- [11] Myers, R.H., Vining, G.G., Giovannitti-Jensen, A. & Myers, S.L. (1992). Variance dispersion properties of second order response surface designs, *Journal of Quality Technology* **24**, 1–11.
- [12] Draper, N.R. (1982). Center points in second-order response surface designs, *Technometrics* **24**, 127–133.
- [13] Giovannitti-Jensen, A. & Myers, R.H. (1989). Graphical assessment of the prediction capability of response surface designs, *Technometrics* **31**, 159–171.
- [14] Borkowski, J.J. (1995). Spherical prediction variance properties of central composite and Box-Behnken designs, *Technometrics* **37**, 399–410.
- [15] Draper, N. (2007). Rotatable designs and rotatability, *Encyclopedia of Quality and Reliability*, John Wiley & Sons, Chichester, pp. xxx–xxx.

JOHN J. BORKOWSKI

Central Composite Designs

Fitting a Response Surface to Data

Suppose an experimenter wishes to examine several predictor variables $x_1, x_2, \dots, x_k$, in order to determine their effects, individual or combined, on a response variable y . A standard way of proceeding is to decide on n (say) specific combinations of levels of the x 's ($x_{1u}, x_{2u}, \dots, x_{ku}$), $u = 1, 2, \dots, n$; at each of these, the value of the response y (y_u for the u th experiment) will be measured. These n experiments, or *runs* as they are usually called, form the *experimental design* or simply the *design*. Further analysis of the experimental results arising from performing the design usually consists of fitting a model, often of polynomial form, to the data, using the least-squares method of regression analysis, and then assessing what the computations reveal. The fitted model could be first order (planar), for example:

$$\hat{y} = b_0 + b_1x_1 + b_2x_2 + \dots + b_kx_k \quad (1)$$

where each x has only a linear effect; or could be second order (quadratic)

$$\begin{aligned} \hat{y} = & b_0 + b_1x_1 + b_2x_2 + \dots + b_kx_k \\ & + b_{11}x_1^2 + b_{22}x_2^2 + \dots + b_{kk}x_k^2 \\ & + b_{12}x_1x_2 + b_{13}x_1x_3 + \dots + b_{k-1,k}x_{k-1}x_k \end{aligned} \quad (2)$$

where all possible second-order terms (power 2, squares and cross-products) have been added to equation (1). These equations are called *fitted equations*, *fitted models*, or *fitted surfaces*. The *hat* (caret) over the y indicates the *predicted value* of y . The b 's are numbers, estimates obtained via the least-squares method, of the *true but unknown* regression coefficients. The fitted models provide (in many cases) a plausible representation of, or an approximation to, the true surface. The models shown above simply represent first- and second-order Taylor's series expansions of a general functional form and are thus expected to embody the main features of the variations in the response. This is often true in practice. The first step in the process is thus to decide what sort of design should be employed (see **Response Surface Methodology**.)

First-order Designs

The simplest good choice to fit a first-order model, equation (1), would be a full two-level factorial or a fractional factorial type design of resolution III or IV (see **Factorial Experiments; Fractional Factorial Designs; Orthogonal Arrays**). A resolution III design would associate (some or all) main effects with two-factor interactions; a resolution IV design would associate (some or all) main effects with three-factor interactions and two-factor interactions with other two-factor interactions. The philosophy here would be that only the main effects (the b_i) were likely to be important in the exploratory phases of an investigation.

Coding the Experimental Conditions

Two levels of a predictor variable can always be coded to -1 and $+1$. For example, if two temperatures $T = 125$ and 175°C were to be used, the coding $(T - 150)/25$ would achieve the coded levels. Note that the coded level 0 (here 150°C) would be the center of the range selected. Often, the center of the entire design would represent the *current best set of conditions within the experimental region* and the experiment would attempt to find the best manner of moving outward to improve the response. With two variables, temperature and pressure say, the coded levels of a 2^2 design would be $(x_1, x_2) = (-1, -1), (-1, 1), (1, -1),$ and $(1, 1)$ and this design would provide enough information to fit a first-order model (a plane) in the two-dimensional experimental space. If the first-order surface were deemed to be adequate, subsequent follow-up analysis would usually consist of following the path of steepest ascent, that is, the straight line in the experimental space in which the predicted response improves the fastest. Often such an analysis would provide rapid improvement in the response and allow the experimenter to move to a better region of the experimental space.

The Central Composite Second-order Design

To fit a second-order model, at least three levels of each x_i must be chosen to permit detection of curvature in the response. A popular and practical

2 Central Composite Designs

design choice, which can be found in many published experimental presentations, is the so-called central (i.e., centrally located about the coded design origin) composite design, which provides three, four, or five levels. A composite design consists of three parts:

1. A “cube” part, that is, a two-level n_c points portion, $(\pm 1, \pm 1, \dots, \pm 1)$, usually of resolution IV or higher (which means that no **main effect** is aliased with any two-factor interaction); plus
2. a “star” part, consisting of $2k$ axial points of the forms $(\pm\alpha, 0, 0, \dots, 0)$, $(0, \pm\alpha, 0, \dots, 0)$, $(0, 0, \pm\alpha, 0, \dots, 0)$, $\dots$, $(0, 0, 0, \dots, 0, \pm\alpha)$, where α will be chosen to match the experimental goals; plus
3. n_0 center points $(0, 0, \dots, 0)$ (*see Center Points*).

In general, this design uses five levels of each predictor variable. Without center points, there are only four levels $(-\alpha, -1, 1, \alpha)$. The omission of center points is unusual, because there would then be no comparison with the current best levels mentioned earlier. When $\alpha = 1$, the general design can be described as a *face-centered central composite design* and has only three levels $(-1, 0, 1)$. In this case, center points must be included so that the full second-order model can be fitted.

Figure 1, reproduced from [1, Chapter 9], shows a composite design in $k = 3$ dimensions with $n_0 = 2$ center points and $\alpha = 2$; response values are shown attached at the various experimental locations. Each factor is at five levels, namely $(-2, -1, 0, 1, 2)$. The statistical analysis of this particular example is described below.

Choice of the Star Distance α

The value of α selected for a particular investigation will depend on the circumstances. Clearly, the width 2α must lie within the chosen experimental range of the factors. A good general question to ask first in considering a choice for α is “What values of α provide design points that cover the experimental region in a way that makes the experimenter feel comfortable?”

By an appropriate choice of α , any second-order central composite design can be made *rotatable*, which means that, in the coded coordinates, information obtained from the design, as measured by $V(\hat{y})$, the variance of the fitted value function, is equally

accurate at all points that lie the same distance from the center of the design (*see Rotatable Designs and Rotatability* for more detail). Rotatability is desirable but not mandatory. Nevertheless the rotatable α values would be useful ones to keep in mind as a starting point. The appropriate values for a selected “single cube plus single star” central composite design are shown in Table 1.

The asterisk indicates that a half fraction of the 2^k design is being used. Adding center points does not affect rotatability.

It would be unusual to choose $\alpha < 1$, unless the factorial levels had been set far too wide in the cube portion. Moreover, such a narrow choice would tend to make the design “less rotatable”.

A composite design is often performed sequentially in two parts. The first part would usually consist of the cube portion plus some center points, thus enabling a first-order model to be fitted and also providing a test for lack of planar fit by comparing the average response at the center point with the average response at the factorial points and dividing by the appropriate **standard error** to create a testable *t*-statistic. A significant value indicates the need to fit a second-order model. Addition of the star portion plus (often) further center points would allow a full second-order model to be fitted using the entire set of data. In such cases, the addition to the model of a *blocking variable* might be desirable if it were felt that the two parts of the experimental design would each occur at a different overall level of response (*see Blocking*). This would be a sensible precaution if the two parts of the design were performed in different weeks, or the operators were different, or two individual reactors were used. Note that performing center points in *both* blocks and comparing the two corresponding sets of response values gives *direct* information on the possible differences between blocks.

Another possible choice of α would be the value that provides orthogonal blocking when the design is performed in more than one block. This is somewhat difficult to explain briefly. One would first need

Table 1 Values of α that make single cube plus single star second-order central composite designs rotatable

k :	2	3	4	5	5*	6	6*
n_c :	4	8	16	32	16	64	32
α :	1.414	1.682	2	2.378	2	2.828	2.378

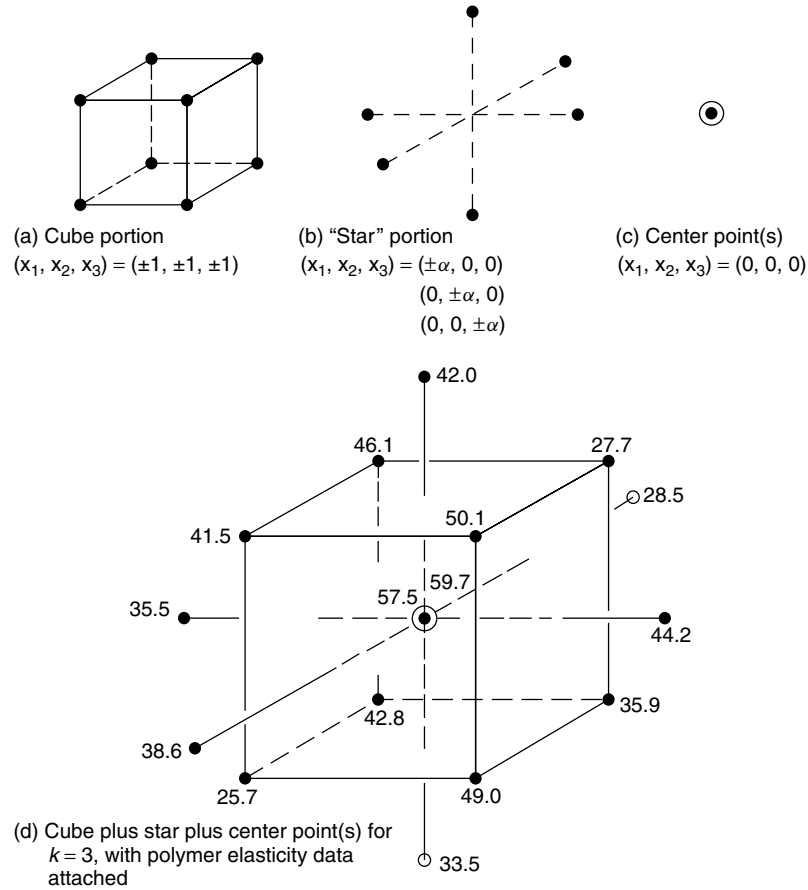


Figure 1 A central composite design in three dimensions ($k = 3$). (a) Cube portion $(\pm 1, \pm 1, \pm 1)$. (b) "Star" portion $(\pm \alpha, 0, 0)$, $(0, \pm \alpha, 0)$, $(0, 0, \pm \alpha)$ with $\alpha = 2$. (c) Center points $(0, 0, 0)$. (d) Cube, plus star, plus two center points with response data values attached

to decide how many blocks to use, then which runs to put in which block, and then how the levels for the blocking variable columns would be chosen. The worked example below has an α value of 2 that achieves orthogonal blocking; the blocking variable column used there, with -1 's and 1 's, is orthogonal to all the columns in the $\mathbf{X}$ matrix generated by the second-order model. (This design is thus not rotatable, which would require $\alpha = 1.682$.) An example in which orthogonal blocking was achieved by a choice of $\alpha = 1.2872$ is mentioned in [2, p. 290]. Second-order central composite designs have been in popular use for over half a century [3], and for excellent reasons. Their theoretical and practical properties make them rank high in comparison with other

suggested designs, using several different criteria (*see Optimal Design*).

Analysis of the Data

The particular experiment of our example was performed sequentially, but not exactly in the way described above. There *were* two blocks, performed in different weeks, but the first consisted of just the cube portion with no center points added, while the second block consisted of the star, at a coded distance of two units from the origin, plus two center points. This had the advantage of providing the same number of runs, namely eight, per block but failed to provide a direct block-to-block center points comparison

4 Central Composite Designs

Table 2 First (runs 1–8) and second (runs 9–16) blocks of experimental data

Number	x_1	x_2	x_3	x_B	y
1	−1	−1	−1	−1	25.74
2	1	−1	−1	−1	48.98
3	−1	1	−1	−1	42.78
4	1	1	−1	−1	35.94
5	−1	−1	1	−1	41.50
6	1	−1	1	−1	50.10
7	−1	1	1	−1	46.06
8	1	1	1	−1	27.70
9	0	0	0	1	57.52
10	0	0	0	1	59.68
11	−2	0	0	1	35.50
12	2	0	0	1	44.18
13	0	−2	0	1	38.58
14	0	2	0	1	28.46
15	0	0	−2	1	33.50
16	0	0	2	1	42.02

at the same levels. The planar model (1) was fitted to the first set of data when it was collected (see observations 1–8 in Table 2), giving rise to the fitted equation (3).

$$\hat{y} = 39.85 + 0.83x_1 - 1.73x_2 + 1.49x_3 \quad (3)$$

The four estimated coefficients 39.85, 0.83, 1.73, and 1.49 all have standard errors of 4.16 based on the residual mean square of $s^2 = 138.3$, revealing a puzzling picture of a very significant mean level and non-significant first-order effects. Two statistical features provide added clarification. If the interaction terms in x_1x_2 , x_1x_3 , x_2x_3 are added to the model, all of their estimated coefficients (values -7.13 , -3.27 , -2.73 respectively) are much larger than the linear effects and, moreover, the estimate 39.85 in equation (3) is actually an estimate of $\beta_0 + \beta_{11} + \beta_{22} + \beta_{33}$ where $\beta_{11}x_1^2 + \beta_{22}x_2^2 + \beta_{33}x_3^2$ are the *pure quadratic* terms that would be added to the model if a full second-order (quadratic) model were appropriate. At the present stage of analysis, we do not have enough data to estimate a full second-order model, however. A logical next step then is to add a “star” set of points (six points for three dimensions) plus a couple of center points, as mentioned earlier. The entire set of data appears in Table 2, together with the addition of a blocking variable denoted by x_B . The latter is allotted the values -1 for the first block of data and 1 for the second block. Actually,

any two distinct numbers would suffice, for example $(0, 1)$, but the specific choice of $(-1, 1)$ has the advantage, in this particular design, of making the blocks column orthogonal to all the other columns needed to fit a full second-order model, which means that the blocking effect can be cleanly separated from the estimated model effects, both numerically and in the subsequent analysis of variance table. Such orthogonality is worth seeking and is often achievable, depending on the specific design used (see **Blocking**).

The second-order model fitted to the complete set of 16 runs takes the form

$$\begin{aligned} \hat{y} = & 57.31 + 1.50x_1 - 2.13x_2 + 1.81x_3 \\ & - 4.69x_1^2 - 6.27x_2^2 - 5.21x_3^2 - 7.13x_1x_2 \\ & - 3.27x_1x_3 - 2.73x_2x_3 + 1.29x_B \end{aligned} \quad (4)$$

and the corresponding analysis of variance table appears as Table 3.

Viewing the entire set of results, we conclude that there is no **lack of fit**, that it is useful to remove a block component of the variation, and that the full second-order model is needed to provide a useful representation of the fitted second-order surface. A “ridge analysis” of the fitted surface can be found in (**Ridge Analysis in Experimental Design**).

Central Composite Designs for Higher Dimensions

For higher dimensions, $k \geq 5$, the number of runs needed for a central composite design can be reduced by choosing the “cube” part to be a *fractional* factorial rather than a full factorial. In order that

Table 3 The analysis of variance table for the experiment

Source of variation	df	SS	MS	$F^{(a)}$
b_0	1	27079.99	—	—
Blocks, orthogonal to model	1	26.63	26.63	8.94
First order, given b_0	3	161.01	53.67	18.01
Second order, given first, b_0	6	1290.61	215.10	72.18
Lack of fit	4	12.60	3.20	1.37
Pure error	1	2.33	2.33	—
Total	16	28573.17	—	—

^(a) Using $s^2 = (12.60 + 2.33)/(4 + 1) = 2.98$ as denominator except for lack of fit F test

all main effects and two-factor interactions (essentially the b_i and the b_{ij}) may be separately estimated, the fractional factorial must usually be of resolution V or higher. Resolution V ensures that main effects are aliased, or associated, with four-factor interactions, at worst, and that two-factor interactions are aliased or associated with three-factor interactions, at worst (see **Fractional Factorial Designs**). In much experimental work, it is typically anticipated (initially, at least, until evidence arises to the contrary) that three- and four-factor interactions are very unlikely to occur. A regression-type extra **sum-of-squares**, lack-of-fit test will allow these sorts of assumptions to be tested; further runs can then be added if required to estimate additional parameters. For tables of central composite designs, see, for example [[1], Chapter 15], [[4], pp. 508–517], and [[5], p. 431]. For more illustrative worked examples, see [[1], Exercises 11.5–11.7, 11.9, 11.16, and 11.18–11.20 together with their corresponding solutions].

Other Composite Design Variations

Other variations on composite designs are the following:

1. Central composite designs that consist of more than two pieces, for example, one cube plus two stars. The two stars could be a doubled star, sharing the same α value, or they could have different α values. Or the cube could be doubled, perhaps using different sizes. The coded dimensions could be chosen to provide rotatability and/or orthogonal blocking, as feasible, and so on.
2. Noncentral composite designs, which add supplementary runs to the cube portion in particular directions in the experimental space, directions chosen by the experimenter perhaps after viewing the data from the cube portion or because of known restrictions in the omitted directions. The idea is to supplement the cube information by working in an interesting and favorable direction.
3. Small composite designs, in which the basic cube portion is fractionated to the point where some of the two-factor interactions are aliased with one another. These would be untangled using the

information obtained from the star portion (see, e.g., [1, Chapter 15]).

Special second-order designs, some of them rotatable, others not, also exist for situations where only three levels (coded $-1, 0, 1$) are possible (see [6], [1, Chapter 15], and **Box–Behnken Designs**). For examples of the use of central composite designs in split plot situations, see [7, p. 170].

Subsequent Analysis

If a second-order surface fits the data well, further analysis could take several forms such as plotting the surface represented by equation (2) if k were small ($k = 2$ or 3), plotting two-dimensional or three-dimensional sections if k were larger; performing canonical reduction, a technique that determines the principal features of the fitted surface for any value of k ; or performing ridge analysis (see **Ridge Analysis in Experimental Design**), which enables the following of a “best path” on the surface in any number of dimensions k , beginning at *any* selected point in the experimental space. In some cases, that selected point might simply be the coded design origin. Higher, lower, or intermediate local optima values of the response can be specified as desirable in the ridge analysis, and it is also possible to track the ridge path onto any designated linear restrictions (restrictive planes in the experimental space) that may be encountered by the initial best path. One such type of linear restriction occurs when the experiment is conducted in mixture ingredients adding to 1 (or any other level specified between 0 and 1) (see [1, Chapters 12 and 19] and **Mixture Experiments**).

It is also possible to fit many other types of special models if they make sense in the specific context. For example, a specific alternative mechanistic model (one that reflects the actual mechanism of the process) might be chosen by an experimenter who has the requisite applicable theoretical knowledge of the experimental mechanism under study (see [4, Chapter 12] and [8]).

References

- [1] Box, G.E.P. & Draper, N.R. (2007). *Response Surfaces, Mixtures and Ridge Analyses*, 2nd Edition, John Wiley & Sons, New York.

- [2] Woods, D.C., Lewis, S.M., Eccleston, J.A. & Russell, K.G. (2006). Designs for generalized linear models with several variables and uncertainty, *Technometrics* **2**, 284–292.
- [3] Box, G.E.P. & Wilson, K.B. (1951). On the experimental attainment of optimum conditions, *Journal of the Royal Statistical Society. Series B* **13**, 1–38; discussion 38–45.
- [4] Box, G.E.P. & Draper, N.R. (1987). *Empirical Model-Building and Response Surfaces*, John Wiley & Sons, New York.
- [5] Wu, C.F.J. & Hamada, M. (2000). *Experiments, Planning, Analysis, and Parameter Design Optimization*, John Wiley & Sons, New York.
- [6] Box, G.E.P. & Behnken, D.W. (1960). Some new three-level designs for the study of quantitative variables, *Technometrics* **2**, 455–475; Corrections, 1961, **3**, 576.
- [7] Goos, P., Langhans, I. & Vandebroek, M. (2006). Practical inference from industrial split plot designs, *Journal of Quality Technology* **2**, 162–179.
- [8] Bates, D.M. & Watts, D.G. (1988). *Nonlinear Regression Analysis and its Applications*, John Wiley & Sons, New York.

Related Articles

Blocking; Box–Behnken Designs; Center Points; Optimal Design; Factorial Experiments; Fractional Factorial Designs; Mixture Experiments; Orthogonal Arrays; Response Surface Methodology; Ridge Analysis in Experimental Design; Rotatable Designs and Rotatability.

NORMAN R. DRAPER

Mixture Experiments

Introduction

A *mixture experiment* involves combining two or more components in various proportions or amounts and then measuring one or more responses for the resulting end products. Other factors that affect the response(s), such as process variables (PVs) and/or the total amount of the mixture, may also be studied in the experiment. A *mixture experiment design* specifies the combinations of mixture components and other experimental factors (if any) to be studied and the response(s) to be measured. *Mixture experiment data analyses* are then used to achieve the desired goals, which may include (a) understanding the effects of components and other factors on the response(s), (b) identifying components and other factors with significant and nonsignificant effects on the response(s), (c) developing models for predicting the response(s) as functions of the mixture components and any other factors, and (d) developing end products with optimal or desired values and uncertainties of the response(s).

Given a mixture experiment problem, a practitioner must consider the possible approaches for designing the experiment and analyzing the data, and then select the approach best suited to the problem. Eight possible approaches are (a) component proportions (CPs), (b) mathematically independent variables (MIVs), (c) slack variable (SV), (d) mixture amount (MA), (e) component amounts (CA), (f) mixture-process variable (MPV), (g) mixture of mixtures, (MoM) and (h) multifactor mixture (MFM). The article provides an overview of the mixture experiment designs, models, and data analyses appropriate for these approaches. For more discussion and examples of these approaches, see Piepel and Cornell [1]. Also, the books by Cornell [2] and Smith [3] are devoted to the design and data analysis of mixture experiments and are recommended for additional information. Piepel [4] summarizes 50 years of mixture experiment research from 1955 to 2004 along with an extensive bibliography.

Before designing a mixture experiment, a practitioner must choose the model form(s) likely to be adequate to achieve the objectives of the experiment, and then design the experiment to adequately support fitting the model form(s). Hence in

the subsequent section for each mixture experiment approach, appropriate models are discussed before experimental designs.

Component Proportions Approach

With the CP approach for mixture experiments, q mixture components and no other factors are studied. The experiment is designed and responses are modeled using the proportions of the mixture components x_i , which are subject to the constraints

$$0 \leq x_i \leq 1, i = 1, 2, \dots, q \text{ and } \sum_{i=1}^q x_i = 1 \quad (1)$$

The experimental region defined by these constraints is a $(q - 1)$ -dimensional simplex. Additional single-component constraints of the form

$$L_i \leq x_i \leq U_i, \quad i = 1, 2, \dots, q \quad (2)$$

and/or multiple-component constraints of the form

$$C_k \leq \sum_{i=1}^q A_{ki}x_i \leq D_k, \quad k = 1, 2, \dots, K \quad (3)$$

generally yield a $(q - 1)$ -dimensional polyhedral experimental region. The constraint in equation (1), that the mixture CPs must sum to one, means that the mixture CPs cannot be varied independently. This fact affects the choice of experimental designs and models for the CP approach (and some other approaches) to mixture experiments.

Scheffé Polynomial Models in Component Proportions

Polynomial models are often successful in modeling response surfaces based on Taylor series theory for approximating functions with polynomial expansions. However, the standard polynomial models must be modified to account for the mixture constraint $\sum_{i=1}^q x_i = 1$ in equation (1). Several commonly used Scheffé [5] polynomial models for CP mixture experiments are

Linear:

$$\eta(\mathbf{x}) = \sum_{i=1}^q \beta_i x_i \quad (4)$$

2 Mixture Experiments

Quadratic:

$$\eta(\mathbf{x}) = \sum_{i=1}^q \beta_i x_i + \sum_{i=1}^{q-1} \sum_{j=i+1}^q \beta_{ij} x_i x_j \quad (5)$$

Special cubic:

$$\begin{aligned} \eta(\mathbf{x}) = & \sum_{i=1}^q \beta_i x_i + \sum_{i=1}^{q-1} \sum_{j=i+1}^q \beta_{ij} x_i x_j \\ & + \sum_{i=1}^{q-2} \sum_{j=i+1}^{q-1} \sum_{k=j+1}^q \beta_{ijk} x_i x_j x_k \end{aligned} \quad (6)$$

Full cubic:

$$\begin{aligned} \eta(\mathbf{x}) = & \sum_{i=1}^q \beta_i x_i + \sum_{i=1}^{q-1} \sum_{j=i+1}^q \beta_{ij} x_i x_j \\ & + \sum_{i=1}^{q-1} \sum_{j=i+1}^q \delta_{ij} x_i x_j (x_i - x_j) \\ & + \sum_{i=1}^{q-2} \sum_{j=i+1}^{q-1} \sum_{k=j+1}^q \beta_{ijk} x_i x_j x_k \end{aligned} \quad (7)$$

Full-quartic and special-quartic models are discussed by Smith [3, Section 3.4]. In all of these model forms, $\eta(\mathbf{x})$ is the expected value of the response at mixture $\mathbf{x} = (x_1, x_2, \dots, x_q)$. The $\sum_{i=1}^q \beta_i x_i$ model terms in equations (4–7) describe *linear blending* of the mixture components, while the remaining terms in equations (5–7) describe *nonlinear blending* of the mixture components in various degrees (i.e., quadratic, special cubic, and full cubic). The nonlinear blending terms are not referred to as *interaction terms*, because curvature and interaction effects of the mixture components are partially confounded as a result of the mixture constraint $\sum_{i=1}^q x_i = 1$. Methods for investigating curvature and interaction effects of mixture components are discussed by Cox [6] and Piepel *et al.* [7].

The coefficient β_i in equations (4–7) represents the expected value of the response at the pure mixture with $x_i = 1$ and $x_j = 0$ for $j \neq i$. Cornell [2, Sections 2.2 and 2.3] discusses the interpretations of other coefficients in mixture experiment models. However, all such interpretations are extrapolations when the experimental region is restricted by constraints of the form in equations (2) and/or (3).

Other Polynomial Models in Component Proportions

Cox [6] proposed using ordinary first- and second-degree polynomial models with the parameters constrained to estimate the main, curvature, and interactive effects of components in a way appropriate to mixture experiments. Cornell [2, Sections 6.7 and 6.8] and Smith and Beverly [8] discuss and illustrate the *Cox mixture polynomial models*. Piepel [9, 10] proposed *component-slope linear mixture (CSLM) models* with parameters representing the slopes of the **response surface** along the component effect directions. The advantage of the Cox and CSLM models over the Scheffé models is that their coefficients provide estimates of component effects. The disadvantage is that the models are overparameterized and involve constraints on the coefficients.

Scheffé quadratic models in equation (5) are adequate for many mixture experiment applications, but in practice only a few quadratic terms may be needed. *Partial quadratic mixture (PQM)* models of the form

$$\eta(\mathbf{x}) = \sum_{i=1}^q \beta_i x_i + \text{subset of} \left\{ \sum_{i=1}^q \beta_{ii} x_i^2 + \sum_{i=1}^{q-1} \sum_{j=i+1}^q \beta_{ij} x_i x_j \right\} \quad (8)$$

contain a subset of squared terms and/or cross-product terms. PQM models contain as a subset the set of reduced Scheffé quadratic models, and thus can more efficiently capture curvature as well as interactive effects of mixture components. See Piepel *et al.* [11] and Smith [3, Section 10.3] for discussion and illustrations of PQM models.

Nonpolynomial Models in Component Proportions

The mixture polynomial model forms previously introduced provide for adequately approximating many true response surfaces for mixture experiments, but there are situations when nonpolynomial model forms are more appropriate.

For CP mixture experiments where one or more components have additive effects, three *models homogeneous of degree one* proposed by Becker [12] are

appropriate

$$\begin{aligned} \text{H1: } \eta(\mathbf{x}) = & \sum_{i=1}^q \beta_i x_i + \sum_{i=1}^{q-1} \sum_{j=i+1}^q \beta_{ij} \min(x_i, x_j) \\ & + \sum_{i=1}^{q-2} \sum_{j=i+1}^{q-1} \sum_{k=j+1}^q \beta_{ijk} \min(x_i, x_j, x_k) \quad (9) \end{aligned}$$

$$\begin{aligned} \text{H2: } \eta(\mathbf{x}) = & \sum_{i=1}^q \beta_i x_i + \sum_{i=1}^{q-1} \sum_{j=i+1}^q \beta_{ij} \frac{x_i x_j}{x_i + x_j} \\ & + \sum_{i=1}^{q-2} \sum_{j=i+1}^{q-1} \sum_{k=j+1}^q \beta_{ijk} \frac{x_i x_j x_k}{x_i + x_j + x_k} \quad (10) \end{aligned}$$

$$\begin{aligned} \text{H3: } \eta(\mathbf{x}) = & \sum_{i=1}^q \beta_i x_i + \sum_{i=1}^{q-1} \sum_{j=i+1}^q \beta_{ij} (x_i x_j)^{1/2} \\ & + \sum_{i=1}^{q-2} \sum_{j=i+1}^{q-1} \sum_{k=j+1}^q \beta_{ijk} (x_i x_j x_k)^{1/3} \quad (11) \end{aligned}$$

The cubic forms of these models are shown, but the quadratic forms (without the last group of terms in each model) are more frequently used in practice. For additional discussion of these models, see Cornell [2, Sections 6.3–6.5] and Snee [13].

For CP mixture experiments where the response variable increases rapidly as CPs approach zero, Scheffé linear and quadratic models with inverse terms

$$\eta(\mathbf{x}) = \sum_{i=1}^q \beta_i x_i + \sum_{i=1}^q \delta_i x_i^{-1} \quad (12)$$

$$\eta(\mathbf{x}) = \sum_{i=1}^q \beta_i x_i + \sum_{i=1}^{q-1} \sum_{j=i+1}^q \beta_{ij} x_i x_j + \sum_{i=1}^q \delta_i x_i^{-1} \quad (13)$$

are appropriate. Inverse terms can be added to other mixture models such as in equations (6) and (7) if needed. It is only necessary to add inverse terms for those components needing them. For CPs that can equal zero, the corresponding x_i^{-1} in equations (12) and (13) are replaced with $(x_i + c_i)^{-1}$, where the c_i values are small positive constants. The forms of Scheffé models with inverse terms can also be modified for use when the response increases rapidly as constrained CPs approach their lower limits. See Cornell [2, Sections 6.1 and 6.2] for further discussion of these topics.

Log-contrast models were introduced by Aitchison and Bacon-Shone [14] for situations in which one or more mixture components are inert or have additive effects on the response. The linear and quadratic forms of the log-contrast models are given by

$$\eta(\mathbf{x}) = \gamma_0 + \sum_{i=1}^q \gamma_i \log x_i \quad (14)$$

$$\begin{aligned} \eta(\mathbf{x}) = & \gamma_0 + \sum_{i=1}^q \gamma_i \log x_i \\ & + \sum_{i=1}^{q-1} \sum_{j=i+1}^q \gamma_{ij} (\log x_i - \log x_j)^2 \quad (15) \end{aligned}$$

which are overparameterized and thus subject to constraints on the parameters. For CPs that can equal zero, the corresponding $\log x_i$ in equations (14) and (15) are replaced with $\log(x_i + c_i)$. See Aitchison and Bacon-Shone [14] and Cornell [2, Sections 6.9 and 6.10] for additional discussion of the constraints on the parameters, interpretations of the parameters, and examples of these models.

Simplex Designs in Component Proportions

Four types of experimental design are commonly used to explore how responses depend on the proportions of mixture components over simplex experimental regions (either the whole simplex region as defined in equation (1) or a simplex-shaped subregion).

A $\{q, m\}$ simplex-lattice design consists of $\binom{q+m-1}{m} = (q+m-1)!/m!(q-1)!$ points with each component having $m+1$ equally spaced values of $0, 1/m, 2/m, \dots, 1$. For example, the $\{3, 2\}$ simplex-lattice design contains the six points $(x_1, x_2, x_3) = (1, 0, 0), (0, 1, 0), (0, 0, 1), (1/2, 1/2, 0), (1/2, 0, 1/2),$ and $(0, 1/2, 1/2)$. This design is pictured in Figure 1(a). As another example, the $\{4, 3\}$ simplex-lattice design contains 20 points including 4 permutations of $(1, 0, 0, 0)$, 12 permutations of $(1/3, 2/3, 0, 0)$, and 4 permutations of $(1/3, 1/3, 1/3, 0)$. A $\{q, m\}$ simplex-lattice design supports fitting a Scheffé polynomial model of degree m (see Cornell [2, Section 2.2]). For example, the $\{3, 3\}$ simplex-lattice design supports fitting the full-cubic model in equation (7).

A *simplex-centroid design* for q components contains $2^q - 1$ points, consisting of all permutations

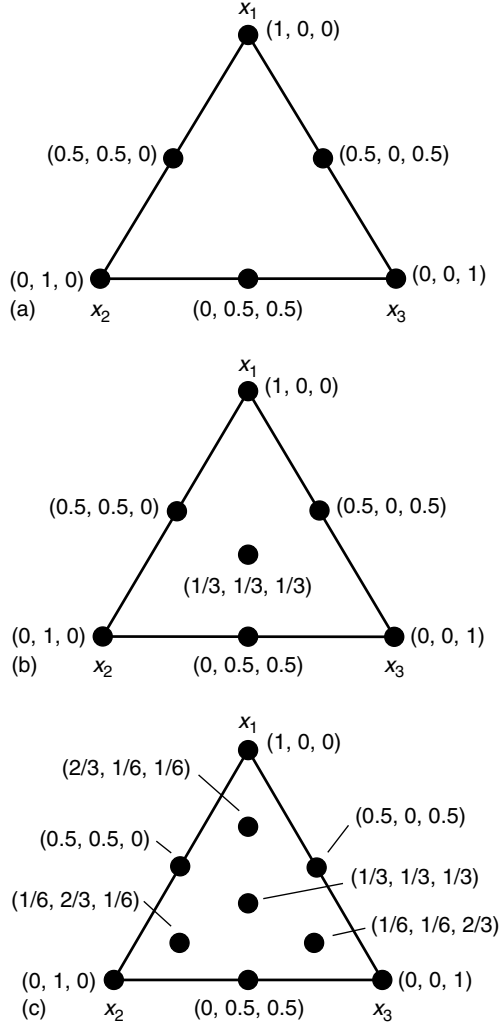


Figure 1 Example designs for three mixture components with the CP approach: (a) {3, 2} simplex-lattice design, (b) simplex-centroid design, and (c) augmented simplex-centroid design

of $(1, 0, 0, \dots, 0)$, $(1/2, 1/2, 0, \dots, 0)$, $(1/3, 1/3, 1/3, 0, \dots, 0)$, ... as well as the centroid $(1/q, 1/q, 1/q, \dots, 1/q)$. Figure 1(b) illustrates this design for $q = 3$. The simplex-centroid design supports fitting the full Scheffé polynomial model in equation (4), as well as the Scheffé linear, quadratic, and special-cubic polynomial models in equations (5–7).

An *augmented simplex-centroid design* adds to the simplex-centroid design the q interior axial points of

the form $((q+1)/2q, 1/2q, 1/2q, \dots, 1/2q)$. This design for $q = 3$ is shown in Figure 1(c).

Snee and Marquardt [15] proposed a *simplex screening design* for q components that contains points along the component axes, including the q vertices of the form $(1, 0, 0, \dots, 0)$, the q axis endpoints of the form $(0, 1/q, \dots, 1/q)$, the q interior axial points of the form $((q+1)/2q, 1/2q, 1/2q, \dots, 1/2q)$, and the overall centroid $(1/q, 1/q, 1/q, \dots, 1/q)$. For $q = 3$, the simplex screening design is the same as the augmented simplex-centroid design shown in Figure 1(c).

When one or more mixture CPs are constrained by lower bounds $x_i \geq L_i$, the constrained experimental region is a simplex-shaped subregion of the whole simplex defined by equation (1). The simplex subregion can be transformed to the whole simplex by L -pseudocomponent transformations

$$x'_i = \frac{x_i - L_i}{1 - \sum_{j=1}^q L_j}, i = 1, 2, \dots, q \quad (16)$$

In such cases, any of the simplex designs described previously can be used. Designs expressed in proportions of the L -pseudocomponents can be transformed into designs expressed in proportions of the original components via

$$x_i = L_i + \left(1 - \sum_{j=1}^q L_j\right) x'_i \quad (17)$$

Constraints of the form $x_i \leq U_i$ can sometimes lead to a simplex-shaped subregion. Cornell [2, Section 4.6] discusses a U -pseudocomponent transformation that can be applied to adapt the simplex designs discussed previously to such situations.

Constrained Designs in Component Proportions

When there are single- and/or multiple-component constraints on the CPs, as in equations (2) and/or (3), the mixture experimental region is generally a convex polyhedron that can be irregular in shape. Designs for constrained CP mixture experiments are generally constructed by *optimal experimental design* methods. These methods involve selecting design points that satisfy the applicable constraints (possibly from a pregenerated candidate set) so as to optimize a selected design optimality criterion.

By analogy with **factorial** and **fractional factorial** designs for nonmixture variables on cuboidal experimental regions, candidate points often consist of all or some of the vertices and dimensional centroids (i.e., 1-D edge centroids, 2-D face centroids, ..., $(q - 2)$ -D face centroids, and the $(q - 1)$ -D overall centroid) of the constrained CP experimental region. Cornell [2, Section 4.9] discusses several algorithms for calculating all or a fraction of the vertices of a polyhedral experimental region defined by single-component constraints as in equation (2). Snee [16] and Piepel [17] discuss algorithms for calculating vertices and dimensional centroids for a constrained region defined by single- and/or multiple-component constraints of the forms (2) and/or (3). Mixtures in the interior of the constrained region, possibly obtained on a grid, are needed as candidate points for some design optimality criteria.

The most commonly used design optimality criteria include

D-optimality: Minimize $|(U'U)^{-1}|$, the generalized variance of the vector of estimated model coefficients.

G-optimality: Minimize $\max_{\mathbf{x} \in R} \{\mathbf{u}'(U'U)^{-1}\mathbf{u}\}$, the maximum standardized prediction variance over the constrained experimental region R .

I-optimality: Minimize $\text{avg}_{\mathbf{x} \in R} \{\mathbf{u}'(U'U)^{-1}\mathbf{u}\}$, the average standardized prediction variance over the constrained experimental region R .

In the preceding notation, $\mathbf{x}$ is a vector of predictor variables (e.g., mixture CPs), $\mathbf{X}$ is the experimental design matrix expressed in the predictor variables, and $\mathbf{u}$ and $\mathbf{U}$ are extensions of $\mathbf{x}$ and $\mathbf{X}$ in the form of the model for which an **optimal design** is desired. Note that *I-optimality* is sometimes referred to as *V-optimality* or *IV-optimality*. These optimality criteria (and others not listed here) are often called *alphabetic design optimality criteria* because letters of the alphabet are used to identify them.

Mathematically Independent Variables Approach

With the MIVs approach for mixture experiments, q mixture components and no other factors are studied. The MIVs approach can be extended to study nonmixture factors (e.g., total amount or PVs) in addition to the mixture components, but such extensions are

not discussed here. An MIVs experiment is designed and responses are modeled using $q - 1$ MIVs to which the q mixture CPs ($x_i, i = 1, 2, \dots, q$) have been transformed. The MIVs are denoted $z_i, i = 1, 2, \dots, q - 1$. Standard designs, models, and other response surface methodology for nonmixture experiments (see **Response Surface Methodology**) can be used on the basis of the $q - 1$ MIVs. For this reason, the MIV approach was widely used many years ago prior to the development of experimental design, modeling, and other methods specifically for mixture experiments. However, it is now recommended that the MIVs approach only be used if the transformation of q mixture CPs to $q - 1$ MIVs is natural or of long-standing practice. For example, in aluminum production the long-term practice has been to use the MIVs approach with bath ratio = NaF/AlF_3 in addition to other separate components such as Al_2O_3 , MgF_2 , and CaF_2 . Various forms of ratio transformations and other transformations to MIVs that are natural for a particular problem are possible. Cornell [2, Section 6.6] discusses the MIV approach using ratios of components. Cornell [2, Chapter 3] also discusses general methods for obtaining MIVs, although those methods typically do not yield transformations natural to any particular problem.

Slack Variable Approach

With the SV approach for mixture experiments, only mixture components and no other factors are studied. The experiment is designed and response(s) modeled using proportions of $q - 1$ of the q mixture components. For a given design in the $q - 1$ components, the proportion of the q th component (called the *slack variable*) is obtained by $x_q = 1 - x_1 - x_2 - \dots - x_{q-1}$ for each point in the design. In this sense the q th component “takes up the slack”. Standard designs, models, and other response surface methods (see **Response Surface Methodology**) are used with the SV approach. For example, the linear and quadratic models with the q th component as the SV are

Linear:

$$\eta(\mathbf{x}) = \alpha_0 + \sum_{i=1}^{q-1} \alpha_i x_i \quad (18)$$

Quadratic:

$$\eta(\mathbf{x}) = \alpha_0 + \sum_{i=1}^{q-1} \alpha_i x_i + \sum_{i=1}^{q-2} \sum_{j=i+1}^{q-1} \alpha_{ij} x_i x_j + \sum_{i=1}^{q-1} \alpha_{ii} x_i^2 \quad (19)$$

These models provide exactly the same fits as would the linear and quadratic Scheffé mixture models in equations (4) and (5). However, reductions of SV models can have different fits than reductions of Scheffé mixture models and can yield misleading conclusions.

The SV approach is often used when the component chosen as the SV (a) makes up most (e.g., $x_q \geq 0.80$) of the mixture, (b) has a negligible effect on the response(s), or (c) serves as a filler or diluent. In all of these cases applications or adaptations of the CP approach are appropriate, whereas the SV approach can result in misleading conclusions. In situation (a), a component that makes up the majority of a mixture often has large effects on the response variables. However, the effects of this component are confounded with the effects of other components in models such (18) or (19). Hence, the SV approach could lead to incorrect conclusions about the effects of other components as well as the SV component. It is better to use the CP approach in this situation. In situation (b), when a component has no effect on a response over the composition region of interest, the CP approach is appropriate, with the response depending on the proportions of the remaining components (i.e., $X_j = x_j / (1 - \sum_{k=1}^{q-1} x_k)$, with $j = 1, 2, \dots, q-1$). In situation (c), a filler or diluent component typically would not blend nonlinearly with the remaining components, but still may affect the response. In such cases, the CP approach using a model homogeneous of degree one (Becker H1, H2, and H3) or a log-contrast model should be used.

Mixture-Amount Approach

When a response variable is expected to depend on the total amount of the mixture as well as the proportions of the components, a MA approach should be used. An MA experiment is designed and the response(s) modeled in terms of the proportions (x_i) and total amount (A) of the mixture components. The

CPs are subject to the basic mixture constraints in equation (1), and possibly to single-component constraints in equation (2) and/or multicomponent constraints in equation (3). The total amount is typically restricted to fall within a specified range, often with two or three levels specified for study. Typically two levels of A are coded as $A' = -1$ and $+1$, while three levels are coded as $A' = -1, 0$, and $+1$.

Mixture-Amount Models

If the amount of a mixture affects a response but not the blending effects of the mixture components, then MA models have the general form

$$\eta(\mathbf{x}) = \{\text{Mixture model}\} + \{\text{Amount model}\} \quad (20)$$

The mixture model can take any permissible form in the CPs (e.g., the forms in equations (4–7)), while the amount model is typically a linear or quadratic polynomial in the amount A (or coded A'). Two example models of this type are

Linear–linear:

$$\eta(\mathbf{x}) = \sum_{i=1}^q \beta_i^0 x_i + \beta_0^1 A' \quad (21)$$

Quadratic–quadratic:

$$\eta(\mathbf{x}) = \sum_{i=1}^q \beta_i^0 x_i + \sum_{i=1}^{q-1} \sum_{j=i+1}^q \beta_{ij}^0 x_i x_j + \beta_0^1 A' + \beta_0^2 A'^2 \quad (22)$$

In equation (21), the mixture components blend linearly and the amount variable has a linear effect on the response. In equation (22), the mixture components blend quadratically and the amount variable has a quadratic effect on the response. In both equations, the component blending is not affected by the amount of the mixture (i.e., the components and amount do not interact).

If the amount of a mixture affects the component blending properties, then MA models have the general form

$$\eta(\mathbf{x}) = (\text{Mixture model})(\text{Amount model}) + \text{Mixture terms} + \text{Amount terms} \quad (23)$$

An example model of this form is

$$\eta(\mathbf{x}) = \left(\sum_{i=1}^q \beta_i x_i \right) (\alpha_0 + \alpha_1 A')$$

$$\begin{aligned}
& + \sum_{i=1}^{q-1} \sum_{j=i+1}^q \beta_{ij}^0 x_i x_j + \beta_0^2 A'^2 \\
& = \sum_{i=1}^q \beta_i^0 x_i + \sum_{i=1}^q \beta_i^1 x_i A' \\
& + \sum_{i=1}^{q-1} \sum_{j=i+1}^q \beta_{ij}^0 x_i x_j + \beta_0^2 A'^2 \quad (24)
\end{aligned}$$

where (a) the linear blending effects of the components are affected by the amount of the mixture, (b) the components have quadratic blending effects that are not affected by the amount of the mixture, and (c) the amount has the same quadratic effect on the response for all mixtures.

Mixture-Amount Designs

Complete MA designs can be formed by running a CP mixture experiment design at each of two or three (or more, if desired) levels of the amount variable. As an example, Figure 2 illustrates a constrained mixture design in three components performed at each of two levels of the amount variable. This design supports fitting a MA model of the form

$$\begin{aligned}
\eta(\mathbf{x}) = & \beta_1^0 x_1 + \beta_2^0 x_2 + \beta_3^0 x_3 + \beta_{12}^0 x_1 x_2 + \beta_{13}^0 x_1 x_3 \\
& + \beta_{23}^0 x_2 x_3 + \beta_{123}^0 x_1 x_2 x_3 + \beta_1^1 x_1 A' + \beta_2^1 x_2 A' \\
& + \beta_3^1 x_3 A' + \beta_{12}^1 x_1 x_2 A' + \beta_{13}^1 x_1 x_3 A' \\
& + \beta_{23}^1 x_2 x_3 A' + \beta_{123}^1 x_1 x_2 x_3 A' \quad (25)
\end{aligned}$$

With A coded to -1 and $+1$, the terms with 0 superscripts describe the component blending at

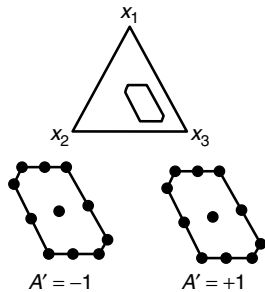


Figure 2 Example of a complete mixture-amount design for three constrained mixture components and two levels of total amount

the average of the two amounts (i.e., at zero-coded amount), while the terms with a 1 superscript describe the effect of the amount variable on the component blending properties relative to the average component blending properties.

Fractional MA designs can be formed with a fraction of a CP mixture design at one or more levels of the total amount variable. For example, the MA design with three mixture components and three levels of total amount shown in Figure 3 has a complete simplex-centroid design at $A' = -1$ and $+1$, but a vertex-only design at $A' = 0$. This design supports fitting the model

$$\begin{aligned}
\eta(\mathbf{x}) = & \beta_1^0 x_1 + \beta_2^0 x_2 + \beta_3^0 x_3 + \beta_{12}^0 x_1 x_2 + \beta_{13}^0 x_1 x_3 \\
& + \beta_{23}^0 x_2 x_3 + \beta_{123}^0 x_1 x_2 x_3 + \beta_1^1 x_1 A' + \beta_2^1 x_2 A' \\
& + \beta_3^1 x_3 A' + \beta_{12}^1 x_1 x_2 A' + \beta_{13}^1 x_1 x_3 A' \\
& + \beta_{23}^1 x_2 x_3 A' + \beta_{123}^1 x_1 x_2 x_3 A' + \beta_1^2 x_1 A'^2 \\
& + \beta_2^2 x_2 A'^2 + \beta_3^2 x_3 A'^2 \quad (26)
\end{aligned}$$

which assumes that the total amount has (a) linear effects on the linear and nonlinear blending properties of the mixture components and (b) a quadratic effect only on the linear blending properties of the mixture components.

Piepel *et al.* [18] discusses the special case of an MA experiment and models when the total amount takes a zero value in addition to nonzero values. These experiments are sometimes referred to as *mixture-amount-zero* (MAZ) experiments.

Component Amounts Approach

With the CA approach for mixture experiments, q mixture components and no other factors are studied. The CA approach can be extended to study nonmixture factors (e.g., total amount or PVs) in addition to the mixture components, but such extensions are not

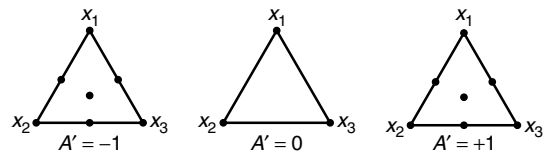


Figure 3 Example of a fractional mixture-amount design with three mixture components and three levels of total amount

8 Mixture Experiments

discussed here. A CA experiment is designed and responses are modeled using amounts (e.g., masses or volumes) of the mixture components, denoted $a_i, i = 1, 2, \dots, q$. The amounts of mixture components are not subject to the constraints in equation (1) for mixture CPs. Hence, the experimental region of interest can be scaled to be cuboidal or spherical, enabling standard statistical designs such as *factorial*, *fractional factorial*, *Plackett–Burman*, **central composite**, *Box–Behnken*, and others to be used. Similarly, standard response surface models, such as

Linear:

$$\eta(\mathbf{x}) = \alpha_0 + \sum_{i=1}^q \alpha_i a_i \quad (27)$$

Linear plus interaction:

$$\eta(\mathbf{x}) = \alpha_0 + \sum_{i=1}^q \alpha_i a_i + \sum_{i=1}^{q-1} \sum_{j=i+1}^q \alpha_{ij} a_i a_j \quad (28)$$

Quadratic:

$$\eta(\mathbf{x}) = \alpha_0 + \sum_{i=1}^q \alpha_i a_i + \sum_{i=1}^{q-1} \sum_{j=i+1}^q \alpha_{ij} a_i a_j + \sum_{i=1}^q \alpha_{ii} a_i^2 \quad (29)$$

and other standard data analysis and graphics methods (see **Response Surface Methodology**) can be used with the CA approach.

When using the CA approach, information about effects of CPs (x_i) and effects of the total amount of the mixture ($\sum_{i=1}^q a_i = A$) are confounded because $a_i = x_i A, i = 1, 2, \dots, q$. If it is natural or desirable to separately understand the effects of the component (x_i) and total amount (A) variables on the response(s), then the MA approach (discussed previously) should be used. On the other hand, if it is natural or desirable, or to understand the effects of the CAs on the response(s), the CA approach is appropriate. Finally, if the total amount of the mixture turns out to have little or no effect on the response, the CA approach could complicate discovering it, so that the MA approach is preferred for such cases. See Piepel [19] for more discussion of the CA approach and how it compares to the MA approach.

Mixture-Process Variable Approach

With the MPV approach, the effects on a response variable of both mixture components and process

variables (a general term for “other factors”) are studied. A MPV experiment is usually designed and responses modeled in terms of the proportions of the mixture components and the levels of the PVs. The CPs are subject to the basic mixture constraints in equation (1), and possibly to single-component constraints in equation (2) and/or multicomponent constraints in equation (3). One or more PVs are typically restricted to fall within specified ranges. Often, two or three values of each PV are specified for study. Typically two levels of a PV would be coded as $z = -1$ and $+1$, while three levels would be coded as $z = -1, 0$, and $+1$.

Mixture-Process Variable Models

If the levels of PVs affect a response but not the blending effects of the mixture components, then MPV models have the general form

$$\eta(\mathbf{x}) = \{\text{Mixture model}\} + \{\text{PVs model}\} \quad (30)$$

The mixture model can take any permissible form in the CPs (e.g., the forms in equations (4–7)), while the PV model is typically a linear, linear plus interactions, or quadratic polynomial in the PVs. An example model of this type with two PVs (z_1 and z_2) is given by

$$\eta(\mathbf{x}) = \sum_{i=1}^q \beta_i^0 x_i + \sum_{i=1}^{q-1} \sum_{j=i+1}^q \beta_{ij}^0 x_i x_j + \beta_0^1 z_1 + \beta_0^2 z_2 + \beta_0^{12} z_1 z_2 \quad (31)$$

In this model, the mixture components blend quadratically and the PVs have linear and interactive effects. However, the component blending is not affected by the PVs (i.e., the components and PVs do not interact).

If the levels of PVs affect the component blending properties (which happens frequently in practice), then MPV models have the general form

$$\eta(\mathbf{x}) = (\text{Mixture model})(\text{PVs model}) + \text{Mixture terms} + \text{PVs terms} \quad (32)$$

An example model of this form is

$$\eta(\mathbf{x}) = \left(\sum_{i=1}^q \beta_i x_i \right) (\alpha_0 + \alpha_1 z_1 + \alpha_2 z_2)$$

$$\begin{aligned}
& + \sum_{i=1}^{q-1} \sum_{j=i+1}^q \beta_{ij}^0 x_i x_j + \beta_0^1 z_1^2 \\
& = \sum_{i=1}^q \beta_i^0 x_i + \sum_{i=1}^{q-1} \sum_{j=i+1}^q \beta_{ij}^0 x_i x_j + \sum_{i=1}^q \beta_i^1 x_i z_1 \\
& \quad + \sum_{i=1}^q \beta_i^2 x_i z_2 + \beta_0^{11} z_1^2 \quad (33)
\end{aligned}$$

where (a) the linear blending properties of the components are affected by linear effects of the PVs, (b) the components have quadratic blending properties that are not affected by the PVs, and (c) one PV (z_1) has a quadratic effect on the response that is the same for all mixtures.

Mixture-Process Variable Designs

A *complete MPV design* is formed by running the same CP mixture experiment design at each point in a PV design (or vice versa). A *fractional MPV design* results from running different mixture designs at each PV design point (or vice versa). As an example, Figure 4 illustrates a simplex-centroid mixture design in three components performed at each of the four combinations in a 2^2 design for two PVs. This 28-point complete MPV design supports exactly fitting a 28-term model of the form

$$\eta(\mathbf{x}) = \beta_1^0 x_1 + \beta_2^0 x_2 + \beta_3^0 x_3 + \beta_{12}^0 x_1 x_2 + \beta_{13}^0 x_1 x_3$$

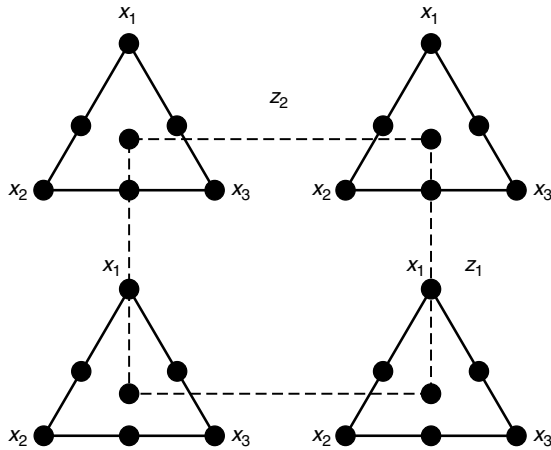


Figure 4 Example of a complete mixture-process variable design with a simplex-centroid mixture design in three components at each of the four combinations of a 2^2 design for two process variables

$$\begin{aligned}
& + \beta_{23}^0 x_2 x_3 + \beta_{123}^0 x_1 x_2 x_3 + (\beta_1^1 x_1 + \beta_2^1 x_2 \\
& + \beta_3^1 x_3 + \beta_{12}^1 x_1 x_2 + \beta_{13}^1 x_1 x_3 + \beta_{23}^1 x_2 x_3 \\
& + \beta_{123}^1 x_1 x_2 x_3) z_1 + (\beta_1^2 x_1 + \beta_2^2 x_2 + \beta_3^2 x_3 \\
& + \beta_{12}^2 x_1 x_2 + \beta_{13}^2 x_1 x_3 + \beta_{23}^2 x_2 x_3 + \beta_{123}^2 x_1 x_2 x_3) z_2 \\
& + (\beta_1^{12} x_1 + \beta_2^{12} x_2 + \beta_3^{12} x_3 + \beta_{12}^{12} x_1 x_2 \\
& + \beta_{13}^{12} x_1 x_3 + \beta_{23}^{12} x_2 x_3 + \beta_{123}^{12} x_1 x_2 x_3) z_1 z_2 \quad (34)
\end{aligned}$$

With z_1 and z_2 coded to -1 and $+1$ the terms with 0 superscripts describe the component blending at the average levels of the PVs (i.e., zero-coded values). Terms with a 1 or a 2 superscript describe the effects of the first or second PV on the component blending properties relative to the average component blending properties. The terms with a 12 superscript describe the interactive effects of the two PVs on the component blending properties relative to the average component blending properties.

Sometimes MPV designs are run as **split-plot** experiments when the mixture blends or the PVs are hard to change. When the mixture components are hard to change, all points in a PV design are run together for each point in a mixture design. When the PVs are hard to change, all points in a mixture design are run together for each point in a PV design. MPV designs run as split-plot experiments require special data analysis methods to account for the restriction on **randomization** of the design points. See Cornell [2, Sections 7.3 and 7.4] and Kowalski *et al.* [20] for additional discussion of the design and analysis of MPV split-plot experiments.

Mixture-of-Mixtures Approach

With the MoM approach, the end product is a mixture of main components (sometimes called *major components*), which are themselves mixtures of sub-components (sometimes called *minor components*). For example, an ink-jet printer image is a mixture of three to six colors of inks, with each ink a mixture of chemicals. In some cases the main components represent categories of components, and the sub-components are kinds of ingredients in each category. Such cases are referred to as *categorized component* (CC) mixture experiments.

There are two types of MoM/CC experiments. In one type (type A), the proportions of the main

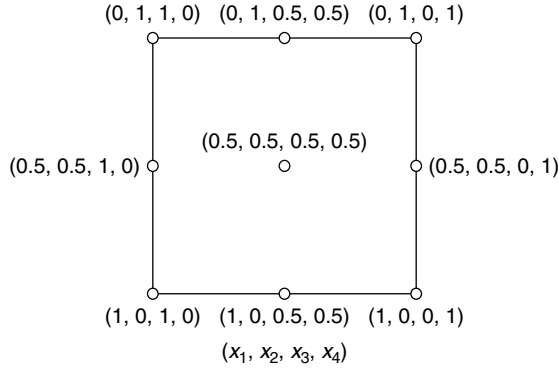


Figure 5 The $\{2, 2; 2, 2\}$ double-lattice design for a mixture-of-mixtures experiment

components have fixed values and only the proportions of the subcomponents are varied. In this type of MoM/CC experiment, only the blending among subcomponents is studied. In the other type (type B) of MoM/CC experiments, the proportions of the main components and the subcomponents are varied. In this type of experiment, blending among main components and among subcomponents is studied. Models and designs for both types of MoM/CC experiments are discussed by Cornell [2, Section 4.13]. In an example discussed by Cornell [2, Section 4.13.2] the main components are two types of resins, with two subcomponents (x_1 and x_2 for the first type, x_3 and x_4 for the second type). The $\{2, 2; 2, 2\}$ double-lattice design used in this example is shown in Figure 5. The corresponding double-Scheffé (quadratic–quadratic) model is

$$\begin{aligned} \eta(\mathbf{x}) = & \gamma_{13}x_1x_3 + \gamma_{14}x_1x_4 + \gamma_{23}x_2x_3 + \gamma_{24}x_2x_4 \\ & + \gamma_{123}x_1x_2x_3 + \gamma_{124}x_1x_2x_4 + \gamma_{134}x_1x_3x_4 \\ & + \gamma_{234}x_2x_3x_4 + \gamma_{1234}x_1x_2x_3x_4 \end{aligned} \quad (35)$$

See Cornell [2, Section 4.13.2] for further discussion of this example. Piepel [21] discusses methods for developing and reducing models for type A MoM experiments with a pharmaceutical tablet formulation example.

Multifactor Mixture Approach

In a *MFM* experiment, mixture experiments are performed for each of two or more independent mixture factors in an experiment. If there are

n mixture factors each having $q_k (k = 1, 2, \dots, n)$ components, then the experimental region is $(\sum_{k=1}^n q_k - n)$ dimensional.

Nigam [22] discusses an MFM example in which the goal is to maximize recombination frequency of a mixed bacterial culture (male and female) when irradiated with short exposures of mixtures of mutagens (x-rays and γ -rays). In this example there are two mixture factors, bacterial strains and mutagens, with each having two components. Note that the two mixture factors are independent and are not themselves major components of a mixture, which differentiates MFM experiments from MoM/CC experiments. Nigam [22] discusses designs and models for MFM experiments.

Optimal Experimental Design

Optimal experimental design methods were discussed previously for the CP approach, but can be used with most of the approaches when needed. Exceptions are the MoM and MFM approaches, where the methods are applicable but not currently implemented in available software. Two common situations when optimal experimental design is needed are (a) the experimental region is irregular in shape because of linear constraints on the variables, and (b) the standard designs for a given approach contain too few or too many design points compared to the desired number.

For optimal experimental design methods that require a set of candidate points, the algorithms of Piepel [16] provide for generating the vertices and various dimensional centroids of any polyhedral region defined by linear inequality constraints on mixture variables, nonmixture variables, or both. Those algorithms only allow for at most one equality (mixture) constraint and hence are applicable to all but the MoM and MFM approaches. Even for those approaches candidate points can be generated separately and combined. However, the optimal design algorithms and software for candidate and candidate-free methods would have to be modified to account for more than one equality (mixture) constraint.

Optimal design methods are very useful for generating nonstandard designs, but there are three cautions to keep in mind. First, optimal designs (especially if constructed using vertices and dimensional centroids as candidate points) often have all (or nearly

all) design points on or near the boundary of the constrained experimental region. It may be desirable to have a substantial fraction of the design points in the interior of the experimental region. *Space-filling designs* (sometimes referred to as **uniform designs**) are one alternative to obtain a more even coverage of the interior as well as the exterior of a constrained experimental region. Piepel [4, Table 2] lists several publications that discuss uniform designs for CP mixture experiments. Other alternatives are the *central composite analog designs* (CCADs) and *layered designs* (LDs) proposed by Piepel *et al.* [23] to assure having design points on different layers of the constrained experimental region. **Ordinary central composite designs** can be used with the MIV, SV, CA approaches, and the PV portion of an MPV experiment. CCADs can be applied with the CP approach or the mixture portion of the MA or MPV approaches. LDs can be applied with any of the mixture approaches.

The second caution is that optimal designs are constructed for a specific model form, and may not be optimal or even good for other model forms. This tends to be the case when a design is optimized for a simpler (smaller) model but a larger model is more appropriate. However, designs optimized for a larger model often tend to be good for reduced forms of that model. Hence, construction of optimal designs for the largest model form that may be needed is generally recommended.

The third caution is that the number of design points must be specified in the optimal design approach. Practitioners should investigate optimal designs with fewer and more points than the desired number and assess the sensitivity of design properties to the number of design points.

Assessing and Comparing Designs

Piepel [4, Section 5] summarizes several analytical and graphical methods for assessing and comparing mixture designs. These include design efficiency measures, collinearity and leverage diagnostics, variance and **bias** dispersion graphs, fraction of design space plots, and other kinds of prediction variance plots. Many of these methods are applicable or can be adapted regardless of the mixture experiment approach used. For discussion of these methods, see **Assessment of Experimental Designs** and the associated references in [4].

Data Analyses for Mixture Experiments

The data analysis process for a mixture experiment usually begins by fitting and assessing models (corresponding to the selected approach) until a model is obtained that does not substantially underfit or overfit the data. Traditional regression methods for assessing model **lack of fit** are applicable to all of the mixture experiment approaches. However, appropriate formulas must be used to calculate entries of the analysis of variance table and related statistics (e.g., R^2) associated with models using the CP, MA, MPV, MoM, and MFM approaches (see Cornell [2, Section 5.2]).

After an adequately fitting model is obtained, it can be used to achieve the desired goals of the experiment. If the goal is to identify important (and unimportant) mixture components in the CP approach, numeric estimates of component effects can be calculated or displayed graphically using response trace plots (see Cornell [2, Section 5.9] and Smith [3, Section 11]). These methods can be adapted for the MA and MPV approaches. If the goal is to select settings of the mixture and/or nonmixture variables to optimize one response variable, standard optimization methods can be used for any of the approaches. If there are two or more response variables, multiresponse optimization methods can be applied for any of the approaches. Cornell [2, Sections 8.13 and 8.14] and Smith [3, Section 12.2] discuss these methods for the CP approach.

Piepel [4, Section 7] lists and provides references for many other data analysis methods applicable for the CP, MA, and MPV approaches. Several of these are also discussed throughout the book by Cornell [2]. Standard data analysis methods from *response surface methodology* are applicable to the MIV, CA, and SV approaches, but can result in misleading conclusions for the SV approach as noted previously.

Choosing a Mixture Experiment Approach

Choosing a mixture experiment approach depends primarily on the types of variables whose effects on the response variables are to be studied in the experiment. If only mixture variables (components) affect the response(s) and are to be studied, then the CP, MIV, and SV approaches are possible. However, the CP approach is recommended in general because the experiment is designed, the data are

analyzed, and conclusions are made directly in terms of the proportions of the mixture variables. The MIV approach can be appropriate if the MIVs are the natural or long-standing variables of interest for a particular problem. The SV approach is not recommended because of the potential for making wrong interpretations and conclusions.

If the total amount of the mixture as well as the CPs affect the response(s), then the MA and CA approaches are possible. The MA approach is recommended when the effects of the components and the effect of the total amount of the mixture are to be studied, understood, and optimized separately. The CA approach is appropriate when the effects of amounts of the individual components are to be studied, rather than effects of the component proportions and total amount.

The MPV approach should be used when the effects on the response(s) of mixture components and PVs are to be studied simultaneously. While it is possible to use the MIV and SV approaches in a MPV situation, doing so is not recommended because these approaches do not directly study all of the mixture components and PVs.

Finally, the MoM and MFM approaches are used in special cases that involve more than one mixture. The MoM approach is appropriate when the components of a primary mixture are in turn mixtures of subcomponents, hence the terminology of *mixtures-of-mixtures*. The MFM approach is appropriate when there are two or more independent factors to be studied in an experiment, and the factors are mixtures of different sets of components.

References

- [1] Piepel, G.F. & Cornell, J.A. (1994). Mixture experiment approaches: examples, discussion, and recommendations, *Journal of Quality Technology* **26**, 177–196.
- [2] Cornell, J.A. (2002). *Experiments with Mixtures: Designs, Models, and the Analysis of Mixture Data*, John Wiley & Sons, New York.
- [3] Smith, W.F. (2005). *Experimental Design for Formulation*, ASA-SIAM Series on Statistics and Applied Probability, ASA, SIAM, Alexandria, Philadelphia.
- [4] Piepel, G.F. (2004). 50 years of mixture experiment research: 1955–2004, in *Response Surface Methodology and Related Topics*, A.I. Khuri, ed, World Scientific Press, Singapore, Chapter 12, pp. 283–327.
- [5] Scheffé, H. (1958). Experiments with mixtures, *Journal of the Royal Statistical Society. Series B* **20**, 344–360.
- [6] Cox, D.R. (1971). A note on polynomial response functions for mixtures, *Biometrika* **58**, 155–159.
- [7] Piepel, G.F., Hicks, R.D., Szychowski, J.M. & Loepky, J.L. (2002). Methods for assessing curvature and interaction in mixture experiments, *Technometrics* **44**, 161–172.
- [8] Smith, W.F. & Beverly, T.A. (1997). Generating linear and quadratic Cox mixture models, *Journal of Quality Technology* **29**, 211–224.
- [9] Piepel, G.F. (2006). A note comparing component-slope, Scheffé, and Cox parameterizations of the linear mixture experiment model, *Journal of Applied Statistics* **33**, 397–403.
- [10] Piepel, G.F. (2007). A component slope linear model for mixture experiments, *Quality Technology and Quantitative Management* **4**(3).
- [11] Piepel, G.F., Szychowski, J.M. & Loepky, J.L. (2002). Augmenting Scheffé linear mixture models with squared and/or crossproduct terms, *Journal of Quality Technology* **34**, 297–314.
- [12] Becker, N.G. (1968). Models for the response of a mixture, *Journal of the Royal Statistical Society. Series B* **30**, 349–358.
- [13] Snee, R.D. (1971). Design and analysis of mixture experiments, *Journal of Quality Technology* **3**, 159–169.
- [14] Aitchison, J. & Bacon-Shone, J. (1984). Log contrast models for experiments with mixtures, *Biometrika* **71**, 323–330.
- [15] Snee, R.D. & Marquardt, D.W. (1976). Screening concepts and designs for experiments with mixtures, *Technometrics* **18**, 19–29.
- [16] Snee, R.D. (1979). Experimental designs for mixture systems with multicomponent constraints, *Communications in Statistics: Theory and Methods* **A8**, 303–326.
- [17] Piepel, G.F. (1988). Programs for generating extreme vertices and centroids of linearly constrained experimental regions, *Journal of Quality Technology* **20**, 125–139.
- [18] Piepel, G.F. (1988). A note on models for mixture-amount experiments when the total amount takes a zero value, *Technometrics* **30**, 449–450.
- [19] Piepel, G.F. (1985). Models and designs for generalizations of mixture experiments where the response depends on the total amount, Ph.D. Dissertation, University Microfilms International, Ann Arbor. (http://www.umi.com/products_umi/dissertations).
- [20] Kowalski, S.M., Cornell, J.A. & Vining, G.G. (2002). Split-plot designs and estimation methods for mixture experiments with process variables, *Technometrics* **44**, 72–79.
- [21] Piepel, G.F. (1999). Modeling methods for mixture-of-mixtures experiments applied to a tablet formulation problem, *Pharmaceutical Development and Technology* **4**, 593–606.
- [22] Nigam A.K. (1973). Multifactor mixture experiments. *Journal of the Royal Statistical Society* **35**, 51–56.
- [23] Piepel, G.F., Anderson, C.M. & Redgate, P.E. (1993). Response surface designs for irregularly-shaped regions (Parts 1, 2, and 3),. *Proceedings of the Section on*

Physical and Engineering Sciences, American Statistical Association, Alexandria, pp. 205–227.

Related Articles

**Assessment of Experimental Designs; Box–Behn-
ken Designs; Central Composite Designs; Factorial**

**Experiments; Fractional Factorial Designs; Opti-
mal Design; Plackett–Burman Designs; Response
Surface Methodology; Split-Plot Designs; Uniform
Experimental Designs.**

GREG F. PIEPEL

Response Surface Methodology

Introduction

Response surface methodology (RSM) is a strategy for exploring the relationship between controllable factors and important response variables. RSM makes extensive use of statistically designed experiments and regression models (*see* **Least-Squares Estimation**). A key element to RSM is the notion of sequential learning (*see* **Sequential Experimentation**): where initial experiments help focus attention on the most important factors, subsequent steps help to bring these factors to settings that achieve good results, and finally efficient statistical modeling techniques are used to estimate how these factors affect the response variables of greatest importance.

RSM was first introduced in a pathbreaking paper by Box and Wilson [1]. That article introduced the ideas of sequential experiments, the central composite design (*see* **Central Composite Designs**) and the use of **sequential experimentation** to find optimal operating conditions for a chemical process. Box continued to develop these ideas in two further articles [2, 3], developed new classes of designs [4, 5] and introduced the concepts of rotatability ([6], **Rotatable Designs and Rotatability**) and minimum **mean squared error** designs [7].

Applications

RSM is one of the most influential and widespread scientific methods proposed in the twentieth century. It has been used in a wide variety of fields such as chemistry, engineering, materials science, and biology, among others. We illustrate RSM with a biotechnology study by Kawaguti, Manrich and Sato [8]. That study aims at finding conditions for achieving a high rate of activity of glucosyltransferase, an enzyme that has important uses in the food industry. The enzyme is produced by bacteria in a culture medium. Various nutrients can be added to the medium and some of these may promote enzyme activity. The experiments described here studied three such nutrients: bacteriological peptone (BP), yeast extract and sugar cane molasses (SCM). The first two factors

were included as sources of nitrogen and the third as a source of carbon.

The Sequential Strategy

In this section we provide an overview of the strategic approach that guides RSM. Most applications begin with a rather large number of potential experimental factors and one or more important process outcomes. The study by Kawaguti, Manrich and Sato focused on just three factors [8], but the original article actually mentions quite a few additional factors, including the basal medium used for the culture, amounts of sucrose, beef and agar that were added to the medium, incubation conditions of the cultures, the speed and duration in which culture flasks were shaken, and the time allotted for fermentation. The ultimate goal of this investigation was to discover how enzyme activity depends on the factors under study and to find good operating conditions.

The following steps are typical in RSM studies:

- screening – identifying which of the factors are really critical for performance;
- effect **estimation** – estimating the effect of each of the important factors;
- steepest ascent – following an estimated direction of steepest ascent to achieve better results;
- response modeling – fitting a more complex regression model that includes nonlinear effects of the factors;
- optimization – using the fitted surface to estimate optimal operating conditions.

The screening step is quite intuitive. Suppose the screening step identifies two key factors and we proceed with these factors to the next steps. An effective analogy to understand the remaining journey is to think of climbing to the top of a mountain. Imagine a contour map in which our two factors form the north–south and east–west axes and the contour lines map regions of high and low response. We cannot see the map. Our only information comes from running experiments that expose the likely height at particular factor settings. Early in the RSM program we are often near the base of the mountain. Our experiments for estimating effects will pick up the direction of the slope that leads us uphill. Experiments along the line of steepest ascent follow this route up until we reach a crest. Then, further

experiments may be used to identify a new direction of steepest ascent. When we approach the peak of the mountain, we conduct more extensive experiments that enable us to model the topography of the contour map. The results of those experiments can be used to help identify the location of the peak. When more than two factors pass the screening phase, the visual analogy of the mountain climber is less immediate, but the same ideas can be applied to explore and map the surface, assess factor effects and discover optimal factor settings.

Many decisions confront the RSM team at each stage of the journey:

1. Should factors be dropped?
2. Should new factors be added?
3. Should the range of one of the factors be moved?
4. Should the range be made narrower? Wider?
5. Should points be added to a given design to accommodate more complex models?

No two experimental teams will make the same decisions. Yet the RSM strategy offers great flexibility, so that there are many routes to success.

The approach described above relies heavily on the ability to “learn as you go”. Results from the screening experiment will be valuable in deciding whether to drop a factor or to add a new factor. Identifying trends that point to better performance will often lead to changing the range of a factor in subsequent steps. Sequential experimentation thus plays a key role in RSM. And it is not surprising that industrial problems led to the birth of RSM; unlike the agricultural experiments that preceded them, industrial experiments could provide the rapid feedback needed for sequential work.

Screening Experiments

The initial phase of many RSM investigations is to identify the most important factors from what may be a long list. The best designs at this stage are ones that enable the team to study a large number of factors relative to the number of experimental runs that are made. Thus it is common to limit the experimental factors to just two levels and to use designs like small 2^{k-p} fractional factorials (see **Fractional Factorial Designs**); and Plackett and Burman designs [9] (see **Plackett–Burman Designs**). See also **Screening Designs**; **Main Effect Designs**. When there are many

candidate factors, one might even want to use a *supersaturated design* (see **Supersaturated Designs**), in which the number of factors is greater than or equal to the **sample size**. Interested readers can consult those entries for details.

Effect Estimation

At this stage, experimental results are used to estimate how changing each factor affects the response variables. Often this stage concentrates on main effects of the factors (see **Main Effects**). Sometimes there are *interactions* (see **Interactions**) between factors; i.e., the effect of a factor on the response is not always the same, but in fact depends on how some other factor has been set. When enough data are available, interactions can be estimated alongside main effects. Again two-level experimental designs are the ones most commonly used. Center points (see **Center Points**), in which all factors are set to an intermediate level, may also be added to the design. Center points allow the RSM team to check for the possibility that a full second-order regression model is needed. Usually more than one **center point** is included, which also permits a “model free” estimate of the experimental variability. Sometimes the effects are estimated from the screening experiment data.

Table 1 shows 11 experimental runs made by [8] in their first experiment, which correspond to a 2^3 design with three center points. In fitting regression models to data from these experiments, it is common to use “coded” versions of the factors in which the high-level is coded to 1 and the low level to -1 . As in this experiment, the intermediate level will typically be exactly between the high and low levels, and so will be coded to zero.

With the 2^3 design, it is possible to estimate all the main effects and interactions of the three factors. Analysis of these data reveals several important features:

1. Most settings give very low enzyme activity, but two runs generated much higher activity.
2. There is a strong **main effect** for SCM.
3. There is a strong main effect for yeast extract Prodex (YEP).
4. There is a very strong interaction effect of SCM and YEP.
5. Effects for BP are weak.

Table 1 Data from the first RSM experiment performed by Kawaguti *et al.* [8]. Coded values for the factors are given in parentheses. Actual levels of all factors are in g/L

Sugar cane molasses (X_1)	Bacteriological peptone (X_2)	Yeast extract Prodex (X_3)	Enzyme activity (Y) (in U/mL)
130.36 (−1)	18.10 (−1)	37.14 (−1)	5.63
219.64 (1)	18.10 (−1)	37.14 (−1)	0.13
130.36 (−1)	41.90 (1)	37.14 (−1)	5.97
219.64 (1)	41.90 (1)	37.14 (−1)	0.20
130.36 (−1)	18.10 (−1)	72.86 (1)	0.63
219.64 (1)	18.10 (−1)	72.86 (1)	0.28
130.36 (−1)	41.90 (1)	72.86 (1)	0.32
219.64 (1)	41.90 (1)	72.86 (1)	2.12
175 (0)	30 (0)	55 (0)	0.42
175 (0)	30 (0)	55 (0)	0.47
175 (0)	30 (0)	55 (0)	0.48

6. The results at the center points are much lower than the average result at the **factorial** points, suggesting that there may be value in fitting a full second-order regression model.

Figure 1 shows a contour plot of the estimated response, as a function of SCM and YEP, when BP is at its intermediate level. The most important conclusion above is the first, which suggests that the experimental regime for these first runs was most likely too wide, placing most runs in regions of the factor space where there was little enzyme activity. The two runs with high activity suggest that there may be great advantage in repeating the experiment, but modifying the settings for the factors.

Steepest Ascent

Let $X_1, \dots, X_k$ denote the k factors (in coded units) in an initial experiment to estimate factor effects. Often the analysis shows that a first-order regression model gives a good fit to the response Y :

$$\hat{Y} = b_0 + b_1X_1 + b_2X_2 + \dots + b_kX_k \quad (1)$$

Here the b_j are the coefficients estimated by fitting a regression model to the data and $\hat{Y}$ is the estimated value of the response.

A first-order regression model is appropriate when the response is affected in a simple and linear fashion by the factors. If, for example, we were to start from the center point of the design and “move away” in a fixed direction, the response

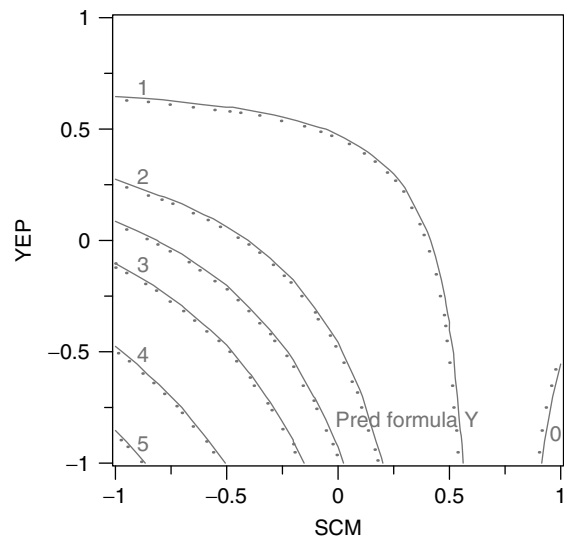


Figure 1 Contour plot of the estimated enzyme activity based on the 11 runs in Table 1. The factor SCM is on the horizontal axis and the factor YEP is on the vertical axis; both are in coded units. The third factor, BP, is set to its intermediate level, 30

would change proportionally to our distance from the center point. The idea of steepest ascent is simply to move in the direction that gives the largest slope in the desired direction (positive if we want to maximize the response, negative if we want to minimize the response). It is easy to find this direction by differentiating the estimated regression equation. The direction of steepest ascent is proportional to the vector of slope coefficients: $(b_1, b_2, \dots, b_k)$. For

4 Response Surface Methodology

Table 2 Data from the second RSM experiment performed by Kawaguti *et al.* [8]. Coded values for the factors are given in parentheses. Actual levels of all factors are in g/L

Sugar cane molasses (X_1)	Bacteriological peptone (X_2)	Yeast extract Prodex (X_3)	Enzyme activity (Y) (in U/mL)
120.24 (−1)	11.07 (−1)	18.10 (−1)	6.34
179.76 (1)	11.07 (−1)	18.10 (−1)	8.28
120.24 (−1)	28.93 (1)	18.10 (−1)	7.56
179.76 (1)	28.93 (1)	18.10 (−1)	8.97
120.24 (−1)	11.07 (−1)	41.90 (1)	7.39
179.76 (1)	11.07 (−1)	41.90 (1)	8.68
120.24 (−1)	28.93 (1)	41.90 (1)	6.63
179.76 (1)	28.93 (1)	41.90 (1)	6.28
100 (−1.68)	20 (0)	30 (0)	4.99
200 (1.68)	20 (0)	30 (0)	7.01
150 (0)	5 (−1.68)	30 (0)	7.06
150 (0)	35 (1.68)	30 (0)	7.47
150 (0)	20 (0)	10 (−1.68)	8.47
150 (0)	20 (0)	50 (1.68)	7.06
150 (0)	20 (0)	30 (0)	8.43
150 (0)	20 (0)	30 (0)	8.23
150 (0)	20 (0)	30 (0)	7.83

example, suppose the estimated equation were $\hat{Y} = 21 + 4X_1 + 7X_2 - 5X_3$. The direction of steepest ascent would be (4, 7, −5), meaning that for each increase of X_1 by 4 coded units, we should increase X_2 by 7 coded units and we should decrease X_3 by 5 coded units. If we wanted to minimize the response, we could proceed in the direction of steepest descent: for each decrease of X_1 by 4 coded units, we should decrease X_2 by 7 coded units and we should increase X_3 by 5 coded units. Note that the two directions are exactly opposite to one another.

Applying the method of steepest ascent will typically lead to “moving” the factor ranges away from those used in the initial experiment to settings that give better results.

For the data we have presented on enzyme activity, there was a strong interaction. Hence a first-order regression model – and steepest ascent – will not be appropriate.

Response Modeling and Optimization

If the RSM team is experimenting in the region of a maximum or a minimum, a first-order regression model will not fit the data well. Instead, a second-order regression model will be needed. The most popular experimental design for fitting a second-order model is the central composite design proposed by

Box and Wilson [1] (see **Central Composite Designs**). This design includes points from a two-level factorial design, points along the “factor axes” and center points. Other popular designs are those developed by Box and Behnken [4, 5] (see **Box–Behnken Designs**).

The biotechnology study used a central composite design at the second stage; the results are listed in Table 2. Dramatic improvement is evident relative to the first phase of experimentation, as all settings now generate strong enzyme activity. The fitted second-order regression model, with all factors in coded form, is

$$\begin{aligned}\hat{Y} = & 8.13 + 0.56SCM - 0.04BP - 0.33YEP \\ & - 0.64SCM^2 - 0.19BP^2 - 0.01YEP^2 \\ & - 0.27SCM \times BP - 0.30SCM \times YEP \\ & - 0.63BP \times YEP\end{aligned}\quad (2)$$

There is a strong quadratic effect of SCM, a linear effect of YEP and strong interactions among all the factors.

There are a number of tools, both analytical and graphical, that can help us to better understand the nature of this fitted model. On the analytical side, second-order surfaces can be characterized algebraically. All have a “stationary point,” found by computing the partial derivatives of the surface with

respect to each factor, in turn, equating the partial derivatives to zero, and solving for the factor levels that simultaneously solve these equations. If the surface has an absolute minimum or maximum, it will occur at the stationary point. The second-order terms (pure quadratics and interactions) can be set up in a matrix. The eigenvalues and eigenvectors of the matrix show how the estimated response changes as the factors are moved away from the stationary point. The eigenvectors define new “axes” in the factor space. Moving along an axis with a positive eigenvalue, the value of the surface increases quadratically; the larger the eigenvalue, the steeper the increase. Similarly, if the eigenvalue is negative, the value of the surface decreases quadratically. The surface has a maximum (minimum) if all the eigenvalues are negative (positive).

Sometimes an eigenvalue is close to zero. This means that the surface has a “ridge”; i.e., a direction in the factor space along which the response stays nearly constant. Ridges can be of great value in a response surface study. They point to a large area in the factor space where similar results can be achieved. For example, suppose we want to maximize **process yield** and have run an experiment in which one of the factors is the viscosity of a blend produced at an intermediate stage of the production process, and another is the baking temperature at the next stage. The fitted regression model has an absolute maximum but also has a ridge corresponding to combinations of viscosity and temperature. Then any settings of viscosity and temperature along the ridge give only slightly lower estimated yield than is found at their optimal settings. This means that we can measure the viscosity after the intermediate stage and, if it is not at its optimal level, we can adjust the temperature to compensate, still achieving near optimal yield.

In the enzyme activity study, the stationary point of the fitted surface for the second experiment is at $SCM = 0.57$, $BP = -0.80$ and $YEP = 0.17$. The eigenvalues of the surface are 0.23, -0.31 and -0.76 . This means that there are two “directions” away from the stationary point in which the estimated surface decreases, i.e. in which the equation from the experimental data predicts lower enzyme activity. However, there is one direction in which there is an estimated increase in activity. That direction is given by the eigenvector that matches the positive eigenvalue, and has coefficients very close to 0, -0.6 and 0.8 for SCM , BP and YEP , respectively. Thus

the fitted surface predicts that we can achieve higher activity by keeping SCM at its stationary value of 0.57 , while modifying the (coded) BP and YEP in a ratio of -3 to 4 units. Had *all* the eigenvalues been negative, the fitted surface would have an absolute maximum at the stationary point.

A full discussion of the analytical properties of second-order response surface models is beyond the scope of this article; for details, see one of the books mentioned earlier. Ridge analysis is another useful method for exploring a second-order regression model; see **Ridge Analysis in Experimental Design**.

Graphic tools for visualizing response surfaces are available in many software packages. We used the contour plotting options in JMP® to explore the surface for the second phase of the enzyme activity experiment. The quadratic dependence on SCM is such that any value between 0 and 1 in coded units (150 – 180 in actual units) gives very good responses. The dependence on BP and YEP is related to the strong interaction between them. Figure 2 shows a contour plot of the estimated enzyme activity, as a function of YEP and BP , when SCM is set at 165 . The best estimated response values occur when YEP is set near 18 (-1 in coded units) and BP near 30 (1 in coded units). The increase in estimated activity as YEP and BP are moved from the center of Figure 2 to the lower right corner exactly matches the trend found by analyzing the eigenvalues of the fitted surface.

Figure 2 shows that the estimated response is very sensitive to the choice of YEP when BP is near 30 , but is rather insensitive to the choice of YEP when BP is set to lower levels (20 or less), with only a modest reduction in estimated enzyme activity. Sometimes this added “robustness” to inexact factor settings is sufficiently important that it may be worthwhile to move operating conditions away from the levels that appear to be “optimal” from the estimated regression.

Discussion

RSM is a strategy for sequentially exploring the relationship between a set of prospective process factors and one or more response variables. It has enjoyed great success in a wide variety of scientific and technological domains.

There are many statistical tools that have been developed under the umbrella of RSM. These include

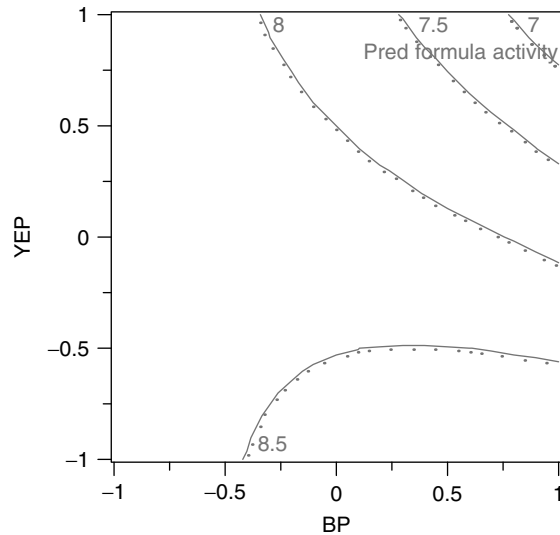


Figure 2 Contour plot of the estimated enzyme activity based on the 17 runs in Table 2. The factor BP is on the horizontal axis and the factor YEP is on the vertical axis; both are in coded units. The third factor, SCM, is set to 165, which gives high estimated activity for all values of the other factors

the central composite and **Box–Behnken** designs for second-order regression models, rotatability in design, fitting and exploring second-order (and higher-order) regression models and ridge analysis. Some researchers will write that they are carrying out a response surface study whenever they use any one of these tools. However, that ignores the very central role of sequential experimentation in RSM. The spirit of RSM is very much in the iterative interplay between the experimental team and the data found at each stage of exploration. As more experiments are run, the RSM team gets a better picture of which factors are important, of the useful ranges of those factors and of how the response depends on the factors in those ranges. In turn, that allows the team to design more penetrating and informative experiments, providing data that further expose the relationships between the factors and the responses. This iterative process continues until the experimenters have achieved the understanding needed to make improvements and achieve goals.

Further Reading

This short article has briefly described the basic concepts and methods of response surface methodology. Readers interested in more detail can find thorough accounts of RSM in the books by Box and Draper [10], Myers and Montgomery [11] and Khuri and Cornell [12]. RSM was the featured topic in the January 1999 issue of the *Journal of Quality Technology*. That issue includes an excellent review by Myers [13], summarizing recent research and setting out future directions, articles by Box and Liu [14] and Box [15] that discuss the role of RSM in the scientific learning process, and thoughtful discussions by several leading statisticians. Finally, readers can find an up-to-date review of RSM research in Myers *et al.* [16].

References

- [1] Box, G.E.P. & Wilson, K.B. (1951). On the experimental attainment of optimum conditions, *Journal of the Royal Statistical Society. Series B* **13**, 1–45.
- [2] Box, G.E.P. (1954). The exploration and exploitation of response surfaces: some general considerations and examples, *Biometrics* **10**, 16–60.
- [3] Box, G.E.P. & Youle, P.V. (1955). The exploration and exploitation of response surfaces: an example of the link between the fitted surface and the basic mechanism of the system, *Biometrics* **11**, 287–323.
- [4] Box, G.E.P. & Behnken, D.W. (1960a). Some new three-level designs for the study of quantitative variables, *Technometrics* **2**, 455–475.
- [5] Box, G.E.P. & Behnken, D.W. (1960b). Simplex-sum designs: a class of second order rotatable designs derivable from those of first order, *Annals of Mathematical Statistics* **31**, 838–864.
- [6] Box, G.E.P. & Hunter, J.S. (1957). Multifactor experimental designs for exploring response surfaces, *Annals of Mathematical Statistics* **28**, 195–241.
- [7] Box, G.E.P. & Draper, N.R. (1959). A basis for the selection of a response surface design, *Journal of the American Statistical Association* **54**, 622–654.
- [8] Kawaguti, H.Y., Manrich, E. & Sato, H.H. (2006). Application of response surface methodology for glucosyltransferase production and conversion of sucrose into isomaltulose using free *Erwinia* sp. Cells, *Electronic Journal of Biotechnology* **9**.
- [9] Plackett, R.L. & Burman, J.P. (1946). The design of optimum multifactorial experiments, *Biometrika* **33**, 305–325.
- [10] Box, G.E.P. & Draper, N.R. (2007). *Response Surfaces, Mixtures, and Ridge Analyses: Empirical Model-building and Response Surfaces*, John Wiley & Sons, New York.

-
- [11] Myers, R.H. & Montgomery, D.C. (2002). *Response Surface Methodology: Process and Product Optimization using Designed Experiments*, John Wiley & Sons, New York.
- [12] Khuri, A.I. & Cornell, J.C. (1996). *Response Surfaces*, Marcel Dekker, New York.
- [13] Myers, R.H. (1999). Response surface methodology – current status and future directions, *Journal of Quality Technology* **31**, 30–44.
- [14] Box, G.E.P. & Liu, P.Y.T. (1999). Statistics as a catalyst to learning by scientific method part I – an example, *Journal of Quality Technology* **31**, 1–15.
- [15] Box, G.E.P. (1999). Statistics as a catalyst to learning by scientific method part II – A discussion, *Journal of Quality Technology* **31**, 16–29.
- [16] Myers, R.H., Montgomery, D.C., Vining, G.G., Borror, C.M. & Kowalski, S.M. (2004). Response surface methodology: a retrospective and literature survey, *Journal of Quality Technology* **36**, 53–77.

Related Articles

Box–Behnken Designs; Center Points; Central Composite Designs; Fractional Factorial Designs; Interactions; Main Effect Designs; Main Effects; Plackett–Burman Designs; Ridge Analysis in Experimental Design; Rotatable Designs and Rotatability; Screening Designs; Sequential Experimentation; Supersaturated Designs.

DAVID M. STEINBERG AND RON S. KENETT

Rotatable Designs and Rotatability

Coding the Experimental Conditions

Two levels of a predictor variable can always be coded to -1 and $+1$. For example, if two temperatures $T = 125$ and 175°C were to be used, the **coding** $(T - 150)/25$ would achieve the coded levels. Note that the coded level zero (here 150°C) would be the center of the range selected. Typically the center of the entire design would represent the *current best set of conditions within the experimental region* and the experiment would attempt to find the best manner of moving outwards to improve the response. With two variables, say temperature and pressure, the coded levels of a 2^2 design would be $(x_1, x_2) = (-1, -1), (-1, 1), (1, -1),$ and $(1, 1)$ and this design would provide enough information to fit a first order model (a plane) in the two-dimensional experimental space. Adding one or more **center points** at $(0, 0)$ would enable us to create a general “**lack of fit** test for planarity”.

For a second-order design, additional experimental points would be needed to provide at least three levels for each variable x_i . These points would also be coded using the same coding formulas. In what follows, we assume the coding has taken place.

The Variance Function of an Experimental Design in General

Suppose, for the moment, that an experimental design with n runs has been performed and the data are available. We then fit a selected model, for example, a polynomial of first or second order, and then interpret the resulting fitted surface. In general regression notation we fit a model of the form

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon} \quad (1)$$

where the n rows of $\mathbf{X}$ all have the same general form but enough distinct values to permit **estimation** of the model under consideration. For a first-order model, the form of the u th row of $\mathbf{X}$ would be $(1, x_{1u}, x_{2u}, \dots, x_{ku})$ and, for a second-order model $(1, x_{1u}, x_{2u}, \dots, x_{ku}, x_{1u}^2, x_{2u}^2, \dots, x_{ku}^2, x_{1u}x_{2u}, x_{1u}x_{3u}, \dots, x_{k-1,u}x_{ku}), u = 1, 2, \dots, n$. The column vector

$\boldsymbol{\beta}$, written out here as a row to save space, is correspondingly labeled $(\beta_0, \beta_1, \beta_2, \dots, \beta_k)$ for first order or $(\beta_0, \beta_1, \beta_2, \dots, \beta_k, \beta_{11}, \beta_{22}, \dots, \beta_{kk}, \beta_{12}, \beta_{13}, \dots, \beta_{k-1,k})$ for second order. Let $r^2 = \mathbf{x}'\mathbf{x} = x_1^2 + x_2^2 + \dots + x_k^2$ where $\mathbf{x} = (x_1, x_2, \dots, x_k)'$ is a general point in the space of the predictor variables. Thus, r is the distance of a general point $\mathbf{x}$ from the coded origin $(0, 0, \dots, 0)$ of the design space. It can be shown that, at any such specified point, the variance of the fitted response is given by

$$V(\hat{y}) = \mathbf{z}'(\mathbf{X}'\mathbf{X})^{-1}\mathbf{z}\sigma^2 \quad (2)$$

where $\mathbf{z} = (1, x_1, x_2, \dots, x_k)'$ for a first-order model and $\mathbf{z} = (1, x_1, x_2, \dots, x_k, x_1^2, x_2^2, \dots, x_k^2, x_1x_2, x_1x_3, \dots, x_{k-1}x_k)'$ for a second-order model where σ^2 is $V(y)$, the variance of a single original observation y . (The σ^2 would usually be estimated from the residual mean square in an analysis of variance table associated with the fitting of the model.) In some presentations, a scaled form of equation (2), multiplied by n/σ^2 is also seen. This scaled form is called V_x , here.

The Concept of Rotatability

If the variance function $V(\hat{y})$ of equation (2) is a function of r^2 , rather than of the individual x values, in other words, if the variance function depends *only* on the squared distance that the general point $\mathbf{x}$ is from the origin of the design, then the design is said to be a *rotatable design*, one that provides response information of the same precision at all points equidistant from the design origin. One could rotate such a design around the origin to any angle (or equivalently rotate the axes around the origin) and the quality of information obtained on $\hat{y}$ would stay constant on spheres around the design origin. If the model being fitted were first order and the design used were rotatable, we would call the design a *first-order rotatable design*. Similarly, if the model were second order and the design used were rotatable, we would call the design a *second-order rotatable design*. Being rotatable is a very desirable characteristic of the design, but is not an absolutely essential one. However, if we are uncertain (as is typical in much experimental work) in what direction in the predictor space product improvements will be found, it is eminently sensible to gather information that is equally good in all directions at the same distance

2 Rotatable Designs and Rotatability

from the center of the design. Knowing how to attain *exact* rotatability allows an experimenter to obtain a design that is at least *approximately* rotatable while perhaps compromising on other design features. By “approximately rotatable,” an inexact concept here, we mean that contours of the variance function $V(\hat{y})$ may be “almost spherical” but not perfectly so. In picturing variance contours, we could plot the scaled variance V_x directly or, alternatively, plot the so-called information function $I_x = 1/V_x$. Note that when the value of V_x is low (a good feature), the information I_x is high. Figures 1 and 2, reproduced from [1, Chapter 14] and drawn originally by Conrad Fung, show information contours that are not rotatable and fully rotatable, respectively. Figure 1 arises from a 3^2 **factorial** design in two factors (x_1, x_2), indicated by the nine dots. Figure 2 arises from a central composite design (a square, plus four axial points plus four center points; see **Central**

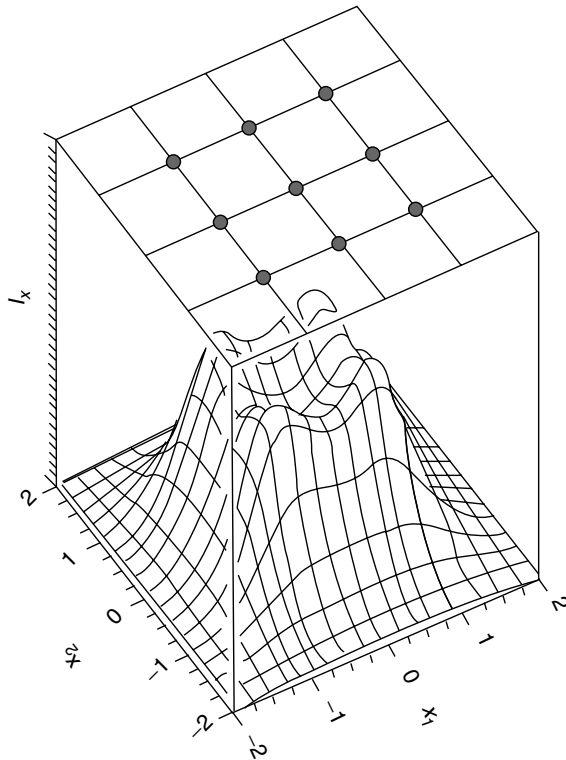
Composite Designs; the four added center points are drawn as separate dots for purposes of illustration only, of course).

Conditions for a Design to be Second-Order Rotatable

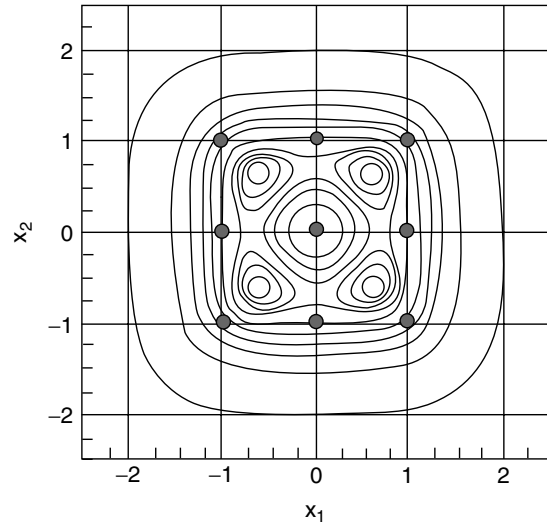
Suppose we have n coded design points $(x_{1u}, x_{2u}, \dots, x_{ku})$, $u = 1, 2, \dots, n$ in an experimental design. The quantity

$$\frac{1}{n} \sum x_{1u}^p x_{2u}^q \dots x_{ku}^t \quad (3)$$

(here and below, the summation is taken over the design points $u = 1, 2, \dots, n$) is called a *design moment* of order $p + q + \dots + t$. Then the design is rotatable of second order if and only if, for all i and j , where $i \neq j$, running from 1 through k ,



(a)



(b)

Figure 1 (a) Information surface for a 3^2 factorial arrangement used as a second-order design. (b) The corresponding information contours, formed by taking horizontal slices of (a). The contours are not circles and so the design is not rotatable

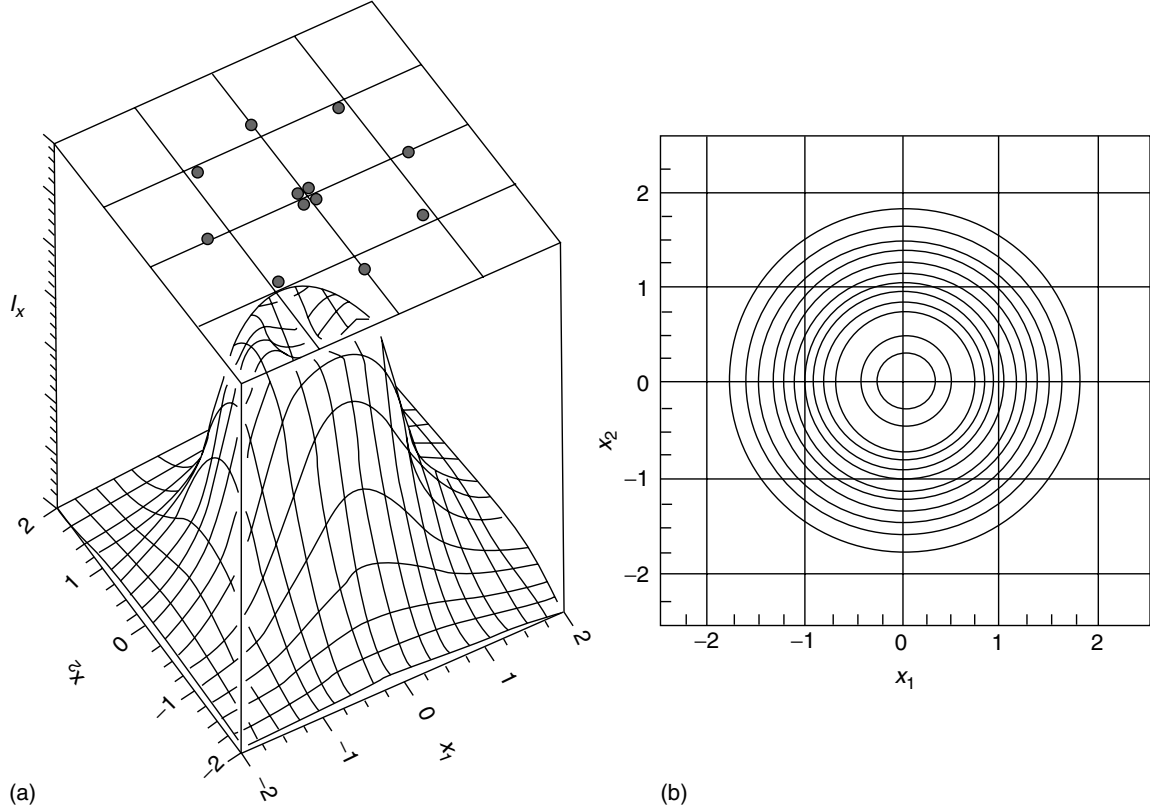


Figure 2 (a) Information surface for a cube plus star plus four center points arrangement used as a second-order design. (b) The corresponding information contours, formed by taking horizontal slices of (a). The contours are circles, and the design is rotatable

the following moment conditions (which define the quantities λ_2 and λ_4) hold:

$$\begin{aligned} \frac{1}{n} \sum x_{iu}^2 &= \lambda_2, \frac{1}{n} \sum x_{iu}^4 \\ &= 3\lambda_4, \frac{1}{n} \sum x_{iu}^2 x_{ju}^2 = \lambda_4 \end{aligned} \quad (4)$$

and provided that all other design **moments** up to and including order 4 are zero (see [2, p. 209]). To avoid singularity in the associated regression analysis, the ratio λ_4/λ_2^2 must exceed $k/(k+2)$, but this can always be achieved by adding center points, if needed. These conditions (equation 4) impose restrictive symmetries on the design in all the dimensions $(x_1, x_2, \dots, x_k)$. The relationship in equation (4) between the two types of fourth order moments (that the pure fourth moment must be three times as big as the mixed fourth moment) will be

applied below to achieve rotatability for central composite designs.

Conditions for a Design to be First-Order Rotatable

Only the conditions $\frac{1}{n} \sum x_{iu}^2 = \lambda_2$, plus “all other design moments up to and including order 2 must be zero” are needed. Thus all orthogonal first-order designs with the same pure second-order moment in all dimensions are first-order rotatable.

Some First-Order Rotatable Designs

A sensible design for fitting a first-order model is a two-level factorial or two-level **fractional factorial** of resolution III or better (i.e., of higher resolution)

(see **Factorial Experiments; Fractional Factorial Designs; Main Effect Designs; Plackett–Burman Designs**) Such a design will estimate all first-order coefficients although, in the case of a resolution III design, some or all main effects will be aliased (confounded, confused; see **Aliasing in Fractional Designs**) with possible two-factor interactions that may exist. Such a design is first-order rotatable, that is, it provides information on the first-order response with the same accuracy at the same distance from the center of the design, in other words on circles/spheres about the origin. For a first-order model, any orthogonal design with all pure second-order moments equal to the same number λ_2 , has a diagonal $\mathbf{X}'\mathbf{X}$ proportional to the unit matrix $\mathbf{I}$, except for the first entry, and is first-order rotatable. In fact, if the n design points are at $(\pm 1, \pm 1, \dots, \pm 1)$, then $V(\hat{y}) = \sigma^2(1 + r^2)/n$ and so $V_x = 1 + r^2$.

Rotatability of the Central Composite Second-Order Design Class

In order to fit a second-order model, at least three levels of each x_i must be chosen in forming a suitable design. A popular and practical design choice (see **Central Composite Designs**), which can be found in many published experimental presentations, is the so-called central composite design which provides four levels (or five if, as is usually done, *center points* are added), in general. A central (i.e., centrally located about the coded design origin) composite design consists of three parts:

1. a “cube” part, that is, a two-level n_c points portion $(\pm 1, \pm 1, \dots, \pm 1)$, usually of resolution IV or higher (which means that no **main effect** is aliased with any two-factor interaction); plus
2. a “star” part, consisting of $2k$ axial points of the forms $(\pm\alpha, 0, 0, \dots, 0)$, $(0, \pm\alpha, 0, \dots, 0)$, $(0, 0, \pm\alpha, \dots, 0)$, $\dots$ $(0, 0, 0, \dots, 0, \pm\alpha)$, where α is not yet specified; plus
3. n_0 center points $(0, 0, \dots, 0)$.

(If the star distance is chosen to be $\alpha = 1$, the design is often called a *face-centered central composite design*.)

For such a design, $\sum x_{iu}^2 = n_c + 2\alpha^2$, $\sum x_{iu}^4 = n_c + 2\alpha^4$, and $\sum x_{iu}^2 x_{ju}^2 = n_c$. Thus, to satisfy equation (4), we set $n_c + 2\alpha^4 = 3n_c$, whereupon $\alpha = n_c^{1/4}$. Some examples of this follow, with the asterisk

denoting that a half fraction of the 2^k design is being used as the “cube” portion of the design:

k	2	3	4	5	5*	6	6*
n_c	4	8	16	32	16	64	32
α	1.414	1.682	2	2.378	2	2.828	2.378

As mentioned previously, Figure 2(a) shows a rotatable composite design in two dimensions with four center points; we see now that $\alpha = 1.414 = 2^{1/2}$. (It happens in this case that the noncentral design points lie on a circle, a two-dimensional sphere; this sphericity feature is *not* true for other examples, as we see from the α numbers shown for other dimensions.)

The literature of experimental design has many papers about rotatable designs and rotatability of orders 1–4. See, for example, the bibliography in [1]. All such designs are technically interesting, but some are too big for typical experimental work. For $k = 4$ and 7, useful rotatable designs exist in the series of Box–Behnken three-level designs (see **Box–Behnken Designs** or [1, Chapter 15] or [3]).

Measuring Second-Order Rotatability for a Design

For comparing two or more second-order experimental designs, a *measure of second-order rotatability* might be useful. Several suggestions appear in the literature; see [4–8]. We recommend the measure in [6], because it does not depend on the orientation of the design in the factor space. Comparisons of the recommended measure against competing measures appear in [6, 8].

Higher-Order Rotatability of a Design

If models of third or fourth order were of interest, rotatable designs of those orders could be used. There are a number of papers in this area, listed in the bibliography of [1]. See also [9] and the references therein. However, with rare exceptions, experimenters usually feel that the large experiments required are impractical. It is usually preferable to fit a second-order model in a moderate size area of the predictor variables space, use **ridge** analysis to move to a better area, and then refit a second-order model there to new data from another experiment.

Graphical Assessment of Response Surface Designs

Consider Figure 2(a) again and think of it as a cake. Imagine a conventional triangular slice cut from this cake. The top edge of such a slice will have an I_x profile that could be plotted in two dimensions against the radius $r = (x_1^2 + x_2^2 + \dots + x_k^2)^{1/2}$. Such a curve could be compared with a similar curve of another competing rotatable design. Alternatively, we could plot the inverse function of I_x , namely V_x , in two dimensions and make comparisons on that basis. Such a plot of V_x is called a *variance dispersion graph* (VDG). When the design is not rotatable, the edges of the cake slices will not all have the same profile, of course, but it would still be useful to look at curves of maximum, minimum, and “spherical” average values of V_x versus r to compare designs (see [10, 11]). Such curves also provide a way to compare, for central composite designs (see **Central Composite Designs**), various choices of the star distance α and also the number of center points. For a similar approach using the **mean squared error** of prediction, which includes model **bias** errors as well as variance errors (see [12]).

References

- [1] Box, G.E.P. & Draper, N.R. (2007). *Response Surfaces, Mixtures and Ridge Analyses*, 2nd Edition, John Wiley & Sons.
- [2] Box, G.E.P. & Hunter, J.S. (1957). Multifactor experimental designs for exploring response surfaces, *Annals of Mathematical Statistics* **28**, 195–241.
- [3] Box, G.E.P. & Behnken, D.W. (1960). Some new three-level designs for the study of quantitative variables, *Technometrics* **2**, 455–475; Corrections, 1961 3: 576.
- [4] Khuri, A.I. (1988). A measure of rotatability for response surface designs, *Technometrics* **30**, 95–104.
- [5] Draper, N.R. & Guttman, I. (1988). An index of rotatability, *Technometrics* **30**, 105–111.
- [6] Draper, N.R. & Pukelsheim, F. (1990). Another look at rotatability, *Technometrics* **32**, 195–202.
- [7] Kshirsager, A.M. & Cheng, J.C. (1996). A new measure of rotatability of response surface designs, *Australian Journal of Statistics* **39**, 83–89.
- [8] Draper, N.R. & Pukelsheim, F. (1997). Kshirsager and Cheng’s rotatability measure (letter to the editor), *Australian Journal of Statistics* **39**, 236–238.
- [9] Draper, N.R. & Pukelsheim, F. (1994). On third-order rotatability, *Metrika* **41**, 137–161.
- [10] Giovannitti-Jensen, A. & Myers, R.H. (1989). Graphical assessment of the prediction capability of response surface designs, *Technometrics* **31**, 159–171.
- [11] Myers, R.H., Vining, G.G., Giovannitti-Jensen, A. & Myers, S.L. (1992). Variance dispersion properties of second-order response surface designs, *Journal of Quality Technology* **24**, 1–11.
- [12] Vining, G.G. & Myers, R.H. (1991). A graphical approach for evaluating response surface designs in terms of the mean squared error of prediction, *Technometrics* **33**, 315–326.

Related Articles

Aliasing in Fractional Designs; Box–Behnken Designs; Central Composite Designs Factorial Experiments; Fractional Factorial Designs; Main Effect Designs; Plackett–Burman Designs.

NORMAN R. DRAPER

Sequential Experimentation

Introduction

R.A. Fisher, the originator of statistical experimental design, once remarked that “The best time to design an experiment is after it’s done”. Wry, but certainly true: the more one knows about what factors are important, what extra-experimental sources of variation might **bias** results, which measurements are difficult to make (i.e., are noisy/imprecise) and which are not – in short, the more one knows about the system under study – the better one can focus an experiment to learn more about it.

There is nothing particularly profound in this, of course. Even the justly maligned but widely used practice of studying many factors by holding all but one constant and varying just this one (OFAT: one-factor-at-a-time experimentation) acknowledges this truism. Each experimental factor is varied to its “optimal” setting, often by looking at whether the last change in the factor made results better or worse. Once the “best” setting for a factor has been found, it is then held there and the next factor of interest is then similarly examined. Of course, “optimal” and “best” must be qualified by scare quotes, as this strategy ignores both experimental noise and the possibility of interactions among the factors, issues that will be discussed further shortly. The bottom line is that more careful approaches are needed.

In any industrial or business process, there are always potentially many factors that might affect the response(s) of interest. The classic **Ishikawa “fish-bone” diagram**, Ishikawa [1], is one well-known quality tool that can be used by those who share process knowledge to identify from dozens to literally hundreds of methodological details, raw material variations, equipment settings, environmental factors, and even maintenance procedures that might impact results.

As a canonical example, consider a measurement process to measure “purity” (of a chemical, food product, or drug, say). Such a process requires a procedure for obtaining and preparing samples to be measured, for storing them over the course of the preparation, to prevent outside contamination or chemical reactions with the environment, and for

obtaining a readout of the prepared samples in an instrument. Typically, the protocols defining these procedures contain many steps – how to sample, how to handle the samples, what reagents to use, what times and temperatures the samples can be stored under, how the instrument is to be set up, calibrated, maintained, and so forth. A long list of individual factors that might affect the process would result.

In reality, for a process to work at all, prior knowledge and experience will generally dictate how to control most, if not all, of these many details satisfactorily. When this is not the case, typically a relatively small subset of process variables – less than, say, 15 or so and perhaps only 1 or 2 – can be identified for empirical investigation and adjustment, that is, experimentation. That is the focus of this article.

If “first principles” and prior experience have whittled the number of process variables down to at most one or two, it is generally not difficult to come up with a reasonably efficient experimental strategy that achieves success. Here, “efficient” simply means with an acceptable expenditure of time and effort. It is important to note, however, that “efficiency” can be defined in a statistical way that formalizes this idea of minimizing experimental effort in order to design “optimally” efficient experiments (*see* **Optimal Design**). While this theory turns out not to be quite so “optimal” in practice, it nevertheless yields very useful insights and approaches that can greatly reduce experimental effort. The strategy of sequential experimentation implicitly uses much of this methodology. So while no claims of statistical optimality will be made, a major advantage of using this approach is that it achieves success with less experimental effort.

But what happens if one cannot identify one or two variables that can be manipulated and cannot achieve a satisfactory result in those that are used. This is often the case, for example, when there are several competing criteria that together constitute “satisfactory” performance: it may be easy to optimize the criteria individually but difficult to find a compromise that is satisfactory for all.

However it comes to pass, when there is a need to investigate the settings of more than just one or two process variables, it is no longer so obvious how to proceed. Following Fisher’s observation, the experimental strategy should be a frugal one. That is, whenever possible (it isn’t always!), one would prefer to gradually accumulate understanding through a sequence of relatively small experiments, rather than

2 Sequential Experimentation

gambling everything on a large, one-shot effort. Not only does this permit later experiments to take advantage of what has been learned, it also allows the experimenter to recover from “mistakes” – misjudgments about what variables should be investigated and over what ranges, for example – before experimental resources (or managerial patience!) have been exhausted. No experimenter who is truly on the frontier of knowledge can avoid making such mistakes, so it makes sense to employ a strategy that quickly reveals and recovers from errors so that more fruitful directions can be pursued.

An Outline of the Strategy

There are three principles that underpin the sequential experimentation strategy.

1. Exploit Occam’s razor, which is an early version (from the fourteenth century English logician and friar, William of Ockham) of the familiar *Pareto principle* that asserts that in any process, about 80% of the process variation is caused by about 20% of the process variables. In this context, it means that one can expect relatively few (sometimes none!) of the experimental variables being considered to meaningfully affect the response. By employing a strategy to efficiently identify and quantify the effects of just these few, one can learn how to achieve the desired performance with great experimental economy. Because time and resources are always limited, this is often crucial.
2. Distinguish signal from noise. To be useful, an experiment must yield results that are replicable. In other words, to infer that observed changes in a response were actually caused by changes in an experimental variable – which can then be used to control the response – those changes must not in fact have been due to uncontrolled and unknown sources of experimental variability. Of course, this is a fundamental precept of science, not merely of statistical experimental design.
3. Allow for interaction (*see Interactions*). Factors often affect results interactively. Example: the concentration of glucose needed for best growth of a cell culture depends on the concentration of certain amino acids present in the growth medium. One cannot specify how glucose affects growth without simultaneously specifying how

much amino acid is present (and conversely). This kind of behavior is clearly not uncommon, and so it is important that multifactor experiments allow for such possibilities (*see Factorial Experiments* for further discussion).

In many respects, these principles seem almost self-evident; but in fact, OFAT experiments violate all of them! They fail the first, because equal effort must be expended to investigate each experimental factor; there is no “information sharing” that allows, for example, *hidden replication* (*see Fractional Factorial Designs*) or elucidation of interactions in the space of *active factors* (various aspects of which are discussed in *Aliasing in Fractional Designs*; *Foldover Designs*; *Fractional Factorial Designs*; *Plackett–Burman Designs*; *Projectivity in Experimental Designs*). They fail the second, because they make poor use of experimental information and therefore are only reliable in sufficiently high signal/noise situations. This can only be achieved by (appropriately!) replicating the individual factor changes a large number of times, thereby incurring excessive experimental effort. They fail the third because interactions cannot be “seen” when only OFAT is changed: by definition, interactions require the simultaneous change of two or more experimental factors to evidence themselves (*see Interactions*).

To gain a rough sense of how ineffective OFAT can be compared to the strategy outlined here, suppose that one has 10 experimental factors to study. By 1, it is reasonable to expect that no more than two or three of them will have large enough effects – either individually or jointly (interactions) – to be of concern. With this in mind, one could use a so-called 12-run Plackett–Burman design (named after the two British mathematicians who discovered it in 1946 – *see Plackett–Burman Designs*), perhaps with a few added center points (*see Center Points*) to determine what the important factors were and how they affect the response of interest. An additional few runs or foldover (*see Foldover Designs*) are occasionally required to resolve possible uncertainties that might exist (e.g., due to aliasing, *see Aliasing in Fractional Designs*). So perhaps 15–20 experimental runs might be required to distinguish the few most active factors (if any!) from noise (2) and gain a general sense of their individual and joint effects (3).

By contrast, to gain the same information on the individual effects using the OFAT strategy, a

minimum of 120 experimental runs (and some center points, perhaps) are required. This is roughly six to eight times the experimental effort! Moreover, even with this large number of runs, there is no information about any possible interactions!

The strategy outlined here was developed over the course of about 70 years (starting about 1920) by Fisher, Frank Yates, Oscar Kempthorne, Gertrude Cox, George Snedecor, George Box, and numerous other statisticians. Interestingly, most of the work prior to 1950 had its roots in agricultural experimentation to improve crop and livestock yields, reduce pests, and so forth. More recent (post-1990) developments by Box, John Tyssedal, Ching-Shui Cheng (*see Projectivity in Experimental Designs*), and others make use of the mathematical properties of Hadamard matrices and other two-level orthogonal arrays (*see Orthogonal Arrays*) originally discovered by the French mathematician, Jacques Hadamard and further developed by Rao and others in the first half of the twentieth century.

The strategy proceeds in stages as follows:

1. Use small two-level screening designs (*see Screening Designs*) to identify the most important few factors of interest and roughly quantify how they affect the response. Fractional factorials (*see Fractional Factorial Designs*) and other designs of projectivity 3 or 4 (*see Plackett–Burman Designs; Projectivity in Experimental Designs*) can be used. Depending on which, full or partial aliasing may require a small number of carefully chosen follow-up runs to resolve any confounding and clearly identify the influential factors. While there is statistical methodology available to guide this choice, for example, using foldover (*see Foldover Designs*) or as in Meyer *et al.* [2] and Box *et al.* [3], in practice, this is rarely necessary. Subject matter expertise and the hierarchy of effects – which asserts that main effects are more likely than two-factor interactions (2 fi's), 2 fi's are more likely than 3 fi's, and so forth – are usually sufficient.
 2. Stage 1 identifies the important factors and quantifies their effects. This provides an approximate direction for improvement that can now be rapidly explored by running several experimental runs in this direction. This is essentially an empirical version of the *method of steepest ascent* (*see Response Surface Methodology*).
- For example, suppose two factors of 10, A and B, have been identified, and the projection of the fitted **response surface** (equivalently, the fitted model coefficients) indicates that the response will improve most rapidly if A is increased and B is decreased in approximately a 2:1 ratio in standardized scaling units (i.e., units in which the ranges of all experimental variables are -1 to $+1$). Suppose the low and high levels of A and B in stage 1 in their natural (unstandardized) units are l_a, l_b, h_a , and h_b respectively, and δ_a and δ_b are $h_a - l_a$ and $h_b - l_b$. Then in stage 2, the values of the other eight experimental factors would be held constant (at the centers of their range or similar “standard” values chosen for economic or other reasons), and A and B pairs of the form $m + t\Delta$ would be run, where $m = ((h_a + l_a)/2, (h_b + l_b)/2)$, the center of the stage 1 design in A and B, and $\Delta = (2\delta_a, -\delta_b)$. To control for the possibility of an additive shift in response between the stages, typically one would include one or more replicates of the center point, $t = 0$, in this sequence. Note that this can be considered a one variable design in the *constructed variable*, Δ . However, unlike OFAT, the variable was constructed based on experiment – it reflects what nature has to say, not prior ignorance.
3. As t is increased, the response may improve for a while, but eventually one will go too far, and it will crash or various physical or economic constraints will make further increments in t impossible. At this point, there are several possibilities.
 - (a) If results are deemed to be satisfactory, you're done: declare victory and move on.
 - (b) A new design region in A and B will have been identified. With A and B at these new levels some of the other eight previously unimportant factors may now have meaningful effects. This is equivalent to the existence of interactions of A and B with the remaining factors over the wider range of A and B from those of stage 1 to the current levels.^a Hence one might wish to return to stage 1 and repeat a screening design in the original factors with A and B either held constant or varied around their new levels.
 4. Stages 1–3 may be repeated until either no further improvement results or other constraints

(technological, economic) prevent further change in the variables. At this point, one will typically have completed a final two-level screening design in the final optimal region and identified a relatively small number of variables that affect the response in this region. At this point, it may be necessary to explore the region of the optimum in greater detail to determine the sensitivity of the response to changes of the critical variables/validate ranges within which to hold these variables to keep the response within desired limits. In currently fashionable language, this would constitute a kind of process *robustness* study (see **Robust Design**; **Tolerance Design**). A wide variety of approaches to do this have been proposed (see **Performance Measures for Robust Design**; **Product Array Designs**; **Tolerance Design**; **Box–Meyer Method for Dispersion Effects**, for example), but the classical and most widely used approach is to add *axial points* to the existing screening design to create a second-order *central composite design* (see **Central Composite Designs**). The fitted response surface can then be both graphically and mathematically explored using, for example, ridge analysis (see **Ridge Analysis in Experimental Design**) to provide often revealing insight into the behavior of the system under study.

A great virtue of this approach is *sequential assembly*, Box *et al.* [3] (see **Central Composite Designs**): the axial and cube portions of the full second-order design comprise separate *blocks* (see **Blocking**), and so any additive shift that may occur is confounded with blocks. Running center points in both the earlier cube portion and the newer axial portion would allow **estimation** of such a shift; but even if this were not the case, both the first- and second-order (including interaction) variable effects would not be corrupted by a systematic shift in response levels. In other words, valid estimates of the effects – and therefore of the response surface – would be obtained even if there were a shift in system behavior between the times the cube portion and axial portion were run. Some of the **robust design** methods – for example, product array designs (see **Product Array Designs**) – that have been proposed do not easily fit into this paradigm and require large designs run *de novo*. While there are certainly situations where this is justifiable,

it may come at a considerable cost. Sequential assembly attempts to avoid those costs.

Although this overall strategy may seem complex, in concept it is actually rather simple. Here is a little constructed example and picture that shows how it works as compared to standard OFAT experimentation. Suppose there is interest in reducing the amount of a certain contaminant – a reaction by-product – that is formed along with the main reaction product in a biochemical reaction. The contaminant is measured as a percentage of the main peak by ion exchange chromatography, and engineers believe that changes (typically reductions) in reaction time and concentration of a reagent should help reduce it. Figure 1 shows a typical OFAT approach: the initial setting has a time of 9.3 s and a concentration of 11.8. The first set of experiments, shown as white circles, vary the time up and down while holding the concentration fixed. The true, but of course unknown to the experimenters, behavior of the system is shown by the shaded contamination contour plot. If results are very precise, these results would show a slight reduction in contamination by reducing the time to about 8.4 s, as indicated by the X on the horizontal line. Then a second set of experiments would be conducted holding the time fixed at this “optimal” value and varying the concentration as shown by the black circles. The optimal concentration would then be found to be about 10.7, indicated by the second X.

While significant improvement results – the off-peak percentage is reduced from around 2.5 to around 1% – one could still clearly do better. Figure 2 shows how. The 2^2 design indicated by the four white circles is run instead. Analysis of these results would show that the best direction for improvement is that indicated by the arrow. A series of subsequent trials in this direction would clearly lead to the optimal region at the lower left. Incidentally, the key feature of this process is that the time and concentration have an interactive effect on the percentage contamination. The 2^2 design can “see” this interaction, but the OFAT experiment cannot.

Of course, there are both variations and enhancements that may be required to handle special situations, some of which will be mentioned in the next section. But these are technical details. One should not lose sight of the fact that the underlying strategy is a straightforward implementation of the three basic principles that are essential to sound and efficient

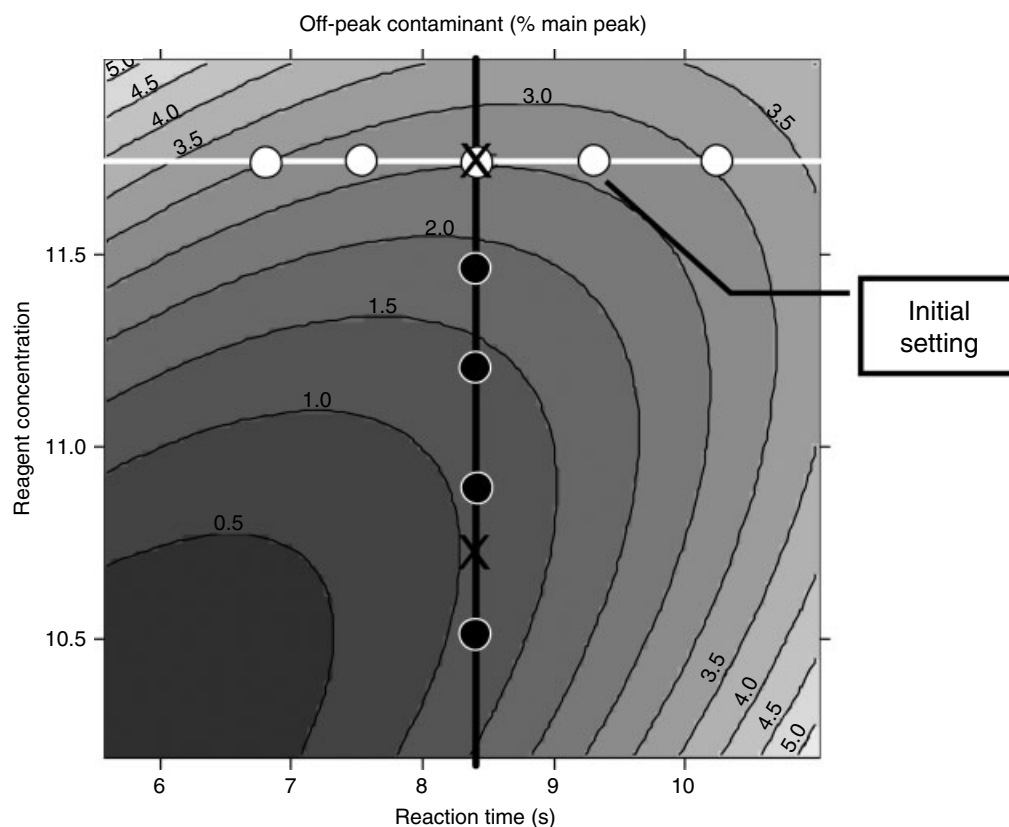


Figure 1 Graphical representation of a one-factor-at-a-time experiment. First, time is varied holding concentration constant; then concentration is varied holding time at the previously found best setting. This leads to the false optimum indicated by the X on the vertical line

experimentation. This is why the strategy has such wide potential application and value.

Comments and Caveats

Given the development and claims of the preceding section, one may well ask: “If the sequential experimental strategy is so effective, why is it so little used in the world of science and technology (as any cursory review of the literature demonstrates)?” While a full discussion of this point could probably fill several volumes,^b it is perhaps more reasonable to frame the issue in terms of the practical problems that one encounters when trying to execute such a strategy. Here are several that come to mind.

1. Screening designs with many factors are complex and difficult to do. They require that the levels

of many variables be simultaneously changed, which can be difficult or even impossible in practice. For example, in an experiment to investigate the effect of time after addition of a reagent and the temperature at which the next step of a chemical process takes place, an exothermic reaction may make it difficult to cool the solution down to a desired low temperature if the time after the reagent addition is too short. The logistics of executing simultaneous changes in many variables may also be difficult in some situations, especially when programming automation software is involved.

2. Two-level screening designs and steepest ascent cannot be used when there are categorical factors with more than two categories that need to be tested. While there are sometimes screening-type designs that can accommodate this situation

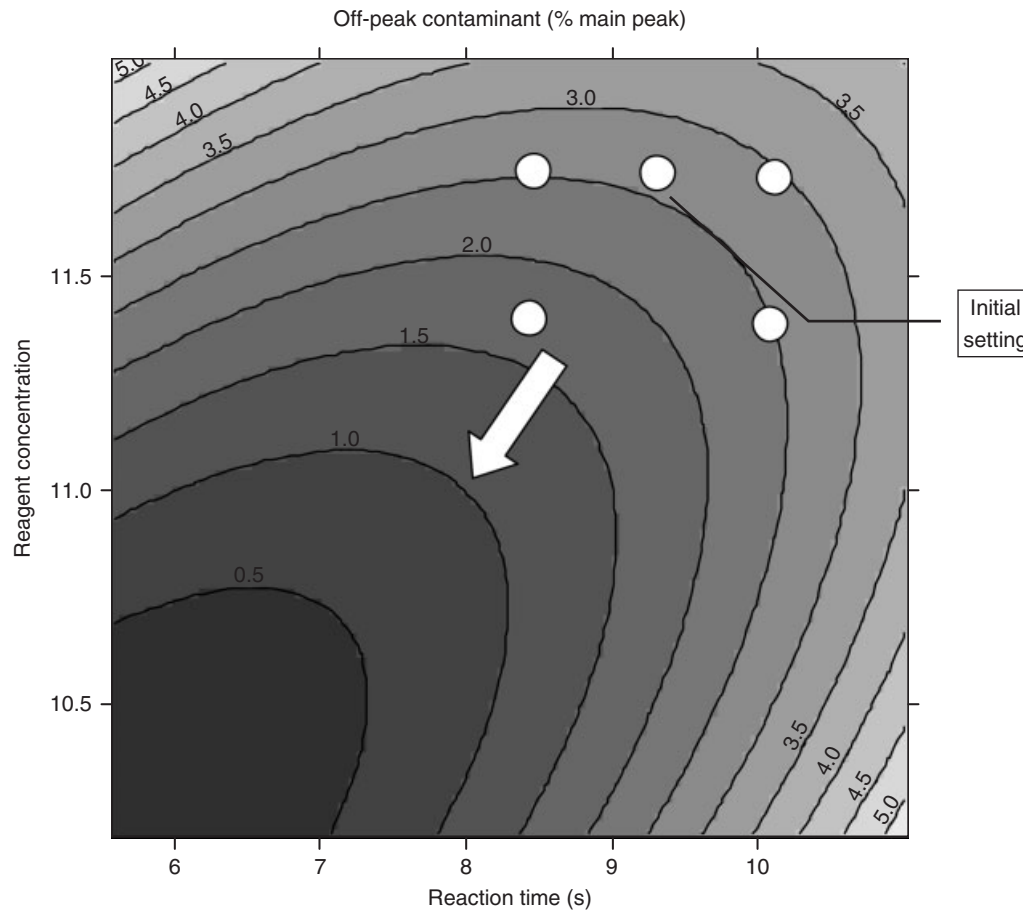


Figure 2 Graphical representation of the sequential experimental strategy. An initial factorial experiment, indicated by the four white circles around the initial setting, shows that the optimal direction of improvement indicated by the arrow. Subsequent experiments in this direction would lead to the true optimum

(see **Orthogonal Arrays**), it is the nature of categorical variables that to understand what happens at any combination of categories, you have to test all possible combinations. For example, to test the possible effects and interactions of fixture type, product type, and bearing material on the performance of an assembly, you have to test all possible combinations of fixture, product type, and bearing material. When there are many such factors with many categories each, factorial experiments (see **Factorial Experiments**) may become large and unwieldy.

3. It can be very difficult to determine workable ranges for all the variables in screening experiments with many factors. One reason for this is

the so-called curse of dimensionality: the volume of the experimental space increases exponentially with the number of variables. So, for example, if one has only two factors to study, in standardized coordinates the corner points where both factors are at their extremes are only about 40% farther from the origin than individual ranges ($\sqrt{2}$ vs 1); but with 10 factors they are more than three times more distant ($\sqrt{10}$ vs 1)! In any real process in which each of the factors are expected to vary more or less randomly and independently within their individual ranges, it is highly unrealistic to expect more than a few of the factors to be simultaneously at one of their extremes. While effects sparsity (i.e., Occam) may be invoked,

in practice such high-order “cube” points are sometimes obviously unworkable. A statistical restatement of this is that high-order interactions are important in this situation. A simple example is when the amounts of two sugars, glucose and mannose, must be varied to determine how to improve the yield of a cell culture. Individually amounts can be zero, but if both are zero simultaneously – the $(-, -)$ level in a simple 2^2 design – the culture cannot grow.

This reasoning suggests that one might want to make the ranges of experimental factors less extreme when there are more of them than when there are fewer. But how much less extreme? And does this exclude settings in which only one or two factors are extreme that should be investigated? See Voelkel [6] for some related ideas. The problem is that there is a lot of volume in which to experiment, and screening designs are very dependent on Occam in such circumstances just to establish workable initial ranges. One approach is to run preliminary experiments at potentially “dangerous” settings, but this can be time consuming and *ad hoc*.

4. One of the cornerstones of experimental design is randomization (see **Randomization in Experimental Designs**). But more often than not it is difficult or impossible to fully randomize. As an example, suppose one is interested in optimizing a plate-based assay for a drug metabolite in blood (these are usually required for **clinical trials** of drugs). Such assays are conducted in 96-well plates, about $10\text{ cm} \times 15\text{ cm}$ in size, in which the wells are arranged in a 8-row $\times$ 12-column matrix. Factors that one would be interested in investigating typically include well-specific components of the biochemistry (pH, buffer type, reagent concentrations) that can be changed well by well within each plate (perhaps with some effort) and plate-specific components like time and temperature of incubation that can only change from one plate to another. It is impossible to fully randomize such experiments, because this would have to occur at the well level, and times and temperatures cannot be changed at that level.

Such an experiment would have to be conducted as a split-plot design (see **Split-Plot Designs**), but with many factors, such designs

are not simple to plan or analyze. When experiments are conducted as split-plot designs but analyzed as if they were completely randomized, it is quite possible to reach incorrect conclusions. For example, conducting a two-level unreplicated **fractional factorial** as a split plot would actually mean that there are no **degrees of freedom** for the main effects of the whole plot factors (however unaliased interactions with subplot factors could be estimated). The simple and appealing graphical analyses described elsewhere in this encyclopedia (see **Half-Normal Plot; Lenth’s Method for the Analysis of Unreplicated Experiments**) are completely inappropriate. By contrast, if sufficient replication is done and main effects predominate, OFAT would work.

5. One of the strengths of screening designs is that they are very efficient: each experimental run provides a lot of information. This is just another way of describing why OFAT designs with many factors require so many more runs to gain equivalent information (and even then only on the main effects). But this blessing of high efficiency can also be a curse: if a run is “lost” due to some kind of experimental glitch – a not uncommon occurrence – a lot of information can be lost with it. Perhaps worse, if some kind of experimental problem leads to a severely corrupted result, *but the experimenter is unaware that this has happened*, the single bad value can distort all the estimates in an unreplicated design. In OFAT, losing a run here or there makes little difference, and a corrupted value can distort only the single factor being varied. When there is less information to be gained, there is less information that can be lost.

This litany of problems may seem to be a fatal indictment of the sequential experimentation strategy. Not so. As indicated above, there are practical approaches to deal with all these issues. But it is important that any exposition or advocacy of statistical experimental design methods deal fully and fairly with the real circumstances that one faces in using them. As a rule, the statistical experimental design literature does not concern itself with such vexing realities. Thus, researchers who encounter them either knowingly or unknowingly when attempting to apply the methods often do not achieve the results that were advertised,

or may even not achieve useful results at all. Attention to detail is key, but without being told about these details, no attention can be paid.

Final Comments

Unfortunately, even this exposition still fails to touch on many important aspects of the sequential design strategy. One of those alluded to in the introduction is that of multiple responses – it is rare that only a single response is of interest; most of the time there are several that “compete” with each other for optimal settings. Often, this is *the* essential problem driving the investigation. Examples: it is easy to increase yield at the **cost of quality** or to improve quality but reduce yield; but it is difficult to do both. It is easy to improve assay precision at the cost of sensitivity or to increase sensitivity at the cost of precision; but it is difficult to do both. One of the real strengths of the sequential design strategy is that it often makes manifest how such trade-offs should be made. Even better, it often provides ways to avoid having to make the trade-offs at all (*see Dispersion Effects; Performance Measures for Robust Design; Product Array Designs; Box–Meyer Method for Dispersion Effects* for some aspects of this in robust design).

Another important omitted issue is the use of response transformations to both simplify the statistical model and provide insight into underlying scientific behavior. Box *et al.* [3, 7] provide details and examples.

Another topic that underlies the practical implementation of experimental design is the interplay among empirical experimentation, theory, and prior experience. Statistical experimental design is usually expository as if all knowledge were in the data, but this is never even close to the case. Experimenters typically work within a firmly established theoretical framework and have a lot of prior experience with the system under study or ones very similar to it. As a consequence, clear results from a well-designed screening experiment often provide sufficient insight to allow the experimenter to cut short further investigation and successfully apply this subject matter knowledge to directly achieve the results that were sought or to recognize inherent constraints that would make further experimentation useless. George Box once described this in an experimental design class

as the ability of statistical design “to catalyze the scientific learning process”. See Box [8] for further discussion.

Finally, like it or not, the strategy of sequential experimentation is inextricably linked to the philosophy of science: how do theory and experiment interact in the “scientific method”? – what constitutes a “replicated” result? – when is a hypothesis confirmed or contradicted by inevitably noisy data? – when is an experimental result sufficiently unexpected to warrant detailed checking or redoing? Because experimentation is an active process to demonstrate **causality** and not merely to assert passive association for the purpose of prediction, it must deal with such questions. This is perhaps the most interesting and even useful aspect of sequential design, because it focuses on how experimentation alters our understanding of how the world works. Unfortunately, such a discussion is beyond the scope of this article, but perhaps this brief mention will pique the reader’s curiosity and stimulate further thought and study.

End Notes

^aAnother equivalent interpretation is that there is meaningful curvature of the response surface over the wider ranges. Interaction is a second-order effect (*see Response Surface Methodology*).

^bIndeed, it already has. See, for example, Kahneman *et al.* [4] and Kuhn [5] for two possible aspects of the issue.

References

Note: It should be evident from the text that practically any reference for experimental design is also relevant to sequential experimentation. Hence, the following should be viewed merely as particularly germane to some of the issues discussed, and not a comprehensive or even sufficient compendium.

- [1] Ishikawa, K. (1986). *Guide to Quality Control*, Asian Productivity Organization, Tokyo.
- [2] Meyer, R., Steinberg, D. & Box, G. (1996). Follow-up designs to resolve confounding in multifactor experiments, *Technometrics* **38**, 303–313.
- [3] Box, G. & Draper, N. (2007). *Response Surfaces, Mixtures, and Ridge Analyses*, 2nd ed., John Wiley & Sons, Hoboken.
- [4] Kahneman, D., Slovic, P. & Tversky, A. (eds) (1982). *Judgment Under Uncertainty: Heuristics and Biases*, Cambridge University Press, Cambridge.

- [5] Kuhn, T. (1996). *The Structure of Scientific Revolutions*, 3rd Edition, The University of Chicago Press, Chicago.
- [6] Voelkel, J. (2005). The efficiencies of fractional factorial designs, *Technometrics* **47**, 488–494.
- [7] Box, G., Hunter, J.S. & Hunter, W.G. (2005). *Statistics for Experimenters: Design, Innovation, and Discovery*, 2nd Edition, John Wiley & Sons, Hoboken.
- [8] Box, G. & Friends. (2006). *Improving Almost Anything: Ideas and Essays*, Revised Edition, John Wiley & Sons, Hoboken.

BERT GUNTER

Dispersion Effects

Introduction

Reducing variation in the performance of industrial products and processes has always been a very important goal. Low variation makes processes predictable, reduces costs, and increases customer satisfaction, among other benefits. Accordingly, variation control and reduction have been the aims of many methodologies: from statistical process control (SPC) (*see Control Charts, Overview*) that boomed after the Second World War to robust design (*see Robust Design*). Since the mid 1980s, robust design has received considerable attention both as a research topic and as a widely used methodology, with excellent results in industrial practice.

Essentially, robust design is a cost-effective method of reducing variation and is based on the use of statistically designed experiments to identify controllable factors that affect both the response level *via* so-called location effects and the response variability *via* dispersion effects. Obviously the idea is to use factors having a dispersion effect to reduce the response variation and factors having a location effect to accomplish an on-target process.

What Dispersion Effects Are and Why Factors Can Affect Variability

Let us call Y the quality characteristic of interest and $X_1, \dots, X_k$ the factors that affect Y , which we can control, and which are usually called *control or design factors*.

The response can then be expressed as

$$Y = f(X_1, \dots, X_k) + \varepsilon \quad (1)$$

Where ε is a random variable assuming $N(0, \sigma^2)$, and where σ^2 is caused by variation in a set of factors that escape our control: for example, environmental factors, variation in raw materials, wear from past use, or simply because they are unknown. These factors are called *noise factors*. Let us label them $N_1, \dots, N_l$. The situation is represented in Figure 1.

It is said that X_i has a dispersion effect when changes in its levels affect the size of σ . This may happen thanks to curvatures in the Y surface caused

by interactions (*see Interactions*) between control and noise factors.

Consider, as an illustration, the simple situation with only one control factor (X) and one noise factor (N) as represented in Figure 2. It is easy to see that the effect on Y of changes in N are much bigger when X is at level -1 than when it is at level 1 . In fact when X is at 0.5 , Y is independent of changes in N . It is clear that X can be used to reduce the variation in the response caused by N ; when this happens we say that factor X has a dispersion effect. Of course, in this case, it also has a location effect.

In summary, it is said that a control factor has a dispersion effect if it interacts with one or several noise factors.

Importance of Dispersion Effects and Examples

As pointed out in the introduction, variation reduction has always been an important goal for all types of processes. W. E. Deming (*see Quality Management, Overview*), one of the most prominent management gurus of the twentieth century, is supposed to have said, "If I would have to reduce my message to managers to a few words, they would be two: reduce variability."

MacKay [1] classifies variation reduction approaches into five generic strategies:

- introducing or tightening output inspection
- introducing or improving feedback control
- reducing variation in process inputs
- introducing or improving feed-forward control
- desensitizing the process to input variation.

Obviously the identification of dispersion effects falls into the last category. It should be noted that this is not only the most economical way to reduce variation, but it can also accomplish variation reduction against types of variability impossible to fight by the other four strategies, especially variation in the conditions under which the product will be used.

Box and Jones [2] inspired by the food industry, give a nice example of this last situation. Consider a producer of a cake mix to be sold in a box containing simple instructions. The quality of the cake when made by the purchaser will depend not only on the ingredients (producer's responsibility), but also on

2 Dispersion Effects

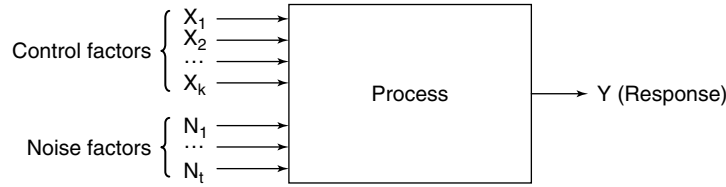


Figure 1 Control and noise factor scheme

how closely the baking instructions printed on the box – oven temperature and time of baking – are followed (outside producer’s control). Clearly, in this case the only possible strategy to reduce the variation in the final product quality characteristics is trying to find control factors with dispersion effects; that is, control factors that interact with time of baking or oven temperature.

To further illustrate the importance of dispersion effects and also how to detect them, we briefly present two well-documented and often studied industrial experiments. The first one includes noise factors in the experiment and the second one tries to detect dispersion effects from an experiment originally conducted to study location effects.

Leaf Spring Experiment

Pignatiello and Ramberg [3] reported an experiment in which the goal was to identify process conditions that would consistently produce leaf springs with a free height of 8 in. Four control factors were included in the experiment: furnace temperature (A), heating

time (B), transfer time (C), and hold down time (D). There was also a noise factor: the temperature of an oil quench (O) in which each spring was submersed at the end of the production process. This temperature was impossible to control during normal production process; it changed with the rate and temperature of the springs submersed, unless a costly investment were made to equip the quench with a heat regulation system.

The aim of the experiment was “. . . to find combinations of the levels of control factors that would yield a free height of 8 in. while being insensitive to the quench oil temperature”. The objective could be better expressed as to identify the location and dispersion effects of controllable factors, and use the knowledge gained to design a process that would consistently produce leaf springs with a free height of 8 in.

The experiment conducted was a 2^{4-1} for the control factors and a 2^1 for the noise factor, shown in Table 1. At each factor combination, three springs were produced in a row; notice that they are not real replicates.

Table 2 summarizes the finding regarding location effects.

The analysis presented by Pignatiello and Ramberg for variability is based on Taguchi’s signal-to-noise ratio (SN) (see **Signal-to-Noise Ratios for Robust Design**), which is a statistic that summarizes the results at each control factor setting over the various noise factor settings. As will be seen later, this type of analysis is called *row summary analysis*. Even though, in general, when SN ratios are used it is difficult or impossible to separate location and dispersion effects, in this case they identify factor B (heating time) as having a big dispersion effect and factors A , B , and D as having significant location effects. Similar results would be obtained by using any usual row summary statistics – the variance and several transformations have been proposed – for

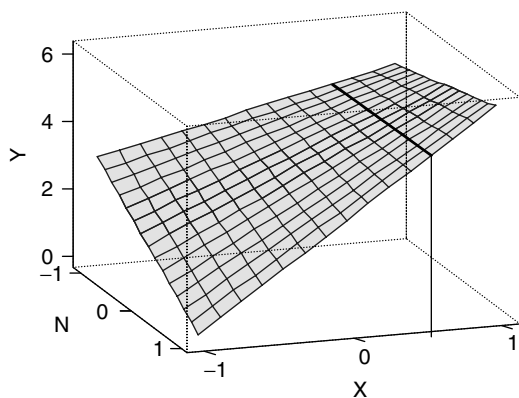


Figure 2 Interaction plot showing the dispersion effect of factor X

Table 1 Design matrix and response for the 2^{4-1} by 2^1 experiment

Control factors				Noise factor						Row summary	
A	B	C	D	-1			1			Mean	S ²
-1	-1	-1	-1	7.78	7.78	7.81	7.50	7.25	7.12	7.54	0.090
1	-1	-1	1	8.15	8.18	7.88	7.88	7.88	7.44	7.90	0.070
-1	1	-1	1	7.50	7.56	7.50	7.50	7.56	7.50	7.52	0.001
1	1	-1	-1	7.59	7.56	7.75	7.63	7.75	7.56	7.64	0.010
-1	-1	1	1	7.94	8.00	7.88	7.32	7.44	7.44	7.67	0.090
1	-1	1	-1	7.69	8.09	8.06	7.56	7.69	7.62	7.79	0.050
-1	1	1	-1	7.56	7.62	7.44	7.18	7.18	7.25	7.37	0.040
1	1	1	1	7.56	7.81	7.69	7.81	7.50	7.59	7.66	0.020

the variability. (Table 1 shows the mean and the variance.)

An alternative analysis was conducted by Steinberg and Bursztyn [4] by looking at the design as if it were a 2^{5-1} , and then studying the interactions between the four controllable factors and the noise factor. Notice that the way the design was constructed the control by noise interactions are free of confounding, although the design is of resolution IV. Table 3 summarizes the results.

This type of analysis has several advantages:

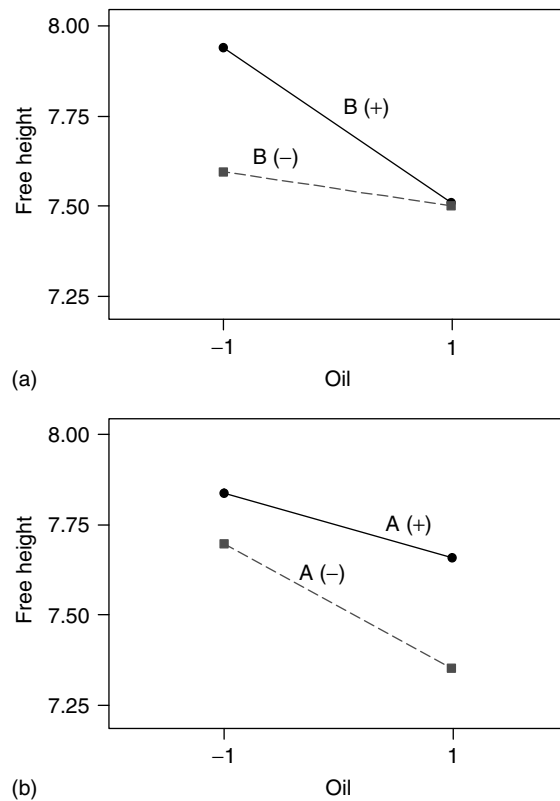
- Allows seeing that the noise factor (oil) has a large **main effect**, indicating that uncontrolled variation during production will cause substantial variation in the free height.
- Identifies B and A (not detected before) as having dispersion effects.
- The study of the control by noise interactions gives a lot of information about the process. In

Table 2 Factors with significant location effects

Factors	Estimated effect
Furnace temperature (A)	0.22
Heating time (B)	-0.18
Hold down time (D)	0.1

Table 3 Factors with significant location effects when analyzed as a 2^{5-1}

Factors	Estimated effect
Oil temperature (O)	-0.26
Furnace temperature (A)	0.22
Heating time (B)	-0.18
Hold down time (D)	0.1
Heating time $\times$ oil (BO)	0.166
Furnace temperature $\times$ oil (AO)	0.084

**Figure 3** Control by noise factor interactions

this case, it can easily be seen in Figure 3 that when working with B at a low level the effect of changing the oil temperature level is less than half the effect of doing so when B is at a high level. Factor B has a strong dispersion effect. Factor A also has a dispersion effect but clearly weaker.

A very simple summary of this case would be that it is possible to simultaneously reduce the variability

4 Dispersion Effects

thanks to the dispersion effects of B and to a lower degree of A, and adjust the mean of the leaf spring's height to 8 in, thanks to factor D, which only has a location effect.

The example illustrates how the detection of dispersion effects is an economical alternative to reduce variability caused by known noise factors.

Taguchi's Welding Experiment

Taguchi and Wu [5] report an arc welding experiment conducted at the National Railroad Corporation in Japan. The experiment was a 16-run, 9-factor, 2^{9-5} **fractional factorial** screening experiment. It has become a reference example of the identification of dispersion effects in unreplicated factorials and thus it has been analyzed many times; see for example: Box and Meyer [6] or Bisgaard and Pinho [7].

The aim of the experiment is to obtain maximum tensile strength (the response) with minimum dispersion. As has been said, the experiment has been analyzed many times by different methods, and all of them agree that factors B (drying method) and C (welding material) concentrate much of the action. Then each method claims further discoveries, but for our purposes let us concentrate on this undisputed conclusion.

Figure 4 shows clearly that both factors B and C affect the location and the level of tensile strength; and also that factor C may affect variability; notice that dispersion is bigger when using material SS41 than when material SB35 is used. Due to the fact that this experiment was conducted without noise factors it is impossible to know the noncontrollable factor that interacts with the material type.

In this case factor C has both location and dispersion effects and, unfortunately, it is not possible to use it at the same time to maximize tensile strength and reduce its variability. A compromise has to be found.

The example illustrates that even when the experiment is not replicated and noise factors are not included in the experiment, it is possible to detect dispersion effects. Of course in this case it is impossible to know which noise factors SB35 neutralizes.

Identification and Estimation of Dispersion Effects

In this section we give a summary overview of the different methods to identify and estimate dispersion effects *via* experimental designs. See reference [8] for a broad panel discussion. In general the designs used are two-level factorial designs (*see* **Factorial Experiments**); however, some authors, Taguchi among them, advocate the use of orthogonal arrays (*see* **Orthogonal Arrays**) that include two- and three-level factorial designs as well as designs with factors at two and three levels.

Designs Including Noise Factors

The basic idea is to include noise factors in the experiment, as in the leaf spring example. This implies that the experimenter has a clear idea of the noise factors to be neutralized and controlled during the experiment, even though they are expensive or impossible to control in normal production conditions.

There are two basic ways to include noise factors in the design, namely, product array (*see* **Product Array Designs**) and single array.

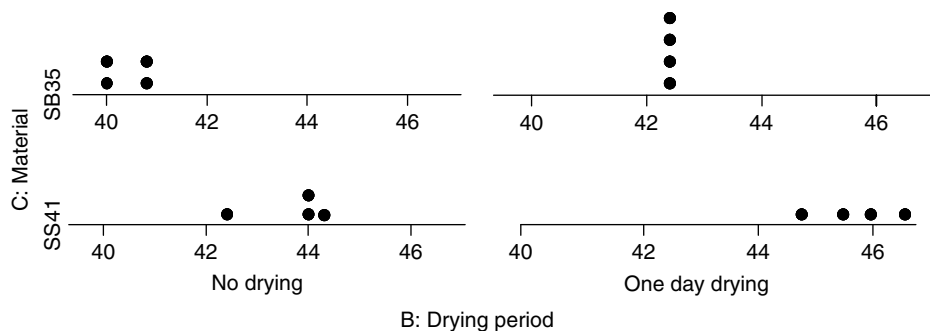


Figure 4 The 2^{9-5} welding experiment projected in the B and C factors space

Product Array. The idea is to set up a design for the control factors and another one for the noise factors – both can be highly fractionated (*see Fractional Factorial Designs*) – and then run all combinations of both designs. The leaf spring design (Table 1) was constructed in this way.

Care should be taken in randomizing (*see Randomization in Experimental Designs*) and running the experiment. It frequently happens that due to constraints on the execution of the trial or for economic reasons the experiment is run in a split-plot way, with all the noise factor levels run in succession at the same level of the control factors (*see Split-Plot Designs*), see Box and Jones [2] and Bisgaard [9].

If the design is run in a completely randomized way, then there are two basic methods of analyzing it. Both have been used and briefly described in the leaf spring example, namely, row summary statistics and response modeling.

Row summary statistics. The procedure begins by computing a statistic that summarizes the responses obtained at each control factor setting over the different levels of noise factors. Table 1 shows two row summary statistics, the mean and the variance. Once the so-called row summary statistic has been computed, it is analyzed in the usual way. Many row summary statistics have been proposed, although this method has been shown to lead under certain conditions to wrong or incomplete answers. This was the case in the leaf spring experiment. Regardless of this fact, the method is widely used, in part thanks to its conceptual simplicity – the method's main advantage – and in part because it is the one advocated by Taguchi and his followers. Another disadvantage is that it does not give information on which noise factor is neutralized by which control factor, thus hindering learning from the process.

Response modeling. The idea is to “integrate” control and noise factors in a single-design matrix and model the response as a function of both control and noise factors; this general approach is favored by many authors, see for example Box and Jones [2], Myers *et al.* [10] and Steinberg and Bursztyn [11]. When this is done, the study of the interactions between control and noise factors, especially through plots – as in the leaf spring example – gives information on the factors having dispersion effects as well as on how they neutralize the variability induced by

noise factors. The model obtained can also be used to derive a model for the variance; however, this is a dangerous approach because the model for the variance is very sensitive to the fitted response model.

Single Array. In this case a factorial experiment is set up, taking care that the design will allow for the **estimation**, without any dangerous confounding, of all two-factor interactions between control and noise factors. See Shoemaker *et al.* [12]. This approach is much more flexible in terms of choosing the effects to be estimated, free of unwanted confounding. Wu and Hamada [13] discuss the choice of single array and give tables for their selection. In general, designs produced by the single-array method cannot be analyzed using row summaries – the rows will not be complete – and thus, they have to be analyzed using the response modeling approach.

Replicated Designs Without Noise Factors

Including noise factors in an experiment is often complicated. As has been pointed out, it requires knowing what they are and also being able to control them during the experiment. It seems then that a possible approach is to conduct “conventional” replicates at each control factor combination indicated by the design matrix and let the noise factors act “naturally”.

The problem of this method is, as noted by Steinberg and Bursztyn [4] that – lacking the “exaggeration” of variability introduced by the deliberate variation of noise factors – they require a big number of replicates to be able to detect dispersion effects. And thus they are not, in general, an economically feasible alternative to designs, including noise factors.

Unreplicated Designs Without Noise Factors

Even though it is clear that the information about the variability of the response contained in an unreplicated factorial design is scarce, there have been many proposals to estimate dispersion effects from unreplicated factorial designs (*see Box–Meyer Method for Dispersion Effects; Bergman–Hynén Method for Dispersion Effects*); a general discussion can be found in Brenneman and Nair [14]. In general this approach should be viewed as a way to try to extract information on dispersion effects from an experiment conducted to estimate location effects rather than

a suitable method to use when the objective is to estimate dispersion effects.

All proposed methods are based on the effect-sparsity principle, which states that in general only a few of the factors included in the experiment will affect the response. The main problem with this approach is that, as has been noted earlier for replicated designs without noise factors, letting the noise factors act naturally gives a very poor **power** for detection of dispersion effects. Obviously not having replicates aggravates this problem. A second difficulty is, as pointed by Box and Meyer [6], the inescapable confusion between location and dispersion effects.

Summary

A factor is said to have a dispersion effect when changing it affects the variability of the response, and this happens when it interacts with factors called *noise factors* that cannot be controlled in normal operating conditions. There are several methods and all of them are based on the use of carefully designed experiments to identify and estimate factors having dispersion effects.

The identification and exploitation of control factors with dispersion effects is, in general, the most economical method and in many occasions – when the noise factors cannot be controlled during product use or normal production – the only one available.

References

- [1] MacKay, R. & Steiner, S. (1997). Strategies for variability reduction, *Quality Engineering* **10**, 125–136.
- [2] Box, G. & Jones, S. (1992). Split-plot designs for robust product experimentation, *Journal of Applied Statistics* **19**, 3–26.
- [3] Pignatiello, J. & Ramberg, J. (1985). Discussion of “off line quality control, parameter design and the Taguchi method” by Kackar R, *Journal of Quality Technology* **17**, 198–206.
- [4] Steinberg, D. & Bursztyn, D. (1998). Noise factors, dispersion effects and robust design, *Statistica Sinica* **8**, 67–86.
- [5] Taguchi, G. & Wu, Y. (1980). *Introduction to Off-Line Quality Control*, Central Quality Control Association, Nagoya.
- [6] Box, G. & Meyer, R. (1986). Dispersion effects from fractional factorial designs, *Technometrics* **28**, 19–27.
- [7] Bisgaard, S. & Pinho, A. (2003–04). Follow-up experiments to verify dispersion effects: Taguchi’s welding experiment, *Quality Engineering* **16**, 335–343.
- [8] Nair, V. (1992). Taguchi’s parameter design: a panel discussion, *Technometrics* **34**, 127–161.
- [9] Bisgaard, S. (2000). The design and analysis of 2k-px2q-r split plot experiments, *Journal of Quality Technology* **32**, 39–56.
- [10] Myers, R., Khuri, A. & Vining, G. (1992). Response surface alternatives to the Taguchi robust parameter design approach, *The American Statistician* **46**, 131–139.
- [11] Steinberg, D. & Bursztyn, D. (1994). Dispersion effects in robust-design experiments with noise factors, *Journal of Quality Technology* **26**, 12–20.
- [12] Shoemaker, A., Tsui, K. & Wu, J. (1991). Economical experimentation methods for robust design, *Technometrics* **33**, 247–257.
- [13] Wu, J. & Hamada, M. (2000). *Experiments. Planning, Analysis, and Parameter Design Optimization*, Wiley Interscience.
- [14] Brenneman, W. & Nair, V. (2001). Methods for identifying dispersion effects in unreplicated factorial experiments: a critical analysis and proposed strategies, *Technometrics* **43**, 388–405.

XAVIER TORT-MARTORELL AND LLUIS MARCO

Performance Measures for Robust Design

Introducing Taguchi's Signal-to-Noise Ratio and Two-Step Procedures

The objective of robust design (*see Robust Design*) is to find the vector of design variable settings that minimizes the quadratic loss function popularized by Taguchi and Wu [2],

$$L(Y, t) = A_0(Y - t)^2 \quad (1)$$

where A_0 is a constant, Y is the process response, and t the desired target value for Y . The average loss may be decomposed as follows:

$$\begin{aligned} E_{L(Y,t)} &= R(\mathbf{C}) \\ &= A_0 E_{\mathbf{N},\epsilon}(Y - t)^2 \\ &= A_0 [\text{Var}_{\mathbf{N},\epsilon}(Y) + (E_{\mathbf{N},\epsilon}(Y) - t)^2] \end{aligned} \quad (2)$$

where $\mathbf{C}$ is a vector of design variables that can be partitioned into two mutually exclusive parts corresponding to “nonadjustment” factors, that is, $\mathbf{d} = (d_1, \dots, d_p)$, and “adjustment” factors, that is, $\mathbf{a} = (a_1, \dots, a_q)$. Note that this partitioning of design variables is a theoretical ideal, and mutually exclusive groupings may not exist in practice. $\text{Var}_{\mathbf{N},\epsilon}(Y)$ and $E_{\mathbf{N},\epsilon}(Y)$ are the variance and mean of the process response over the random sources of variability represented by $\mathbf{N}$ and ϵ . $\mathbf{N}$ is a vector of uncontrollable “noise” variables (*see Factorial Experiments*) and ϵ represents purely random variation. For ease of notation we will simply write $\text{Var}(Y)$ and $E(Y)$ from here onward.

The adjustment factors making up $\mathbf{a}$ are those that affect the mean, that is, $E(Y)$, independently of $\text{Var}(Y)$, whereas the nonadjustment variables that make up $\mathbf{d}$ are assumed to affect only the variance $\text{Var}(Y)$. Equation (2) indicates that the minimization of average quadratic loss is mathematically equivalent to the minimization of the sum of process variance and the square of the deviation of process mean from the target.

Taguchi [3] introduced a family of performance measures called *signal-to-noise ratios* (SNRs; *see Signal-to-Noise Ratios for Robust Design*) whose specific form depends on the desired response

outcome. The case where the response has a fixed nonzero target is called the *nominal-the-best* (NTB) case and is defined as follows:

$$SNR = 10 \log_{10} \left[\frac{\bar{Y}^2}{S_Y^2} \right] \quad (3)$$

where $\bar{Y}$ and S_Y^2 are respectively the sample mean and variance of the response estimated at each test combination of design variables selected for experimentation. Responses having unique smaller-the-better or larger-the-better targets are respectively called the *STB* and *LTB* cases. We have only presented the formula of the SNR for the NTB case because it most transparently conveys the concept of a “signal”, that is, $\bar{Y}$, being compared to systemic background variability or “noise”, that is, S_Y^2 .

To accomplish the objective of minimal expected squared-error loss for the NTB case, Taguchi [3] proposed the following two-step optimization procedure:

1. Calculate and model the SNRs and find the values of the nonadjustment factor settings, that is, $\mathbf{d}$, which maximize the SNR.
2. Shift mean response to the target by changing the values of the adjustment factor(s) $\mathbf{a}$.

Although he provided a performance measure, that is, the SNR, and a related optimization scheme, that is, the two-step procedure, it was not explained how these two entities collectively achieve the goal of minimizing average quadratic loss.

Relating Performance Measure to Optimization: The Performance Measure Independent of Adjustment (PerMIA)

León *et al.* [1] showed that under certain circumstances Taguchi's use of the SNR, within the context of his two-step optimization procedure, does in fact minimize expected quadratic loss. Specifically this happens when the response conforms to a **transfer function** of the following form

$$Y = E_Y(\mathbf{a}, \mathbf{d})\epsilon(\mathbf{N}, \mathbf{d}) \quad (4)$$

where $\mathbf{N}$ are the noncontrollable noise variables and $(\mathbf{a}, \mathbf{d})$ are the adjustment and nonadjustment factors. They observed that for this class of models, the

2 Performance Measures for Robust Design

SNR of equation (3) can be minimized independently of the mean. This was a rigorous articulation of when the use of Taguchi's SNR within his two-step optimization procedure accomplished the stated goal of minimizing expected quadratic loss. They caution that when a response is not compliant with the functional form of equation (4), the use of the SNR in the two-step procedure of Taguchi does not necessarily minimize expected quadratic loss.

In order to provide a consistent framework within which to employ Taguchi's two-step procedure, they defined a new entity called the *performance measure independent of adjustment*, that is, the PerMIA. The PerMIA is a performance measure, that is, a function, which can be minimized independently of the mean and is specifically defined for the particular class of models being considered. Assuming that a PerMIA can be defined for the situation under consideration, use of a PerMIA within the framework of the two-step procedure guarantees the minimization of expected quadratic loss. This is because, in the first step, settings of the nonadjustment factors are chosen to minimize the PerMIA and, in the second step, the mean is shifted as closely as possible to its target value by choosing corresponding values for the adjustment factors. Because of the defining properties of the PerMIA, the adjustment of the mean in the second step does not affect the value of the PerMIA. When the process under consideration follows the model of equation (4), the SNR of equation (3) is a PerMIA but for other models it is not. This explained why use of the SNR in all situations does not always yield results that minimize the value of expected quadratic loss.

In particular León *et al.* [1] indicate that for an additive transfer function, that is,

$$Y = E_Y(\mathbf{a}, \mathbf{d}) + \epsilon(\mathbf{N}, \mathbf{d}) \quad (5)$$

that $\text{Var}(Y)$ is a PerMIA and should be used in place of the SNR. León *et al.* [1] conclude by proposing a modification of Taguchi's two-step procedure which first finds values of the nonadjustment factors $\mathbf{d}$ which minimize the PerMIA and, conditional on those values of $\mathbf{d}$, finds the values of the adjustment factors $\mathbf{a}$ that minimize expected quadratic loss.

Nair and Pregibon [4] advocate a three-phase process for robust design consisting of exploration, modeling, and optimization. Their performance measure is process response which they assume is a function of the design and noise variables, that is, $\mathbf{C}$ and $\mathbf{N}$

referred to in equation (2). They reduce the robust design problem to minimization of Var_Y subject to the constraint that E_Y is held as close as possible to the target level t .

They suggest that process variance can typically be modeled as a function of adjustment factors ($\mathbf{a}$) and nonadjustment factors ($\mathbf{d}$) as follows:

$$\text{Var}(Y)(\mathbf{C}) = f_1(E(Y)(\mathbf{a}, \mathbf{d}))\epsilon(\mathbf{N}, \mathbf{d}) \quad (6)$$

This model implies that minimization of variance relies primarily on the nonadjustment factors ($\mathbf{d}$) with residual variability determined by the distance between the mean and the target. They recommend a transformation-based approach which allows the separation of the **dispersion effects**, that is, $\epsilon(\mathbf{d})$, and the location effects, that is, $f_1(E(Y)(\mathbf{a}, \mathbf{d}))$. The specific transformation is identified in the exploratory phase with plots of the response data at each treatment combination supplemented with mean–variance plots. Their graphical examination plots the means and variances on a log–log scale which is approximately linear with slope k for the common situation where

$$f_1(E_Y(\mathbf{a}, \mathbf{d})) = [E_Y(\mathbf{a}, \mathbf{d})]^k \quad (7)$$

which represents an important specific case of equation (6). The values of k suggest the appropriate transformation with a value of $k = 0$ indicating a unity transformation due to the lack of a strong relationship between mean and variance.

Box [5] introduced a generalized transformation of response approach that applies to commonly found relations between location and dispersion. We remind the reader that León *et al.* [1] demonstrated that Taguchi's SNR is a PerMIA when Var_Y is proportional to E_Y . Box defines a PerMIA as

$$P(\mathbf{d}) = \frac{\text{Var}_Y(\mathbf{a}, \mathbf{d})}{f_2[(E_Y(\mathbf{a}, \mathbf{d}))]} \quad (8)$$

under the assumption that the function $f_2[(E_Y(\mathbf{a}, \mathbf{d}))]$ can be found for which the resulting PerMIA is a function of only the nonadjustment design variables $\mathbf{d}$. Because $f_2[(E_Y(\mathbf{a}, \mathbf{d}))]$ is not known *a priori*, Box provides a data analytic technique called the λ plot to empirically find it. The λ plot assumes a transformation of the data, that is, $Y^{(\lambda)}$, such that a value of λ that achieves maximum separation of location and dispersion effects of the process response can be graphically identified.

On a strategic level, Box [5] asserted that the two crucial aspects of robust design are the choice of an appropriate performance criterion and its **estimation**. He states that experimentation, use of simple graphical techniques and sequential testing are the best ways to pursue robust design rather than the use of pre-ordained performance criteria such as the SNR. Box also argued that the SNRs for the STB and LTB cases are inadequate summaries of data and extremely inefficient measures of location.

León and Wu [6] extend the work of León *et al.* [1]. The first thing they explain is how a two-step procedure converts a constrained robust design problem into an unconstrained one. They continue by articulating when a simple adjustment of mean to target can be substituted for solving for values of the adjustment variables **a** that minimize expected loss conditional on the values of the nonadjustment variables **d** that minimize the PerMIA. They introduce a corresponding procedure called the *constrained two-step procedure* (C2P) to simplify the identification of a PerMIA and solve the constrained minimization problem with a second step that simply adjusts the process mean.

Lastly, León and Wu [6] cite the fallacy of assuming a quadratic loss function and mention asymmetric losses around a target as one instance where the quadratic loss function is inappropriate. For non-quadratic loss functions they introduce general dispersion, location, and off-target measures for use in a two-step process. This generalization of the loss function and related two-step procedure includes identification of adjustment functions, that is, functions of design factors used to make adjustments of process location. Though Taguchi used the process mean as an adjustment function, they point to the median as a possible alternative. They apply these new techniques in a number of examples featuring additive and multiplicative models with nonquadratic loss functions.

In an extension of León and Wu's discussion of nonquadratic loss functions [6], Moorhead and Wu [7] develop modeling and analysis strategies for a general loss function where the quality characteristic follows a location-scale model. Their procedure adds a third step to the traditional two-step process, an adjustment step which moves the mean to the side of the target with lower cost. Although limited by its need for a location-scale model, this procedure builds upon the familiar two-step procedures used for quadratic functions. An alternative approach for

asymmetric loss that considers both univariate and multivariate cases is presented in Joseph [8].

Concluding Remarks

In defining the PerMIA, León *et al.* [1] provided a unifying framework in which a performance measure explicitly achieves the optimization goal of minimizing expected quadratic loss via the two-step procedures promoted by Taguchi. León and Wu [6] showed how the PerMIA represents a family of performance measures appropriate for nonquadratic and nonsymmetric loss functions as well. Box [5] and others, including Nair and Pregibon [4], Vining and Myers [9], Shoemaker *et al.* [10], Tsui [11–14], Tsui and Li [15], Lunani *et al.* [16], and Bérubé and Wu [17] posit that using process response or its transformation as a performance measure provides a more flexible approach to robust design than the SNR-based techniques of Taguchi.

References

- [1] León, R.V., Shoemaker, A.C. & Kackar, R.N. (1987). Performance measure independent of adjustment: an explanation and extension of Taguchi's signal to noise ratio, *Technometrics* **29**, 253–285.
- [2] Taguchi, G. & Wu, Y. (1980). *Introduction to Off-Line Quality Control*, Central Japan Quality Control Association, Nagoya.
- [3] Taguchi, G. (1986). *Introduction to Quality Engineering: Designing Quality into Products and Processes*, Kraus International Publications, White Plains.
- [4] Nair, V.N. & Pregibon, D. (1986). A data analysis strategy for quality engineering experiments, *AT&T Technical Journal* **65**, 73–84.
- [5] Box, G.E.P. (1988). Signal to noise ratios, performance criteria and transformations, *Technometrics* **30**, 1–31.
- [6] León, R.V. & Wu, C.F.J. (1992). Theory of performance measures in parameter design, *Statistica Sinica* **2**, 335–358.
- [7] Moorhead, P.R. & Wu, C.F.J. (1998). Cost-driven parameter design, *Technometrics* **40**, 111–119.
- [8] Joseph, V.R. (2004). Quality Loss functions for nonnegative variables and their applications, *Journal of Quality Technology* **36**, 129–138.
- [9] Vining, G.G. & Myers, R.H. (1990). Combining Taguchi and response surface philosophies: a dual response approach, *Journal of Quality Technology* **22**, 38–45.
- [10] Shoemaker, A.C., Tsui, K.-L. & Wu, C.F.J. (1991). Econometric experimentation methods for robust parameter design, *Technometrics* **33**, 415–428.

4 Performance Measures for Robust Design

- [11] Tsui, K.-L. (1992). An overview of Taguchi method and newly developed statistical methods for robust design, *IIE Transactions on Quality and Reliability* **24**(5), 44–57.
- [12] Tsui, K.-L. (1996). A multi-step analysis procedure for robust design, *Statistica Sinica* **6**, 631–648.
- [13] Tsui, K.-L. (1999a). Robust design optimization with multiple characteristics, *International Journal of Production Research* **37**, 433–445.
- [14] Tsui, K.-L. (1999b). Response model analysis of dynamic robust design experiments, *IIE Transactions on Quality and Reliability* **31**, 1113–1122.
- [15] Tsui, K.-L. & Li, A. (1994). Analysis of smaller and larger the better robust design experiments, *International Journal of Industrial Engineering* **1**, 193–202.
- [16] Lunani, M., Nair, V.N. & Wasserman, G.S. (1997). Graphical methods for robust design with dynamic characteristics, *Journal of Quality Technology* **29**, 327–338.
- [17] Bérubé, J. & Wu, C.F.J. (1998). Signal-to-noise ratio and related measures parameter design optimization, Technical Report 321, University of Michigan.

Related Articles

Robust Design; Signal-to-Noise Ratios for Robust Design.

TERRENCE E. MURPHY AND
KWOK-LEUNG TSUI

Product Array Designs

Introduction

Robust parameter design or **robust design** is the centerpiece of Taguchi's systematic approach to quality improvement (*see Robust Design*), and it consists of three major steps. The first step is to identify parameters (or factors) that affect the performance of a system. Taguchi proposed classification of the factors into three types: control factors (or design factors), noise factors, and signal factors. A system with at least one signal factor is referred to as a *dynamic system*. The primary goal of robust design is to find settings of control factors that simultaneously optimize a system's performance and make it robust to variations caused by noise and signal factors. The second step is to investigate the selected factors through experimentation and to identify optimal control settings to achieve the goal of robust design. Statistical design and analysis of experiments play a critical role in this step. The last step is to conduct confirmatory experiments to ensure that, when the optimal control settings are used, the quality of the system under investigation has indeed been improved.

This article gives a brief review of product arrays, which are among the most popular experimental plans in the second step of robust design. Product arrays were originally proposed by Taguchi and were given a different name (i.e. *inner-outer arrays*). Some researchers and engineers use the term *cross* or *crossed arrays*. The name *product array* is used throughout this article. Though Taguchi's original idea of using robust design for quality improvement is widely considered important, many of his proposed methods have been criticized and more efficient and effective alternatives have been developed by statisticians and other researchers in quality engineering. This is also the case with regard to selecting experimental plans for robust design experiments. In this article, both Taguchi's methods and other new alternatives are discussed.

The organization of the article is as follows. The section titled "Product Arrays" first reviews **factorial** designs and orthogonal arrays, and then defines and discusses product arrays. The section titled "Analysis of Parameter Design Experiments" briefly introduces several popular modeling approaches for analyzing data generated from robust design experiments and

discusses the advantage of product arrays from the modeling perspective. The section titled "Optimal Selection of Product Arrays" discusses the optimal selection of product arrays. The section titled "Compound Orthogonal Arrays" introduces compound orthogonal arrays, a generalization of product arrays. The section titled "Product Arrays with Signal Factors" briefly discusses product arrays involving signal factors.

Product Arrays

Factorial Design and Orthogonal Array

Suppose there are k factors with $l_1, l_2, \dots, l_k$ levels, respectively, in an experiment. These factors have in total $l_1 \cdot l_2 \cdots l_k$ different level combinations or settings. If all the settings are used, the experiment is said to have a full factorial design (or a $l_1 \cdots l_2 \cdots l_k$ design). When $l_1 = l_2 \cdots = l_k = l$, i.e. all the factors have the same number of levels, the full factorial design is then said to be an l^k design. For example, when $l = 2$, we have a full factorial 2^k design. After data are collected from an experiment with a full factorial design, the analysis of variance method is often used to assess the factorial effects of the factors, which include main effects, two-factor interactions, and so on (*see Analysis of Variance*). If an effect involves r factors, it is said to be of order r where $1 \leq r \leq k$; a **main effect** only involves one factor, so it is of order 1.

The major advantage of using a full factorial design is that it guarantees the independent **estimation** of all the factorial effects. When the number of factors increases, the full factorial design quickly becomes infeasible. In practice, a subset of the factor settings is usually selected and used in the experiment. The selected subset is referred to as a **fractional factorial design**. In a fractional factorial design, factorial effects become entangled with each other and are not independently estimable. Nevertheless, a fractional factorial design is sufficient for the study of factorial effects of lower order such as main effects and two-factor interactions, when it is judiciously selected and effects of higher order are assumed to be negligible. The orthogonal arrays proposed by Rao [1] are an example of such designs. An **orthogonal array** with N rows, k columns (or factors), and strength t is an N by k array whose rows are the factor settings such that for any t columns, all

2 Product Array Designs

the possible level combinations of the corresponding t factors appear an equal number of times. The orthogonal array is denoted by $OA(N, l_1 l_2 \cdots l_k, t)$. If some of the factors have the same number of levels, they are usually grouped together in the notation. For example, if among the k factors, k_1 factors have s_1 levels, k_2 factors have s_2 levels, and so on, then the orthogonal array is denoted as $OA(N, s_1^{k_1} s_2^{k_2} \cdots, t)$. An $OA(N, s_1^{k_1} s_2^{k_2} \cdots, t)$ is sometimes denoted by $L_N(s_1^{k_1} s_2^{k_2} \cdots)$ when its strength is equal to two. See **Orthogonal Arrays** for more details.

Orthogonal arrays of strength two, three, or four are commonly used in practice. Under reasonable assumptions about higher-order factorial effects, these orthogonal arrays are sufficient for investigating main effects and/or two-factor interactions. For example, if interactions are assumed to be negligible, then an orthogonal array of strength two guarantees the independent estimation of main effects; if interactions between three or more factors are assumed to be negligible, an orthogonal array of strength four guarantees that both main effects and two-factor interactions are estimable. The estimability of an orthogonal array of strength three is between those of strengths two and four. For general results regarding the estimability of orthogonal arrays of strength t , the reader can consult Hedayat *et al.* [2] and Dey and Mukerjee [3].

Orthogonal arrays can be classified into symmetrical arrays and asymmetrical arrays, according to whether the factors have the same number of levels or not. They can also be classified into regular arrays and nonregular arrays. In general, regular arrays can be constructed by defining relations between factorial effects. Regular symmetrical arrays are among the most popular experimental designs in practice. Two examples are 2^{k-p} and 3^{k-p} designs, where k is the number of factors, 2 or 3 is the number of levels and p is the fraction index. The resolution of a regular symmetrical array is defined to be the lowest order of the effects that are entangled or aliased with the grand mean. It can be shown that the resolution is equal to the strength of the array plus 1. Hence, a 2^{k-p} or 3^{k-p} design with resolution R is an $OA(2^q, 2^k, R-1)$ or $OA(3^q, 3^k, R-1)$, respectively, where $q=k-p$. The minimum aberration criterion proposed by Fries and Hunter [4] is generally used to select optimal regular symmetrical arrays (see **Minimum Aberration**). For a comprehensive account of 2^{k-p} and 3^{k-p} designs and extensive tables of optimal designs,

the reader is referred to Wu and Hamada [5]. General orthogonal arrays have also been tabulated; for example, Hedayat *et al.* [2] includes a fairly comprehensive table (i.e. Table 12.7) of orthogonal arrays of strength 2 with up to 100 runs, and Sloane maintains a library of over 200 orthogonal arrays online (<http://www.research.att.com/~njas/oadir/index.html>).

Product Arrays

As discussed in the introduction, there are three possible types of factors present in a robust design experiment. For ease of discussion, we focus on experiments involving only control and noise factors in the sections titled “Product Arrays”, “Analysis of Parameter Design Experiments”, “Optimal Selection of Product Arrays”, and “Compound Orthogonal Arrays”, and briefly discuss the dynamic cases (i.e. those involving signal factors) in the section titled “Product Arrays with Signal Factors”. Product arrays were originally proposed by Taguchi to accommodate control factors and noise factors in robust design experiments.

Suppose there are k control factors and m noise factors in a robust design experiment. Let $OA_c(N_c, \dots, t_c)$ be an orthogonal array for the k control factors. Because the number of levels of each factor can be arbitrary, it is left unspecified in the notation. Denote the i th row (i.e., the i th control setting) of $OA_c(N_c, \dots, t_c)$ as c_i for $1 \leq i \leq N_c$. Then $OA_c(N_c, \dots, t_c)$ can be considered a collection of control settings, that is, $OA_c(N_c, \dots, t_c) = \{c_i : 1 \leq i \leq N_c\}$. Similarly, let $OA_n(N_n, \dots, t_n)$ be an orthogonal array for the m noise factors. Denote the j th row of $OA_n(N_n, \dots, t_n)$ as n_j . Then $OA_n(N_n, \dots, t_n) = \{n_j : 1 \leq j \leq N_n\}$. The two arrays $OA_c(N_c, \dots, t_c)$ and $OA_n(N_n, \dots, t_n)$ are referred to as the *control array* and the *noise array*, respectively. Taguchi proposed the use of the combined settings (c_i, n_j) with $1 \leq i \leq N_c$ and $1 \leq j \leq N_n$ for the robust design experiment. Clearly, these combined settings are the Cartesian products of the control settings of OA_c and the noise settings of OA_n . Therefore, the resulting design is named a *product array*. The product array is denoted by $OA_p(N_c \times N_n, \dots, (t_c, t_n))$, and it is equal to

$$OA_c(N_c, \dots, t_c) \times OA_n(N_n, \dots, t_n)$$

$$= \{(c_i, n_j) : 1 \leq i \leq N_c, 1 \leq j \leq N_n\} \quad (1)$$

In Taguchi's original terminology, the control array, noise array, and product array are referred to as the *inner array*, *outer array*, and *inner-outer array*, respectively. The product array OA_p is indeed an orthogonal array with $N_c N_n$ rows and $k + m$ columns.

We present a real robust design experiment as an illustrative example. A robust design experiment for improving an injection-molding process was reported in Engel [6]. There were seven control factors, namely, cycle time (A), mold temperature (B), cavity thickness (C), holding pressure (D), injection speed (E), holding time (F), and gate size (G), and three noise factors, namely, percentage regrind (a), moisture content (b), and ambient temperature (c). Each factor had two levels labeled -1 and $+1$, respectively. The goal of the experiment was to determine process parameter settings at which percent shrinkage would be consistently close to a target value. A regular 2^{7-4} design with resolution III (i.e. $OA_c(2^3, 2^7, 2)$) was selected as the control array and a regular 2^{3-1} design with resolution III (i.e. $OA_n(2^2, 2^3, 2)$) was selected as the noise array. The product array is $OA_p(2^3 \times 2^2, 2^{10}, (2, 2)) = OA_c(2^3, 2^7, 2) \times OA_n(2^2, 2^3, 2)$, whose layout together with the data generated from the experiment is given in Table 1.

The c_i 's and n_j 's in the above table indicate the control and the noise settings, respectively. The experiment has 32 runs in total and one observation at each combined setting of the control and noise factors.

A robust design experiment using a product array can be conducted in two different ways. First, it can

be conducted as a completely randomized experiment in which the run order needs to be completely randomized and the factors need to be reset after each run. Second, it can be conducted as a **split-plot** experiment in which the run order is subject to some restrictions. For example, when readjusting the control factors means rebuilding a different prototype system, which can be both time consuming and costly, it is practically more appealing to build the prototype system according to a fixed control setting, and then test the system across all the noise settings. Therefore, the control setting will not be reset between these runs. An experiment run in this manner is known as a *split-plot experiment*, and the analysis needs to take account of this structure; for details, see **Split-Plot Designs**, Box and Jones [7], and Bisgaard and Steinberg [8]. Bingham and Sitter [9] studied the optimal selection of split-plot designs for robust design experiments. In the remainder of this article, we focus on completely randomized robust design experiments only.

Compounding Noise Factors

A major criticism on using product arrays for robust design experiments is that their size tends to be large. For a product array $OA_p(N_c \times N_n, \dots, (t_c, t_n))$, the total number of runs is $N_c N_n$. When the number of control and noise factors increases, $N_c N_n$ may become too large for the product array to be practical. Therefore, there is a need to limit the size of a product array. According to Taguchi, robust design needs to investigate as many control factors as possible, so there might not be much room to reduce the size of a control array. Taguchi saw the opportunity

Table 1 Product Array and Shrinkage Data, Injection Molding Experiment

Runs		A	B	C	D	E	F	G	n_1	n_2	n_3	n_4
									a	b	c	
									-1	-1	$+1$	$+1$
									-1	$+1$	-1	$+1$
									-1	$+1$	$+1$	-1
1-4	c_1	-1	-1	-1	-1	-1	-1	-1	2.2	2.1	2.3	2.3
5-8	c_2	-1	-1	-1	$+1$	$+1$	$+1$	$+1$	0.3	2.5	2.7	0.3
9-12	c_3	-1	$+1$	$+1$	-1	-1	$+1$	$+1$	0.5	3.1	0.4	2.8
13-16	c_4	-1	$+1$	$+1$	$+1$	$+1$	-1	-1	2.0	1.9	1.8	2.0
17-20	c_5	$+1$	-1	$+1$	-1	$+1$	-1	$+1$	3.0	3.1	3.0	3.0
21-24	c_6	$+1$	-1	$+1$	$+1$	-1	$+1$	-1	2.1	4.2	1.0	3.1
25-28	c_7	$+1$	$+1$	-1	-1	$+1$	$+1$	-1	4.0	1.9	4.6	2.2
29-32	c_8	$+1$	$+1$	-1	$+1$	-1	-1	$+1$	2.0	1.9	1.9	1.8

of size reduction in the noise array. He proposed to compound noise factors as one hyperfactor and select a small number of representative levels of this hypernoise factor as the noise settings. The success of using compounded noise factors in robust design requires the judicious selection of the hypernoise factor settings as well as domain knowledge and understanding of the noise effects. There has not been much study about this technique in the literature, though it has been occasionally used in practice. One systematic investigation of the advantages and disadvantages of this technique can be found in Hou [10].

Analysis of Parameter Design Experiments

After data are collected from a robust design experiment, statistical methods are used to analyze the data and find the control settings for robust optimization. These methods primarily follow three different modeling approaches, which are the signal-to-noise (SN) ratio modeling approach proposed by Taguchi (*see Signal-to-Noise Ratios for Robust Design*), the location-dispersion modeling approach, and the response modeling approach. Statisticians and researchers in quality engineering have shown that while the SN ratio modeling approach is effective under some scenarios, it can be seriously flawed in other cases. Many close followers of Taguchi's quality philosophy and practice however still solely rely on this approach. Product arrays might have been proposed particularly for the SN ratio modeling approach. Nonetheless, they are also well suited for the other two approaches. This section briefly reviews these three approaches and then discusses the major advantage of product arrays from an analysis perspective.

SN Ratio Modeling

The SN ratio summarizes performance at each possible control setting. The particular SN ratio used depends on the nature of the problem; *see* Phadke [11] or **Signal-to-Noise Ratios for Robust Design** for details. A common goal of robust design is to achieve a particular nominal level and in this setting the SN ratio is defined to be $\eta = \ln(\mu^2/\sigma^2)$, where μ and σ^2 are the mean and variance of the response of the system under study. Suppose an $OA_p(N_c \times N_n, \dots, (t_c, t_n))$ is used in a robust

design experiment and the obtained data are $\{y_{ij} : 1 \leq i \leq N_c, 1 \leq j \leq N_n\}$, where y_{ij} is the observed response at (c_i, n_j) for $1 \leq i \leq N_c$ and $1 \leq j \leq N_n$. From the data, the SN ratio at the control setting c_i ($1 \leq i \leq N_c$) can be estimated by $\hat{\eta}_i = \ln(\bar{y}_i^2/s_i^2)$, where $\bar{y}_i$ and s_i^2 are the sample mean and variance of $\{y_{ij}\}_{1 \leq j \leq N_n}$. It is clear that the use of the noise array in the experiment is to systematically introduce and evaluate variations attributable to the noise factors at different control settings. The SN ratio $\hat{\eta}_i$ reflects the magnitude of the mean response (i.e. the signal) relative to the variation at the control setting c_i (i.e. the noise). Taguchi proposed to select control settings that maximize the SN ratio to make the system robust to the noise factors.

Once the SN ratios $\{\hat{\eta}_i\}_{1 \leq i \leq N_c}$ are available, regression analysis is used to identify the so-called dispersion effects, which are the control effects that significantly affect the SN ratio (*see Dispersion Effects*). In a separate regression analysis using $\{\bar{y}_i, 1 \leq i \leq N_c\}$ as the responses, the control effects that contribute significantly to the response mean can also be identified. Effects that affect the mean response but do not affect the SN ratio are called the *adjustment effects*. For the nominal-the-best problem, Taguchi suggested the following two-step procedure for robust optimization: (a) Select the settings of the control factors involved in the dispersion effects to maximize the SN ratio $\hat{\eta}$; (b) Select the settings of the factors involved in the adjustment effects to bring the mean response on target. The use of this procedure can be justified by the performance measures independent of adjustment (**PerMIA**), which were proposed by Léon *et al.* [12] and further studied by Léon and Wu [13]; *see Performance Measures for Robust Design* for details. The procedure may however fail to identify dispersion or adjustment effects under other scenarios; *see* Bérubé and Wu [14] for a simulation study. Statisticians and probably the majority of researchers in quality engineering seem to agree that the unquestioning use of the SN ratio for robust design is not desirable; more discussions can be found in Nair [15], Steinberg [16], and Wu and Hamada [5].

Location-Dispersion Modeling

Because the success of robust design hinges on the understanding of how control factors affect the mean μ and variance σ^2 of a system's performance, it is

desirable to directly model μ and σ^2 as functions of the control effects. This idea leads to the location-dispersion modeling approach, which fits regression models for $\{\bar{y}_i\}$ and $\{\log s_i^2\}$ separately and uses the resulting models for robust optimization. Modeling of response mean and variance separately, however, suffers from an obvious drawback that the **heteroscedasticity** inherent in the problem has not been taken into account. A more efficient approach is to fit the mean and variance jointly. Such approaches have already been studied in the statistics literature; see Davidian and Carroll [17, 18] and Nelder and Pregibon [19]. Engel [6] proposed a model to jointly fit response mean and variance in robust design, which is $E(y) = \mu = x^\tau \beta$, $\text{Var}(y) = \phi V(\mu, \theta)$, and $\log(\phi) = x^\tau \gamma$, where β and γ are parameter vectors and θ is usually a scalar, and x is the vector of control effects. Engel then proposed the use of either pseudonormal likelihood or extended quasi-likelihood for parameter estimation. The fitted model is then used for robust optimization; see Engel [6] for more details.

Response Modeling

In a robust design experiment, both the control and noise factors are varied systematically and the responses at different combined settings are observed. Therefore, it is possible to directly model the response as a function of all the factorial effects including the control effects, the noise effects, and their interactions. Subsequently, the fitted response model can be used to derive models for the response mean and variance, which can be further used for robust optimization. This approach is referred to as the *response modeling approach*. The response modeling approach was hinted as early as 1985 by Eastering [20], and was investigated in detail by Welch *et al.* [21], Shoemaker *et al.* [22], and Vining and Myers [23].

Let x and z denote the control and noise factors, respectively, and let $f(x)$, $g(z)$, $h(x)$, and $k(z)$ denote vectors of control effects and noise effects accordingly. A general model that relates a response y to the control and noise effects is $y = \alpha_0 + f(x)^\tau \alpha + g(z)^\tau \beta + h(x)^\tau \Gamma k(z) + \epsilon$, where α and β are coefficient vectors, Γ is a coefficient matrix of the interactions between the control and noise factors, and ϵ is a random error. This general model induces the models for the response mean and variance, which are, respectively, $\mu = E(y) =$

$\alpha_0 + f(x)^\tau \alpha$, and $\sigma^2 = \text{Var}(y) = \beta^\tau \text{Var}(g(z))\beta + h(x)^\tau \Gamma \text{Var}(k(z))\Gamma h(x)$. Given the data from a robust design experiment, the coefficients α , β , and Γ can be estimated, using the general response model. Plugging these estimates in μ and σ^2 gives the fitted mean and variance models, which can be further used for robust optimization. The advantage of the response modeling approach is multifold. First, it can reveal how control and noise factors interact with each other to affect a system's performance. Second, more efficient experimental strategies such as combined arrays can be used for robust design. Third, **response surface** methodologies such as sequential design and **central composite** design can be incorporated into robust design. See **Robust Design** for more details on this approach.

Modeling Advantage of Product Arrays

As mentioned earlier, the product structure of a product array makes it possible to calculate the response mean and variance at any control setting. Therefore, it is well suited for the SN ratio modeling approach. As a matter of fact, the product structure has its merits for the other two approaches as well.

For the location-dispersion modeling approach, Rosenbaum [24, 25] showed that a product array $OA_p(N_c \times N_n, \dots, (t_c, t_n))$ with $t_c = t_n = 2$ ensures that variance calculated at each control setting does not confound control-by-control interactions with any control-by-noise interactions. Hence, when modeling variance as a function of control effects, the dispersion effects can be identified. Furthermore, Rosenbaum extended product arrays to compound arrays that have more **power** in finding location and dispersion effects. Compound orthogonal arrays will be discussed in the section titled "Compound Orthogonal Arrays". For the response modeling approach, the advantage of a product array is stated in the following theorem, which is a modified version of Theorem 10.1 in Wu and Hamada [5]. For ease of presentation, both the control and noise arrays of the product array are two-level regular designs.

Theorem Suppose a 2^{k-p} design is chosen for the control array OA_c , a 2^{m-q} design is chosen for the noise array OA_n , and the product array is $OA_p = OA_c \times OA_n = 2^{k-p} \times 2^{m-q}$.

1. If $\alpha_1, \alpha_2, \dots, \alpha_u$ are u estimable control effects in OA_c and $\beta_1, \beta_2, \dots, \beta_v$ are v estimable noise

effects in OA_n , then the factorial effects $\alpha_i, \beta_j, \alpha_i\beta_j$ for $1 \leq i \leq u, 1 \leq j \leq v$ are estimable in OA_p .

2. All the km control-by-noise two-factor interactions are not confounded with other main effects or two-factor interactions.

Part 2 of the above theorem implies that control-by-noise two-factor interactions can be estimated under the assumption that interactions of order higher than two are negligible. This is considered the major advantage of a product array for the response modeling approach.

Optimal Selection of Product Arrays

Owing to the product structure, the estimation properties of a product array are completely determined by those of its control and noise arrays. Control and noise arrays with good estimation properties lead to good product arrays. Therefore, given the size and the numbers of control and noise factors, an optimal product array is the product of a control array and a noise array, which are optimal in their own rights. For regular orthogonal arrays, the minimum aberration criterion is used to define and select optimal designs; and for nonregular orthogonal arrays, minimum aberration type of criteria have been developed lately (for example, see Xu [26, 27]). In practice, many practitioners are not aware of these better designs and still use standard choices such as those advocated by Taguchi. Extensive tables of minimum aberration two-level, three-level, and mixed-level designs with economical size can be found in Wu and Hamada [5], and some optimal nonregular orthogonal arrays can be found in Xu [26, 27].

Compound Orthogonal Arrays

By relaxing the product structure, Rosenbaum [24, 25] proposed compound orthogonal arrays as a generalization of product arrays. Compared with product arrays, compound orthogonal arrays can have higher strengths for given size and better run-size economy for given strengths. A product array $OA_p(N_c \times N_n, \dots, (t_c, t_n))$ is a collection of (c_i, n_j) where c_i and n_j are the i th row and j th row of the control array $OA_c(N_c, \dots, t_c)$ and the noise array $OA_n(N_n, \dots, t_n)$, respectively. Let $c_i \times$

$OA_n = \{(c_i, n_j), 1 \leq j \leq N_n\}$. Then $OA_p(N_c \times N_n, \dots, (t_c, t_n)) = \bigcup_{1 \leq i \leq N_c} c_i \times OA_n$. Clearly when forming the product array, the same noise array is used for each control setting. Rosenbaum suggested that it should be advantageous to relax this restriction and allow the use of different noise arrays at different control settings. For example, at two different control settings c_i and c_j , two different noise arrays denoted by OA_{n,c_i} and OA_{n,c_j} can be considered, where the subscripts c_i and c_j indicate that these arrays depend on the corresponding control settings. The collection of $c_i \times OA_{n,c_i}$ ($1 \leq i \leq N_c$) forms an orthogonal array, which Rosenbaum called a *compound orthogonal array*. Owing to the allowance of using different noise arrays, compound orthogonal arrays provide more flexibility in accommodating control factors and noise factors, which leads to the advantage stated in the beginning of this section.

Rosenbaum proposed various ways to construct compound orthogonal arrays. Hedayat and Stufken [28] and Zhu *et al.* [29] have further studied compound arrays that have optimal estimation properties. Some useful two-level compound orthogonal arrays have been tabulated in these two references. Compound orthogonal arrays have not yet been adopted by practitioners in real robust design experiments.

Product Arrays with Signal Factors

As mentioned in the introduction, a system with at least one signal factor is referred to as a *dynamic system*. Taguchi [30] referred to the robust design of such systems by the term *dynamic problems*. The focus of a dynamic problem is to optimize the relationship between a system performance measure and the signal factor and simultaneously make the relationship robust to the noise factors. Taguchi defined the dynamic SN ratio and proposed its use for robust optimization of dynamic systems. This can be considered a generalization of the SN ratio modeling approach discussed in the section titled “SN Ratio Modeling”. The response modeling approach has also been developed for dynamic problems by Miller and Wu [31]. Other analysis methods for dynamic problems can be found in Grove and Davis [32], Lunani *et al.* [33], and McCaskey and Tsui [34]. See **Signal–Response Systems** for more details on dynamic systems.

The product arrays discussed in the previous sections do not include signal factors. There are three

possible methods to incorporate signal factors in a product array. For ease of discussion, we assume there is only one signal factor, denoted by G . The first method is to cross a product array OA_p of control and noise factors with G , which leads to an orthogonal array $OA_p \times G$. In other words, for each combined setting of the control and noise factors, the system is tested across the experimental levels of G . The second method is to add the signal factor to the noise array, which leads to a noise–signal array $OA_{n,g}$; and then it is crossed with a control array to generate a product array $OA_c \times OA_{n,g}$. This way was suggested by Taguchi and was often used in practice. The third method is to add the signal factor to the control array, which results in a control–signal array denoted by $OA_{c,g}$; and then it is crossed with the noise array to generate a product array $OA_{c,g} \times OA_n$. There does not exist much study on optimal selection of product arrays involving signal factors, especially for the last two types. See Wu and Hamada [5] for more discussions.

References

- [1] Rao, C.R. (1947). Factorial experiments derivable from combinatorial arrangements of arrays, *Journal of the Royal Statistical Society* **9**, 128–139.
- [2] Hedayat, A.S., Sloane, N.J.A. & Stufken, J. (1999). *Orthogonal Arrays: Theory and Applications*, Springer, New York.
- [3] Dey, A. & Mukerjee, R. (1999). *Fractional Factorial Plans*, John Wiley & Sons, New York.
- [4] Fries, A. & Hunter, W.G. (1980). Minimum aberration 2^{k-p} designs, *Technometrics* **22**, 601–608.
- [5] Wu, C.F.J. & Hamada, H. (2000). *Experiments: Planning, Analysis, and Parameter Design Optimization*, John Wiley & Sons, New York.
- [6] Engel, J. (1992). Modeling variation in industrial experiments, *Applied Statistics* **41**, 579–593.
- [7] Box, G.E.P. & Jones, S. (1992). Split-plot designs for robust product experimentation, *Journal of Applied Statistics* **19**, 3–26.
- [8] Bisgaard, S. & Steinberg, D. (1997). The design and analysis of $2^{k-p} \times S$ prototype experiments, *Technometrics* **39**, 52–62.
- [9] Bingham, D. & Sitter, R.R. (2003). Fractional factorial split-plot designs for robust parameter experiments, *Technometrics* **45**, 80–89.
- [10] Hou, X.S. (2002). On the use of compound noise factor in parameter design experiments, *Applied Stochastic Models in Business and Industry* **18**, 225–243.
- [11] Phadke, M.S. (1989). *Quality Engineering Using Robust Design*, Prentice Hall, Eaglewood Cliffs.
- [12] Léon, R.V., Shoemaker, A.C. & Kacker, R.N. (1987). Performance measurements independent of adjustment (with discussion), *Technometrics* **29**, 253–285.
- [13] Léon, R.V. & Wu, C.F.J. (1992). A theory of performance measures in parameter design, *Statistica Sinica* **2**, 335–358.
- [14] Bérubé, J. & Wu, C.F.J. (2000). Signal-to-noise ratio and related measures in parameter design optimization: an overview, *Sankhya, Series B* **62**, 417–432.
- [15] Nair, V.N. (1992). Taguchi's parameter design: a panel discussion, *Technometrics* **34**, 127–161.
- [16] Steinberg, D.M. (1996). Robust design: experiments for improving quality, *Handbook of Statistics*, S., Ghosh & C.R. Rao, eds, Elsevier Science, Vol. 13, pp. 199–240.
- [17] Davidian, M. & Carroll, R.J. (1987). Variance function estimation, *Journal of the American Statistical Association* **82**, 1079–1092.
- [18] Davidian, M. & Carroll, R.J. (1988). A note on extended quasilielihood estimation, *Journal of the Royal Statistical Society. Series B* **50**, 74–82.
- [19] Nelder, J.A. & Pregibon, D. (1987). An extended quasilielihood function, *Biometrika* **74**, 221–232.
- [20] Easterling, R.G. (1985). Discussion of the Paper by Kacker, *Journal of Quality Technology* **17**, 191–192.
- [21] Welch, W.J., Yu, T.K., Kang, S.M. & Sacks, J. (1990). Computer experiments for quality control by parameter design, *Journal of Quality Technology* **22**, 15–22.
- [22] Shoemaker, A.C., Tsui, K.L. & Wu, C.F.J. (1991). Economical experimentation methods for robust design, *Technometrics* **33**, 415–427.
- [23] Vining, G.G. & Myers, R.H. (1990). Combining Taguchi and response surface philosophies: a dual response approach, *Journal of Quality Technology* **22**, 38–45.
- [24] Rosenbaum, P.R. (1994). Dispersion effects from fractional factorials in Taguchi's method of quality design, *Journal of the Royal Statistical Society. Series B* **56**, 641–652.
- [25] Rosenbaum, P.R. (1996). Some useful compound dispersion experiments in quality design, *Technometrics* **38**, 354–364.
- [26] Xu, H. (2002). An algorithm for constructing orthogonal and nearly-orthogonal arrays with mixed levels and small runs, *Technometrics* **44**, 356–368.
- [27] Xu, H. (2003). Minimum moment aberration for nonregular designs and supersaturated designs, *Statistica Sinica* **13**, 691–708.
- [28] Hedayat, A.S. & Stufken, J. (1999). Compound orthogonal arrays, *Technometrics* **41**, 57–61.
- [29] Zhu, Y., Zeng, P. & Jennings, K. (2007). Optimal compound orthogonal arrays and single arrays for robust parameter design experiments, *Technometrics*, to appear.
- [30] Taguchi, G. (1987). *System of Experimental Design*, UNIPUB, White Plains, Vols 1 and 2.
- [31] Miller, A. & Wu, C.F.J. (1996). Parameter design for signal-response systems: a different look at Taguchi's dynamic parameter design, *Statistical Science* **11**, 122–136.

8 Product Array Designs

- [32] Grove, D.M. & Davis, T.P. (1992). *Engineering Quality and Experimental Design*, Longman Scientific and Technical, Essex.
- [33] Lunani, M., Nair, V. & Wasserman, G. (1997). Graphical methods for robust design with dynamic characteristics *Journal of Quality Technology* **29**, 327–338.
- [34] McCaskey, S.D. & Tsui, K.-L. (1997). Analysis of dynamic robust design experiments, *International Journal of Production Research* **35**, 1561–1574.

Further Reading

Taguchi, G. (1991). *Taguchi Methods, Vol. 1: Research and Development; Vol. 3 Signal-to-Noise Ratio for Quality Evaluation*, ASI Press, Dearborn.

YU ZHU

Robust Design

Introduction

Robust design is a methodology by which the controllable factors (inputs) of a system are set in order to achieve consistent and desirable response (output) from the system. The defining characteristic of robust design is the emphasis on achieving a consistently desirable response, not just a response that is desirable on average. “Robust” in this context means “insensitive” or “unaffected” and thus the goal is to find settings of the controllable factors of the system that produce results that are not only desirable but also unaffected by environmental variation or by simple repeated use of the system. Here a brief introduction and tutorial of robust design will be given. More extensive reviews of robust design methodology and literature can be found in Ankenman and Dean [1], Steinberg [30], Wu and Hamada [35, Chapter 10], and Box, Hunter, and Hunter [9, Chapter 13].

In the statistical literature, robust design is a natural successor to response surface methodology (RSM) (see **Response Surface Methodology** and [24]). Box and Wilson [10] wrote the seminal paper for RSM and its title was “On Experimental Attainment of Optimum Conditions”. In that paper, as with most of the literature on RSM, the focus was on optimizing the expected value of the response. Specifically, they strove to maximize the expected yield of a chemical process. Conceptually, robust design can also be thought of as experimental attainment of optimum conditions, where the optimization now includes not only achieving a desirable expected value of the response but also achieving low variance of the response over time and across all environmental conditions that are likely to occur.

The high-level procedure of robust design follows the same basic set of steps as RSM, it is only the specific design and analysis methods that must change due to the emphasis on robustness. The high-level procedure, applicable for both RSM and robust design, involves four basic steps. Step 1 is to study the system to determine the factors (inputs) and response (output) of interest. Step 2 is to design experiments to develop a model of the output as a function of the factors. Step 3 is to analyze the data

from the experiments. Step 4 is to set the controllable factors to desirable settings. The procedure may iterate through steps 1–3 several times before a final decision is reached for step 4.

For clarity, a simple manufacturing example will be used. Suppose that a plastic structural beam is produced by extrusion followed by a knife cut. See Figure 1 for a diagram of the system.

Robust design procedure for the plastic beam example:

- Step 1: Factor and response identification.
Initially, one might identify the factors of interest to be (x_1) knife speed, (x_2) knife temperature, (x_3) extruder screw speed, (x_4) knife position (distance from extruder die), (z_1) ambient temperature, (z_2) percentage of recycled plastic in the material, (z_3) barrel temperature of the extruder, and potentially many more. There could be many responses of interest such as the throughput of the system, the length of the beam, or the surface roughness of the end of the beam where the knife cuts off the extrusion. For this example, the surface roughness of the end, (y) , will be used as the response of interest. And clearly, lower roughness is better.
- Step 2: Experimental design.
After selecting the factors of interest, experiments need to be designed that allow for modeling of the roughness with respect to the factors. The experiments will involve setting the factors at various levels and then producing a number of beams at each of these conditions.
- Step 3: Data analysis.
After data are collected (measurements of the surface roughness of the beams produced in the experiments), the roughness and the variability of the surface roughness are modeled as a function of the factors.
- Step 4: Factor setting.
Given the models created in step 3, settings of all controllable factors should be selected to produce consistently low surface roughness at the end of the beams.

Since step 1 is concerned with the specific knowledge of the system in question and step 4 is a direct result of the models developed in step 3, the

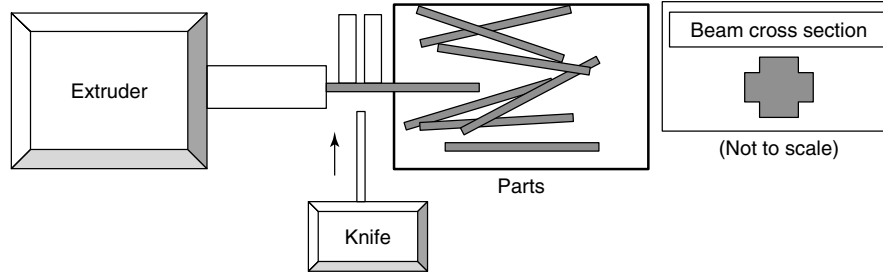


Figure 1 An extruder and knife system for producing a plastic structural beam

robust design literature is essentially the development of specific methods for steps 2 and 3. In the robust design literature, there are two branches of methodology.

The first branch of methodology, called here *robust design for variability from replication*, is used when the variability of interest is caused by mere repeated use of the system. In this scenario, the experimental design techniques are relatively simple. Typically, the designs used are standard first- or second-order experimental designs that are replicated multiple times. After the experiment is run, separate models for the mean and variance of the response are created as functions of the factors in the experiment (see [34]). Numerical optimization methods are then used to determine the settings of the factors that achieve a good balance between achieving a desirable mean response and minimizing the response variation (see [12 or 18]).

The second branch of methodology, called here *robust design for variability from experimentally controllable noise factors*, is used when there exist experimentally controllable noise factors in the system. To understand noise factors, it is first important to understand design factors (sometimes called *control factors*). *Design factors* are factors that are controllable by the system designer. The levels of these factors are not subject to (much) variability during the experiment or even when the system is put into its standard operation mode. For example, in the plastic beam experiment the speed of the knife blade is controlled by a servomotor and can be set with considerable accuracy; thus, knife speed would be considered a design factor. *Noise factors* are factors that affect the response of interest but are randomly varying during the standard operation of the system. In the plastic beam experiment, although the barrel

temperature of the extruder may be affected by various zone temperature settings, it may vary randomly due to other uncontrollable influences that are present both during the robust design experiment and during typical use of the extruder. This would make it a noise factor, but since it cannot be controlled it simply adds to the replication error of the response. *Experimentally controllable noise factors* are factors that, although randomly varying during typical use of the system, can be controlled during the experiments in robust design procedure (Step 2). Again using the plastic beam example, the percentage of recycled plastic might be able to be controlled during the robust design experiment. However, as the process is put into production, this percentage may be subject to random variation from the plastic supplier. If this is the case, then the percentage of recycled plastic is an experimentally controlled noise factor and one of the goals of the robust design procedure would be to make the surface roughness “robust to” (i.e., “unaffected by”) the percentage of recycled plastic in the raw material. Since it can be controlled during the experiment, this factor can be included in the model for the mean and the variance of the surface roughness.

In the next two sections, each of these two branches of robust design methodology is described and explained. Each section begins with a model and proceeds to discuss experimental design and analysis techniques that are appropriate to accurately fitting that model and then using the fitted model to achieve a robust design.

Robust Design for Variability from Replication

When there are no experimentally controllable noise

factors in a system, the most reliable way to measure the variability of the system is to run replicate trials at the various settings of the design factors in the experiment. There are a few available methods for estimating variability without replication (*see Bergman–Hynén Method for Dispersion Effects; Box–Meyer Method for Dispersion Effects*). These methods are less reliable, but can be used if the experimental budget precludes the use of replicates. The goal of robust design in this scenario is to find settings of the design factors that allow for a desirable mean response with low variability. In a typical RSM procedure, the mean response is assumed to depend on the factors, but the variability of the response is assumed to be constant. If these assumptions hold, setting the factors to optimize the mean is in fact the best choice since variability cannot be minimized. However many engineers, most notably Genichi Taguchi (*see* [32]), recognized that often the variability of the system also depends on the design factor settings (*see Dispersion Effects; Performance Measures for Robust Design*). As this became more widely considered, it became clear that the variability of the system can cause substantial quality problems even when the mean response is optimized. This led to the explosion of interest and literature on the robust design procedure.

Since the robust design procedure calls for both the mean and the variability of the response, say y , to depend on the levels of k design factors, the model for the i th response in a robust design experiment is as follows:

$$y_i = \mu(\mathbf{x}_i) + \sigma(\mathbf{x}_i)\varepsilon \quad (1)$$

where $\mathbf{x}_i$ is a k vector such that the j th element of $\mathbf{x}_i$ is the level of the j th design factor in the i th experimental trial, $\mu()$ and $\sigma()$ are unknown functions, and ε is a random error term that follows a standard **normal distribution** (i.e., $\varepsilon \sim N(0, 1)$). Often, $\mu(\mathbf{x}_i)$ is modeled as a first-order linear model so that $E(y_i) = \mu(\mathbf{x}) = \mathbf{x}'\boldsymbol{\beta}$, where $\boldsymbol{\beta}$ is a k vector of parameters. Specifically, the j th element of $\boldsymbol{\beta}$, β_j , is the increase (or decrease if $\beta_j < 0$) in the expected value of the response as factor j is increased by one unit; this is called the *location effect of factor j* . This model is easily extended to a second- or higher-order model by adding additional terms as linear regressors. Since factors are often thought to have a multiplicative effect on the standard deviation of the response, the model $\text{Sdev}(y_i) = \sigma_i(\mathbf{x}_i) = \exp(\mathbf{x}_i'\boldsymbol{\gamma})$

is often used to express the dependence of the variability of the response on the factors (*see Dispersion Effects; Generalized Linear Models* for other models of variance effects, often called *dispersion effects*). Here $\boldsymbol{\gamma}$ is a k vector of parameters where the exponential function of the j th element, $\exp(\gamma_j)$, is the factor by which the standard deviation of the response increases (or decreases if $\gamma_j < 0$) as factor j is increased by one unit; this is called the *dispersion effect of factor j* . For example, if $\exp(\gamma_2) = 3$, then the standard deviation of the response would triple each time factor 2 is increased by one unit. If $\exp(\gamma_2) = -3$, then the standard deviation of the response would be reduced to a third of its value each time factor 2 is increased by one unit. (For notational convenience, $\mathbf{x}$, $\boldsymbol{\beta}$, and $\boldsymbol{\gamma}$ are expanded to $k + 1$ length vectors by including a 0th term; β_0 , and γ_0 represent the constant terms in the linear models and $x_0 = 1$.)

In this scenario, standard response surface designs such as fractional factorials or central composite designs (*see Fractional Factorial Designs; Central Composite Designs*) are appropriate since the primary requirement of the design is that it supports the linear model being fit. If a first-order model is being fit, a two-level resolution III or IV **fractional factorial** design is typically recommended (*see Factorial Designs, Resolution of* and [24], p. 179). If a second-order model is being fit, a central composite or Box–Behnken design is typically recommended (*see Box–Behnken Designs; Central Composite Designs*; and [24], pp. 321–350). One of the considerations in choosing the design is that each design point will need to be replicated to support the model for the standard deviation. Unless the experimental budget is very tight, running only two replications is not recommended as the estimates of the standard deviation at each design point will be very poor. Typically, between three and five replications are taken. More may be needed if the random error is very large.

By selecting only the design factors for the plastic beam experiment, a simple example of robust design for variability from replication can be constructed. If a first-order model is to be fit for both the mean and the standard deviation of the surface roughness, then a four factor, resolution IV, fractional factorial with eight design points could be run and replicated three times as shown in Tables 1 and 2. The high and low levels of each factor are represented by +1 and -1 respectively. The surface roughness data provided are

4 Robust Design

constructed for the purpose of demonstration and thus no units are given.

The analysis begins by calculating the row averages and natural log of the row standard deviations across the three replications for each of the eight design points or *trials*. Standard linear regression can be used to fit the average response vector, $\bar{\mathbf{y}}$, to the expected value model, $\bar{\mathbf{y}} = \mathbf{X}\boldsymbol{\beta} + \mathbf{e}_M$, and then to fit the natural log of the standard deviation to the variability model $\ln(\text{Stdev}(\mathbf{y})) = \mathbf{X}\boldsymbol{\gamma} + \mathbf{e}_{sd}$, where $\mathbf{X}$ is the design matrix shown in Table 1 with an additional column of ones appended for the constant, $\mathbf{e}_M$ and $\mathbf{e}_{sd}$ are random error vectors. In this case, the estimates for the location and dispersion effects are $\hat{\boldsymbol{\beta}}' = (79.40 \ -2.58 \ -8.40 \ -0.32 \ 2.10)$ and $\hat{\boldsymbol{\gamma}}' = (-0.23 \ -1.07 \ -0.01 \ -0.03 \ 1.06)$, respectively.

Deciding which of these effects are statistically significant can be a challenging problem. The simplest analysis is to assume that $\mathbf{e}_M$ and $\mathbf{e}_{sd}$ are normally distributed and then to use either a “normal plot of effects” (see **Half-Normal Plot**) or some version of Lenth’s method (see **Lenth’s Method for the Analysis of Unreplicated Experiments** and [36]) to select the significant effects for each model. The

normal plot of effects method is simply a normal **probability plot** of all estimable effects (including estimable interactions, see **Interactions**). Any factor effect that is significantly different than zero will appear as an **outlier** on this normal plot (see [8] pp. 128–131). In Figures 2 and 3, the normal plot of effects for the location and dispersion effects are shown. The literature, notably Nair and Pregibon [26], Vining and Myers [34], Nelder and Lee [27], and Engel and Huele [14], provides a number of more sophisticated methods for determining the significance of the effects that do not necessarily assume normally distributed error and account for the fact that the common regression assumption of constant variance of observations is violated if there are any significant dispersion effects. These methods are very worthwhile if there are no obvious conclusions that can be drawn from the simple normal plotting method. However, if the normal plot method provides very clear delineation of significant location and dispersion effects, it is unlikely that any of these more sophisticated methods will come to substantially different conclusions.

In this example, chiefly because it is contrived, conclusions are easily obtained and it can be seen that factors 1, 2, and 4 have significant location effects, and factors 1 and 4 have dispersion effects. Knife speed (x_1) has a negative effect on both the surface roughness and the variability of the surface roughness and thus should be set to its high level. Since variability of a response is often positively correlated with the mean of the response, it is common to have factors that have both a location and a dispersion effect. Knife temperature (x_2) should clearly be set to its high level since it has a negative effect on the surface roughness and seems to have little, if any, effect on the variability. The knife position (x_4) is a

Table 1 The design matrix for the robust design experiment for replication error

<i>Trial number</i>	x_1	x_2	x_3	x_4
1	−1	−1	−1	−1
2	+1	−1	−1	+1
3	−1	+1	−1	+1
4	+1	+1	−1	−1
5	−1	−1	+1	+1
6	+1	−1	+1	−1
7	−1	+1	+1	−1
8	+1	+1	+1	+1

Table 2 The three replications, the sample mean, and natural log of the standard deviation for each design point

<i>Replication 1</i>	<i>Replication 2</i>	<i>Replication 3</i>	$\bar{\mathbf{y}}$	$\ln(\text{Stdev}(\mathbf{y}))$
87.2	88.4	89.5	88.4	0.140
90.2	89.0	88.7	89.3	−0.231
83.8	74.8	66.6	75.1	2.152
66.0	66.0	66.1	66.0	−2.852
91.0	95.7	87.7	91.5	1.391
82.0	82.0	81.8	81.9	−2.159
73.1	73.4	72.0	72.8	−0.305
69.1	69.9	71.1	70.0	0.007

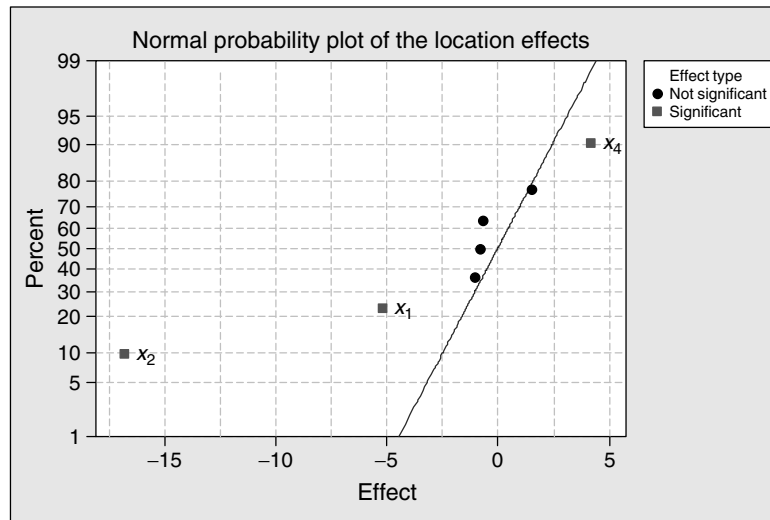


Figure 2 A normal plot of the location effects

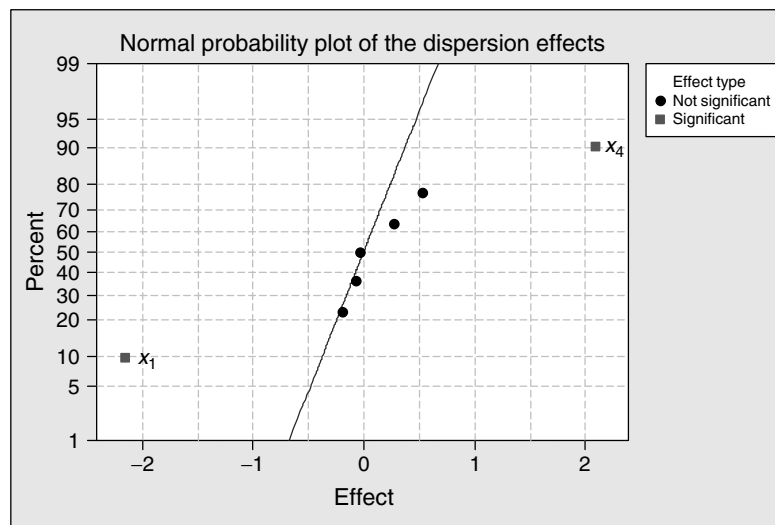


Figure 3 A normal plot of the dispersion effects

particularly interesting effect since it has a positive effect on variability, but no significant effect on the expected value of the surface roughness. Clearly then, it should be set at its low level to reduce the variability. The extruder screw speed (x_3) has little or no effect on surface roughness or its variability. This factor should probably be set the high level to increase production rate.

In many cases, especially when using second-order models, the proper factor settings for optimizing both

mean and variability are not so clear. In fact, many times setting a factor to its high level may improve the location of the response mean but also increase the variability. In such cases, trade-offs must be made to find factor settings that give a desirable mean response and still reduce variability. Once again there is a substantial literature not only on how to formulate the trade-off problem, but also what algorithms will efficiently and effectively find good or even optimal solutions to the problem posed.

6 Robust Design

Formulation of the trade-off problem depends heavily on whether the response is “the smaller, the better” (as in the case of surface roughness), “the larger, the better” (as in the case of throughput of the system), or “target value is best” (as in the case of length of the plastic beam). In Vining and Myers [34], the first paper to formally pose the problem in this way, they chose the following formulations:

1. “The smaller, the better”: minimize $\mu(\mathbf{x})$, subject to $\sigma(\mathbf{x}) = \sigma_0$, where σ_0 is chosen to be acceptably small.
2. “The larger, the better”: maximize $\mu(\mathbf{x})$, subject to $\sigma(\mathbf{x}) = \sigma_0$, where σ_0 is chosen to be acceptably small.
3. “Target value is best”: minimize $\sigma(\mathbf{x})$, subject to $\mu(\mathbf{x}) = \mu_0$, where μ_0 is the target value.

They show how an optimization method provided by Myers and Carter [22] can provide solutions to the posed problems. Subsequent to this paper, many authors weighed in not only on the formulation of the problems but also on various mathematical programming techniques for solving these problems. Examples of these papers are Del Castillo and Montgomery [13], Lin and Tu [19], Copeland and Nelson [11], Kim and Lin [17], Del Castillo, Fan, and Semple [12], Fan [15], Tang and Xu [33], and Köksoy and Doganaksoy [18].

Robust Design for Variability from Experimentally Controlled Noise Factors

When noise factors (which by definition are uncontrolled in the standard operating conditions of the

system) can be controlled in a robust design experiment, a different set of experimental design and analysis tools are needed to achieve steps 2 and 3 of the robust design procedure. These tools focus on reducing the variability of the response that is related to the uncontrolled variability of the noise factors during the standard operation of the system (i.e., after the robust design experiment is finished and the system is returned to a production mode.)

Again, Genichi Taguchi (see [31]) was the person who recognized that including noise factors in the experiment allows for design factors to be set not only to optimize the mean response, but also to minimize the variability of the response due to those noise factors. He called this practice *parameter design*, but here it is called *robust design*. A simple way to think about robust design with noise factors is to study the experimental design recommended by Taguchi (and many others) for these experiments (see Table 3). Taguchi advocated what is called an inner and outer (sometimes called *product* or *crossed*) array design (see **Product Array Designs**). This design begins by separating the noise factors from the design factors and creating an experimental design (usually an **orthogonal array**) for only the design factors. Taguchi calls this the *inner array* or *design array* and it can be thought of as a representative sample of the possible designs that could be chosen for the system. The design factors are listed in columns so that each design point in the inner array corresponds to a row. Next, a separate *outer array* or *noise array* is constructed for the noise factors. The outer array can be thought of as a representative sample of all likely noise conditions (levels of the noise factors) that the system is likely to encounter in its standard operation. Each noise factor appears in a row so that

Table 3 A crossed (inner and outer) array experiment

Design array trial number	Noise array trial number				1	2	3	4	$\bar{y}$	
					z_2					
	x_1	x_2	x_3	x_4	z_1					
1	−1	−1	−1	−1		81.9	55.5	119.2	98.4	88.8
2	+1	−1	−1	+1		71.8	81.7	96.8	101.3	87.9
3	−1	+1	−1	+1		85.7	56.7	93.5	62.8	74.7
4	+1	+1	−1	−1		56.4	66.1	68.7	68.5	64.9
5	−1	−1	+1	+1		95.0	71.5	120.2	97.6	96.1
6	+1	−1	+1	−1		67.0	59.9	95.4	91.1	78.4
7	−1	+1	+1	−1		71.8	52.5	93.8	62.6	70.2
8	+1	+1	+1	+1		65.0	67.1	79.7	80.4	73.1

each design point in the noise array is a column. The overall design is setup like a matrix of observations that runs the entire outer array for the noise factors at each design point for the design factors in the inner array. In the plastic beam example, (z_1) ambient temperature and (z_2) percentage of recycled plastic in the material are both subject to variability in the actual process, but will be controlled noise factors in the experiment. Table 3 shows an inner and an outer array design that could be used for this robust design experiment. The observations are shown in the table. So for example, 65.0 is in the last row and the first column and it is the surface roughness observation for a beam that was produced with all the design factors at their high level and all the noise factors at their low levels. As in the previous section, the surface roughness data provided are constructed for the purpose of demonstration and thus no units are given.

By looking at the crossed array experiment, one can easily visualize the goal of a robust design experiment that includes noise factors. Specifically, the goal is to choose settings of the design factors that give a consistently desirable response across all noise conditions (a row of observations). It is possible to treat this problem in the same way as the “robust design for variability from replication”, and try to simply minimize the standard deviation across all the noise conditions while setting the mean response to a desirable value. However, calculation of a standard deviation across a row implies that observations in the row are random draws from a stable distribution. This is clearly not true since the settings of the noise factors are not randomly chosen, but intentionally set to various high and low levels across the row. Thus, calculating a standard deviation discards the information that is available about the levels of the noise factors for each of these four trials in the row.

A better analysis is offered by a model of the response with respect to both control and noise factors. This approach can be found in Shoemaker *et al.* [29], Myers *et al.* [23], Lucas [20], and Myers *et al.* [25]. Since the noise factors will not be under the control of the designer in the standard operating mode of the system, the noise factors in the model must be treated in a different way. In particular, the level of each noise factor in standard operating mode can be thought of as a random variable with some distribution. Historical data can be used to determine a model distribution for each noise factor. Adding the noise factors to the model for the

experimental data gives the following:

$$y_i = \mu(\mathbf{x}_i, \mathbf{z}_i) + \sigma(\mathbf{x}_i, \mathbf{z}_i)\varepsilon \quad (2)$$

where $\mathbf{z}_i$ is a k vector such that the t th element of $\mathbf{z}_i$ is the level of the t th noise factor in the i th experimental trial. Although other models are possible, a first-order model with two-factor interactions is the most common model to be used for robustness experiments. This is primarily due to the key role of design-by-noise factor interactions. To highlight this role, the model with k design factors and m noise factors will be written as follows:

$$\begin{aligned} y = & \beta_0 + \sum_{s=1}^k \beta_s x_s + \sum_{s=1}^{k-1} \sum_{t \neq s} \beta_{st} x_s x_t + \sum_{t=1}^m \alpha_t z_t \\ & + \sum_{s=1}^{m-1} \sum_{s \neq t} \alpha_{st} z_s z_t + \sum_{s=1}^k \sum_{t=1}^m \lambda_{st} x_s z_t + \sigma(\mathbf{x}, \mathbf{z})\varepsilon \end{aligned} \quad (3)$$

where the λ_{st} parameter is associated with the interaction between the s th design factor and the t th noise factor. There is a β parameter associated with each **main effect** and each two-factor interaction for the design factors and an α parameter associated with each main effect and each two-factor interaction for the noise factors. When noise factors are included in the robust design experiment, it is often thought that these noise factors are the primary causes of variability. Therefore, it will be assumed here that when noise factors are included in the experiment, the random error does not depend on the factor levels (i.e., $\sigma(\mathbf{x}_i, \mathbf{z}_i) = \sigma$). In operating mode, each noise factor, z_t , will have a variance, $\text{Var}(z_t)$. For ease of discussion, it will be assumed here that noise factor distributions are independent of each other (so that $\text{Cov}(z_s, z_t) = 0$ for all s and t) and that none of the noise factors interact with each other (i.e., $\alpha_{st} = 0$ for all s and t). Thus, the variance of the response can be expressed as:

$$\begin{aligned} \text{Var}(y) = & \text{Var} \left[\sum_{t=1}^m \alpha_t z_t + \sum_{s=1}^k \sum_{t=1}^m \lambda_{st} x_s z_t \right] + \sigma^2 \\ = & \text{Var} \left[\sum_{t=1}^m \left(\alpha_t + \sum_{s=1}^k \lambda_{st} x_s \right) z_t \right] + \sigma^2 \\ = & \sum_{t=1}^m \left(\alpha_t + \sum_{s=1}^k \lambda_{st} x_s \right)^2 \text{Var}(z_t) + \sigma^2 \end{aligned} \quad (4)$$

From equation (4), it can be seen that if there are no design-by-noise factor interactions, (i.e., $\lambda_s = 0$ for all s), then the variance of the response does not depend on the design factor settings ($\mathbf{x}$). Thus, without design-by-noise factor interactions, all settings of the design factors are equally robust to the noise factors and the design factors should be set to optimize the mean response as in traditional RSM.

However, if there are design-by-noise factor interactions (i.e., $\lambda_s \neq 0$ for some s), then there must exist some set of design factor settings that reduce the variance of the response. For example, if there exists a set of design factor levels ($x_1, x_2, \dots, x_k$), such that $\sum_{s=1}^k \lambda_{st} x_s = -\alpha_t$ for each t , the variance of the response is reduced to σ^2 , the replication error. Since the existence of design-by-noise factor interactions is critical to reducing variability, the focus of robust design experiments is often heavily focused on **estimation** of these particular interaction effects.

In addition to the robustness, the mean or average response is also of interest. Taking the same assumptions as before (i.e., $\text{Cov}(z_s, z_t) = 0$ and $\alpha_{st} = 0$ for all s and t), and letting the noise variables be coded such that the mean value in production mode is 0 (i.e., $E(z_t) = 0$), then the expected value of y is as follows:

$$E(y) = \beta_0 + \sum_{s=1}^k \beta_s x_s + \sum_{s=1}^{k-1} \sum_{t \neq s} \beta_{st} x_s x_t \quad (5)$$

Optimization techniques can again be used to search for the settings for the design factors that will produce a desirable mean response with low variability (see [21, 28]). For example, suppose that we wish to solve the problem in equation (6) for the plastic beam example, where y represents the surface roughness:

$$\begin{aligned} \text{Minimize}_{\mathbf{x}} : \quad & E(y) = \beta_0 + \sum_{s=1}^k \beta_s x_s + \sum_{s=1}^{k-1} \sum_{t \neq s} \beta_{st} x_s x_t \\ \text{Subject to :} \quad & \text{Var}(y) = \sum_{t=1}^m \left(\alpha_t + \sum_{s=1}^k \lambda_{st} x_s \right)^2 \\ & \times \text{Var}(z_t) + \sigma^2 \leq \sigma_0^2, \text{ where } \sigma_0^2 \\ & \text{is chosen to be acceptably small.} \end{aligned} \quad (6)$$

To solve equation (6), estimates for the variance of the noise factors in standard operating mode must be

obtained from historical data. In addition, the robust design experiment must provide estimates of the α , β , and λ parameters. A more in-depth discussion of the issues of experimental design will be given later, but note for now that a crossed array can provide estimates of all these parameters as long as the design array is of at least resolution V and the noise array is of resolution III. If interactions among the design factors are believed to be unimportant (i.e., $\beta_{st} \approx 0$ for all s and t), then a design array of resolution III is sufficient.

In the plastic beam example, it will be assumed that the only interactions are design-by-noise interactions. To estimate parameters from the robust design experiment, the regression model $\mathbf{y} = \mathbf{W}\boldsymbol{\theta} + \boldsymbol{\epsilon}$ is used, where $\boldsymbol{\theta}' = (\boldsymbol{\beta}' \boldsymbol{\alpha}' \boldsymbol{\lambda}')$ and $\mathbf{W}$ is a 32×15 linear regression model matrix that is a function of the design array and the noise array. It has 32 rows corresponding to the number of observations in the crossed array. It has one column of ones for the constant, a column for each of the four design factors, a column for each of the two noise factors and a column for each of the eight design-by-noise interactions. More generally, let k and m be the number of design factors and noise factors respectively and let p and q be the number of design points in the design and noise arrays respectively. The $\mathbf{W}$ matrix is a $pq \times (1 + k + m + km)$ made up of four parts as follows $\mathbf{W} = [\mathbf{1} \ \mathbf{X} \otimes \mathbf{1}_q \ \mathbf{1}_p \otimes \mathbf{Z} \ \mathbf{U}]$, where $\otimes$ is the Kronecker product, $\mathbf{1}$ is a column of ones, $\mathbf{X}$ is the design matrix of the design array, $\mathbf{Z}$ is the design matrix for the noise array, and $\mathbf{U}$ is a model matrix with columns formed by element-by-element multiplication of the each of the k design factor columns in $\mathbf{X} \otimes \mathbf{1}_q$ with each of the m noise factor columns in $\mathbf{1}_p \otimes \mathbf{Z}$. If design-by-design factor or noise-by-noise factor interactions were supported by the experimental design, these could also be appended onto the $\mathbf{X}$ and $\mathbf{Z}$ matrices respectively. Least-squares estimates for the parameters are now obtained by $\hat{\boldsymbol{\theta}} = (\hat{\boldsymbol{\beta}}' \hat{\boldsymbol{\alpha}}' \hat{\boldsymbol{\lambda}}')' = (\mathbf{W}'\mathbf{W})^{-1} \mathbf{W}'\mathbf{y}$. All of these estimates could be inserted into the optimization problem in equation (6). However, many of these effects may not be significant in which case the problem can be substantially simplified.

A normal plot of all the effects from the robust design experiment, shown in Figure 4, shows that both noise factor main effects and the main effects of design factors x_1 , x_2 , and x_4 are significant. The only design-by-noise factor interactions that are significant

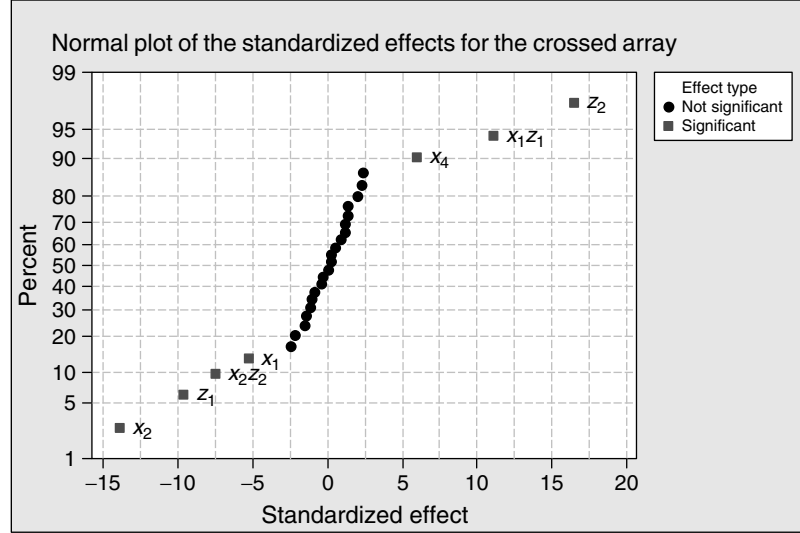


Figure 4 A normal plot of the effects for the crossed array

are the interactions between x_1 and z_1 and the interaction between x_2 and z_2 . In Table 4, values for the historical variances of the noise factors are given. After reducing the regression to include only the significant effects, the mean square error from the regression can be used as an estimate of σ^2 ; in this case $\hat{\sigma}^2 = 22.5$.

Using the values from Table 4 and the reduced model, the robust design problem becomes

$$\begin{aligned} \text{Minimize : } E(y) &\approx 79.24 - 3.18x_1 \\ &\quad - 8.53x_2 + 3.69x_4 \end{aligned}$$

$$\begin{aligned} \text{Subject to : } \text{Var}(y) &\approx (\hat{\alpha}_1 + \hat{\lambda}_{11}x_1)^2 \text{Var}(z_1) \\ &\quad + (\hat{\alpha}_2 + \hat{\lambda}_{22}x_2)^2 \text{Var}(z_2) + \hat{\sigma}^2 \\ &\approx (-5.88 + 6.84x_1)^2(6) \\ &\quad + (10.14 - 4.59x_2)^2(4) \\ &\quad + 22.5 \leq 100. \end{aligned} \quad (7)$$

If it is assumed that $E(z_i) = 0$ and that the design factors can only be set to values between -1 and 1 , then the following logic can be used to solve equation (7). If $x_1 = 0.86$, then the first term in the variance model is zero. The second term cannot be set to zero unless $x_2 > 1$, so to minimize variance and stay within the acceptable range, $x_2 = 1$, this leads

to $\text{Var}(y) \approx 145.71 \not\leq 100$. The mean surface roughness is reduced by setting $x_1 = x_2 = 1$ and $x_4 = -1$. Clearly, to reduce the variance below 100, one needs to explore the idea of further increasing x_2 , the knife temperature. However, if the constraints are observed, the robust design would be $x_1 = 0.86$, $x_2 = 1$, and $x_4 = -1$, alternatively, one could choose the settings, $x_1 = x_2 = 1$ and $x_4 = -1$ which would have a slightly lower mean but would have slightly higher variance. The trade-off between low mean and low variance should be made by the engineers knowledgeable about the consequences of surface roughness in the final application of the plastic beams. Again the conclusions here are relatively clear since this example was constructed for discussion purposes. However, if the models in equation (7) were more complicated, optimization techniques could again be used to solve the problem.

A simple, graphical method that can often be used to reduce the variance function involves making

Table 4 Historical variances for the noise variables in standard operating mode

Noise variable	Historical variance in operating mode
z_1	6
z_2	4

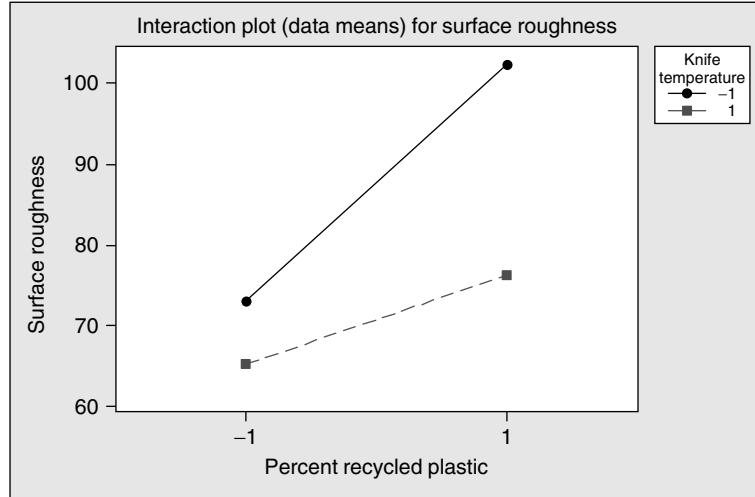


Figure 5 The interaction plot for knife temperature and percent recycled plastic

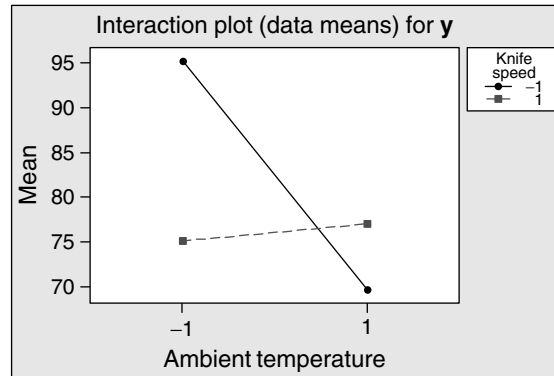


Figure 6 The interaction plot for knife speed and ambient temperature

and interpreting the design-by-noise factor interaction plots. In Figures 5 and 6, the interaction plots for the two significant design-by-noise interactions are shown. Figure 5 shows that the effect of changing ambient temperature, z_2 , from the low level to the high level is substantially smaller if the knife speed, x_2 , is set to its high level. It suggests that if $x_2 > 1$, the effect of ambient temperature might even be further reduced. Figure 6 suggests that the effect of percentage of recycled plastic may be close to zero if the knife temperature is set just below its high level. These are exactly the same conclusions that were made when equation (7) was solved. Thus for many cases, sophisticated optimization methods are not

necessary, simply looking at normal plot of effects and the interaction plots for the significant design-by-noise factor interactions may be sufficient to find the design factor settings that minimize the effects of the noise factors and thus give a robust design.

Experimental Design with Noise Factors

There are three primary considerations that are important when designing a robustness experiment that includes controllable noise factors. These are as follows:

1. Estimability of the design-by-noise factor interactions.
2. Restricted randomization (split-plot designs) (*see Randomization in Experimental Designs; Split-Plot Designs*).
3. Precision of the design-by-noise factor interaction estimators.

In standard experimental designs like fractional factorials, it is usually assumed that main effects are the most important effects, two-factor interaction effects are the next most important and so forth. Design criteria like resolution and minimum aberration (*see Minimum Aberration, Factorial Designs, Resolution of*, and [16]) were developed under this assumption and seek to keep the estimators of these important effects from being confounded together. However, as we have previously seen, estimation of the design-by-noise factor interactions is a primary goal of a robust design experiment. This means that design-by-noise factor interactions are more important than other two-factor interactions and maybe even as important as the main effects. This realization has led to a large literature on designs that allow for estimation of all design-by-noise factor interactions. It is well known that crossed arrays (where both design and noise arrays are at least resolution III) will never have design-by-noise factor interactions aliased (*see Aliasing in Fractional Designs*) with any main effects or any other two-factor interactions. Therefore, as Taguchi originally suggested, crossed arrays are good design for robust design experiments. However, Box and Jones [7] and Borkowski and Lucas [5] noted that there may be designs with fewer trials, often called *combined arrays*, that also protect the design-by-noise interactions from confounding. Wu and Hamada [35, pp. 484–492] provide a table of efficient robust design experiments that do not confound design-by-noise factor interactions with main effects or two-factor interactions.

The second major design consideration has to do with the order in which the trials in an experiment are run and how that affects the analysis of the data. Traditional analysis of experiments assumes that the trials in the experiment are run in a random order. However, in many (if not most) experiments, the order of the trials is determined by convenience. For example, observe the plastic beam experiment in Table 3, it is easy to imagine that all of the design factors can be changed from their low to high levels

in just a few minutes. The noise factors are different. When changing the ambient temperature, it might take an hour or two before all the equipment in the room actually reaches a steady-state temperature. Similarly, changing from a low percentage of recycled plastic to a high percentage or *vice versa* might require 15–20 min to allow all the plastic that is currently in the extruder to be replaced by the new blend. Thus, it would require much less time to run the 32 trials if the ambient temperature and the percent recycled plastic are both set to the low level and then all eight observations in the first column are taken (in random order). Next, one might increase the percentage of recycled plastic and then run all eight trials in the third column. Each of the other two columns of observations would then be taken in turn.

Experiments like this where certain factors are held constant for portions of the experiment are called *split-plot experiments* (*see Split-Plot Designs; [7, 3]*). Split-plot experiments require a somewhat more complicated analysis since some effect estimators have high variance and other effect estimators have lower variance. Specifically, the estimators of the effects of the factors that are run in a random order will have lower variance than the estimators of the effects of the factors that are held constant for large parts of the experiment. If the plastic beam experiment is run in the order suggested above, the noise factor effect estimators will have high variance since they are held constant while the entire column of observations is taken. The design factor effect estimators and the design-by-noise factor interactions will have lower variance. Thus, running this design in a convenient order has not only reduced the amount of time that it takes to run the experiment, but it also allows certain effects to be estimated more precisely. If the normal plot method of analysis is used, then the only difference in the analysis is that both the noise factor main effects and the noise-by-noise factor interaction should be removed from the normal plot since they are known to have higher variance than the other effects (*see [4]*).

In some situations, it may be the design factors that are held constant while all the noise factor settings across the row are run in a random order. In this case, it would be the design factor main effect estimators and all the design-by-design factor interactions that would have higher variance and would need to be removed from the normal plot of effects. If there are enough of these removed effects, they can be put on

a separate normal plot since all of these estimators have the same (higher) variance. Whether it is the design factors or noise factors that are held constant, it is clear that crossed arrays will often be run as split-plot experiments. However, several authors have shown that there often exist designs smaller than the corresponding crossed array that can also be run as split-plot designs. Bingham and Sitter [2, 3] provide tables of efficient split-plot designs and advice for their use and analysis.

Finally the last consideration for design of robustness experiments occurs when the experiment is being run as a split-plot experiment. The issue is whether the design-by-noise factor estimates have the higher or lower variance. Since the design-by-noise factor interactions are so important for finding a robust design, it would be best if the estimators for these effects had low variance. A crossed array, when run as a split-plot design (with either the design factor settings or the noise factors settings held constant), will always provide estimates of the design-by-noise factor interactions at the lower variance (see [7]). Other split-plot designs may also achieve this goal, but choosing a crossed array ensures this property.

In summary, crossed arrays, as originally suggested by Taguchi for robust experiments, have remarkably good properties for estimating design-by-noise factor interactions without confounding and with high precision when run as a split-plot design. Thus, when running a robust design experiment with controllable noise factors, a crossed array will be relatively easy to design and will have good design properties. However, by searching the literature for a more sophisticated design, one can often find a design that has fewer observations than the crossed array and many, if not all, of its good design properties.

References

- [1] Ankenman, B.E. & Dean, A.M. (2003). Quality improvement and robustness via design of experiments, in *Handbook of Statistics*, R. Khattree & C.R. Rao, eds, Elsevier Science, Amsterdam, Vol. 22, Chapter 8, pp. 263–317.
- [2] Bingham, D. & Sitter, R.R. (2003). Fractional factorial split-plot designs for robust parameter experiments, *Technometrics* **45**, 80–89.
- [3] Bingham, D. & Sitter, R.R. (1999). Minimum-aberration two-level fractional factorial split-plot designs, *Technometrics* **41**, 62–70.
- [4] Bisgaard, S. (2000). The design and analysis of $2^{k-p} \times 2^{q-r}$ split plot experiments, *Journal of Quality Technology* **32**, 39–56.
- [5] Borkowski, J.J. & Lucas, J.M. (1997). Designs of mixed resolution for process robustness studies, *Technometrics* **39**, 63–70.
- [6] Box, G.E.P. & Jones, S. (1990). Designing products that are robust to the environment, Technical Report 56, Center for Quality and Productivity Improvement, University of Wisconsin.
- [7] Box, G.E.P. & Jones, S. (1992). Split-plot designs for robust product experimentation, *Journal of Applied Statistics* **19**, 3–26.
- [8] Box, G.E.P. & Draper, N.R. (1987). *Empirical Model Building*, John Wiley & Sons, New York.
- [9] Box, G.E.P., Hunter, W.G. & Hunter, J.S. (1978). *Statistics for Experimenters*, John Wiley & Sons, New York.
- [10] Box, G.E.P. & Wilson, K.B. (1951). On the experimental attainment of optimum conditions, *Journal of the Royal Statistical Society. Series B* **13**, 1–45.
- [11] Copeland, K.A. & Nelson, P.R. (1996). Dual response optimization via direct function minimization, *Journal of Quality Technology* **28**, 331–336.
- [12] Del Castillo, E., Fan, S. & Semple, J. (1997). The computation of global optima in dual response systems, *Journal of Quality Technology* **29**, 347–353.
- [13] Del Castillo, E. & Montgomery, D.C. (1993). A nonlinear programming solution to the dual response problem, *Journal of Quality Technology* **25**, 199–204.
- [14] Engel, J. & Huele, A.F. (1996). A generalized linear modeling approach to robust design, *Technometrics* **38**, 365–373.
- [15] Fan, S. (2000). A generalized global optimization algorithm for dual response systems, *Journal of Quality Technology* **32**, 444–456.
- [16] Fries, A. & Hunter, W.G. (1980). Minimum aberration 2^{k-p} designs, *Technometrics* **22**, 601–608.
- [17] Kim, K. & Lin, D.K. (1998). Dual response surface optimization: a fuzzy logic approach, *Journal of Quality Technology* **30**, 1–10.
- [18] Köksoy, O. & Doganaksoy, N. (2003). Joint optimization of mean and standard deviation using response surface methods, *Journal of Quality Technology* **35**, 239–252.
- [19] Lin, D.K. & Tu, W. (1995). Dual response surface optimization, *Journal of Quality Technology* **27**, 34–39.
- [20] Lucas, J.M. (1994). How to achieve a robust process using response surface methodology, *Journal of Quality Technology* **26**, 248–260.
- [21] Miró-Quesada, G. & Del Castillo, E. (2004). Two approaches for improving the dual response method in robust parameter design, *Journal of Quality Technology* **36**, 154–168.
- [22] Myers, R.H. & Carter, W.H. (1973). Response surface techniques for dual response systems, *Technometrics* **15**, 301–317.

-
- [23] Myers, R.H., Khuri, A.I. & Vining, G.G. (1992). Response surface alternatives to the Taguchi robust parameter design approach, *The American Statistician* **46**, 131–139.
 - [24] Myers, R.H. & Montgomery, D.C. (2002). *Response Surface Methodology: Process and Product Optimization Using Designed Experiments*, 2nd Edition, John Wiley & Sons, New York.
 - [25] Myers, W.R., Brenneman, W.A. & Myers, R.H. (2004). A dual response approach to robust parameter design for a generalized linear model, *Journal of Quality Technology* **37**, 130–138.
 - [26] Nair, V.N. & Pregibon, D. (1988). Analyzing dispersion effects from replicated factorial experiments, *Technometrics* **30**, 247–257.
 - [27] Nelder, J.A. & Lee, Y. (1991). Generalized linear models for the analysis of Taguchi type experiments, *Applied Stochastic Models and Data Analysis* **7**, 107–120.
 - [28] Rajagopal, R., Del Castillo, E. & Peterson, J.J. (2005). Model and distribution robust process optimization with noise factors, *Journal of Quality Technology* **37**, 210–222.
 - [29] Shoemaker, A.C., Tsui, K.-L. & Wu, C.F.J. (1991). Economical experimentation methods for robust design experiments, *Technometrics* **33**, 415–427.
 - [30] Steinberg, D.M. (1996). Robust design: experiments for improving quality, in *Handbook of Statistics*, S. Ghosh & C.R. Rao, eds, Elsevier Science, Amsterdam, Vol. 13, Chapter 7, pp. 199–240.
 - [31] Taguchi, G. (1986). *An Introduction to Quality Engineering*, American Supplier Institute, Dearborn.
 - [32] Taguchi, G. (1985). Quality engineering in Japan, *Communications in Statistics: Theory and Methods* **14**, 2785–2801.
 - [33] Tang, L.C. & Xu, K. (2002). A unified approach for dual response surface optimization, *Journal of Quality Technology* **34**, 437–447.
 - [34] Vining, G.G. & Myers, R.H. (1990). Combining Taguchi and response surface philosophies: a dual response approach, *Journal of Quality Technology* **22**, 38–45.
 - [35] Wu, C.F.J. & Hamada, M. (2000). *Experiments: Planning, Analysis, and Parameter Design Optimization*, John Wiley & Sons, New York.
 - [36] Ye, K.Q., Hamada, M. & Wu, C.F.J. (2001). A step-down Lenth method for analyzing unreplicated factorial designs, *Journal of Quality Technology* **33**, 140–152.

BRUCE E. ANKENMAN

Signal–Response Systems

Introduction

A signal–response system is one where the function of the system requires a response variable to take different values as a result of changes in the signal factor. How well the system performs is determined by the characteristics of the relationship between the signal factor, M , and the response, Y . A number of different types of systems fall into this classification as illustrated by the following three examples.

Example 1: Alashe *et al.* [1] describe an experiment that investigated the engine idling performance of vehicles. Of particular interest was the relationship between the fuel flow (signal) and indicated mean effective pressure or IMEP (response). The IMEP is a measurement of the average pressure exerted on the piston during each power stroke and thus is related to the power generated by the engine. If the engine is efficiently converting fuel to power over the entire practical range of fuel flow then the relationship between fuel flow and IMEP should be approximately linear with an optimal slope that depends on the technical specifications of the engine. The goal of the experiment was to configure the engine so that the fuel flow/IMEP relationship was as close as possible to the optimal relationship.

Example 2: DeMates [2] describes an experiment conducted to improve the performance of an injection-molding machine. The machine was used for different applications that required different amounts of material to be delivered to the mold. It was necessary to be able to control the amount of material delivered by the machine and it was decided that this could best be done by adjusting the high injection pressure. Thus high injection pressure was designated as the signal factor and the amount (weight) of material delivered as the response. Ideally, the operator would be able to set the injection pressure so that exactly the desired amount of material was delivered for a wide range of target weights. Of course in practice there will always be some variability in the amount of material delivered. The goal of the experiment was to set the controllable aspects of the injection-molding machine so that the injection pressure could be used to control the amount of

material delivered with as little variability from the target values as possible.

Example 3: Most **measurement systems** fall into the signal–response classification. For example, atomic-absorption spectroscopy is a standard method of measuring the concentration of metal ions in liquid samples that is used for a wide range of metals. For a particular metal the atomic-absorption spectrometer focuses a beam of UV light through a flame and into a detector. The liquid sample is aspirated into the flame and if the metal is present the flame will absorb light of a wavelength specific to that metal. The amount of light absorbed increases as the concentration of metal in the sample increases. A **calibration curve** is created by measuring the rates of absorption for a set of standards which are solutions that contain known concentrations of the metal. The amount of metal in the sample is estimated by comparing the measured rate of absorption for the sample to the calibration curve. In this example, the concentration of metal in solution is the signal, the rate of absorption is the response and the calibration curve is the signal–response relationship. The characteristics of the calibration curve determine the precision of the estimated concentration of the metal in the sample.

This article starts by discussing how the performance of signal–response systems can be assessed (“Measuring the Performance of a Signal–Response System”), it then discusses approaches to designing and analyzing signal–response systems (“Robust Parameter Design for Signal–Response Systems”), a case study is presented to illustrate these approaches (“Case Study”), and it concludes by giving some selected references for further reading (“Concluding Remarks”).

Measuring the Performance of a Signal–Response System

In investigating a signal–response system, a fundamental issue is how to evaluate the performance of the system. The selection of a suitable measure of performance will depend on the function and possibly the specific features of the system being studied. To assist in this discussion we will consider a simple statistical model where during normal operation of the system the response is a random variable whose

2 Signal–Response Systems

mean and (possibly) variance are functions of the signal factor:

$$E(Y) = f(M) \text{ and } V(Y) = \sigma^2(M) \quad (1)$$

Consider the three examples of signal–response systems from the previous section.

In Example 1 there is a specified target function, say $f_T(M)$, for the signal–response relationship. Thus it is desirable to make the actual signal–response relationship as close as possible to $f_T(M)$ and to make $\sigma^2(M)$ as small as possible over the range of values required for M . A common approach in such situations is to use a quadratic loss function to penalize for departures from the target function – see McCaskey and Tsui [3]. Thus the performance measure becomes the mean square error $[f(M) - f_T(M)]^2 + \sigma^2(M)$, averaged over the desired range of M .

Example 2 can be described as a multiple-target system. For multiple-target systems the signal factor is used to adjust the system to accommodate different target values for the response. There is no specific target function as the shape of the signal–response relationship is not of direct concern but the signal–response relationship must allow the entire set of target values for the response to be achieved. The main objective is to reduce the variance in the response over the set of target values. Thus, an obvious choice is to use the variance, $\sigma^2(M)$, averaged over the set of desired target values with the caveat that the entire set can be achieved.

Example 3 is a measurement system. A measurement system uses the measured value of the response Y to predict the value of the signal M . The size of prediction error depends on the characteristics of the signal–response relationship and the goal is to find settings of the control factors that make these predictions as precise as possible. Miller and Wu [4] show that for a linear measurement system with constant variance – $Y = \beta_0 + \beta_1 M + \epsilon$, $\epsilon \sim N(0, \sigma^2)$ – the expected size of the Fieller’s interval is a monotonically decreasing function of β_1^2/σ^2 . That is, as β_1^2/σ^2 increases the estimates become more precise. Thus, β_1^2/σ^2 is an obvious choice for the performance measure. Although having a linear signal–response relationship and constant variance simplifies the calibration problem, these are not essential for measurement systems. More generally, one could base a performance measure on the “criterion for technical

merit” for measurement systems devised by Mandel [5]: $(f'(M))^2/\sigma^2(M)$. For example, the average of this quantity over the desired range for the signal factor could be used.

Robust Parameter Design for Signal–Response Systems

The investigation of signal–response systems through designed experiments was originally discussed by Taguchi [6] as a special case of robust parameter design (*see Robust Design*). Taguchi identified two types of applications which he called *static* problems and *dynamic* problems. The difference can be illustrated using Example 2 which, as described, is what Taguchi would classify as a dynamic characteristic. However, if the injection-molding machine was only used to produce a single part and the goal was to reduce variation around the target weight then it would be a static characteristic. Miller and Wu [4] used the term *signal–response systems* instead of dynamic characteristics as they felt this was a more appropriate description. In simple terms, the objectives of a **robust design** experiment applied to a signal–response system are to optimize the relationship between the expected response and the signal factor while minimizing the variability in the observed response. Typically, a set of control factors (factors whose levels can be set by the operator) is identified and an experiment is conducted to find the best settings for these factors. Usually the impact that the control factors have on the relationship between the expected response and the signal factor can be assessed with greater precision than their impact on variability. Thus it can be very advantageous to include noise factors in the experiment. Noise factors (*see Factorial Experiments*) represent suspected sources of variability in the response. They are factors that vary during the normal operation of the system but can be set at fixed levels for the purposes of an experiment. By identifying such factors and deliberately varying their levels in an experiment, the chances of identifying settings for the control factors that reduce the impact of these (noise) factors is greatly enhanced.

Two distinct approaches for conducting parameter design experiments have emerged. The first approach defines a specific criterion to be used to evaluate the performance of the system and then an experiment

is conducted with the goal of modeling this criterion as a function of the control factors. This model is used to determine the best settings for the control factors. This is the approach proposed by Taguchi who referred to the performance criteria as (*dynamic*) *signal-to-noise ratios* (see **Signal-to-Noise Ratios for Robust Design**). Conceptually, the experiment is the marriage of two parts. First, there is a primary experiment where the chosen criterion is the response and the levels of the control factors are varied according to a selected design usually referred to as the *inner array*. For each run in the inner array, a secondary experiment is run where the levels of the signal and noise factors are varied according to an “outer array” – the idea is that the outer array runs should encompass the range of conditions that the system is expected to experience. The overall design matrix for the experiment is created by combining each run in the inner array with every run in the outer array (see **Product Array Designs**). The analysis uses a two-stage procedure. The first stage is to estimate the criterion for each run (combination of control factors) in the inner array using the set of observations taken by varying the signal and noise factors according to the outer array. The second stage treats these estimates as the observed response for the inner-array experiment and uses standard statistical techniques to assess the impact of the control factors on the criterion. Clearly, this approach can be applied to criteria other than Taguchi’s signal-to-noise ratios. Miller and Wu [4] use the term *performance measure modeling* to describe the general approach of modeling a chosen performance criterion directly as a function of the control factors.

The second approach is to model the response variable as a function of all of the factors (control, noise, and signal). This fitted model is used to determine the settings of the control factors that will optimize the performance of the system. Often this can be done without identifying a specific performance measure but, if necessary, a criterion can be chosen and the fitted response model can be used to evaluate it for different settings of the control factors. In designing an experiment for this approach, a single design array can be used to set the levels of all of the experimental factors (control, noise, and signal). For reasons of economy the experimenter may wish to use a **fractional factorial** design and having a single design array provides more flexibility in selecting such a design. Note that the fitted model is to be used

to assess how the control and noise factors affect the signal–response relationship and to find settings of the control factors that reduce the effect of the noise factors. As a result, it is important that the design allows signal $\times$ control, signal $\times$ noise, control $\times$ noise, and signal $\times$ control $\times$ noise interactions to be estimated in addition to main effects.

Miller and Wu [4] proposed a modification of the response modeling approach, which they named *response function modeling*. The key idea is to treat the signal–response relationship as the response and to model this relationship as a function of the control and noise factors. An experiment designed for this strategy, should use a single array to vary the control and noise factors and then for each run in that array vary the signal factor over a selected set of levels. The analysis employs a two-stage procedure. For each run in the control/noise array a parametric model is fitted for the signal–response relationship using the set of observations obtained by varying the signal factor. Thus, for each run a set of estimated parameters is obtained. The second stage uses standard techniques to investigate how the control and noise factors affect these parameters. For example suppose that the signal–response relationship is linear and has constant variance. For each run in the control/noise array a linear model of the form $Y = \beta_0 + \beta_1 M + \epsilon$ where $\epsilon \sim N(0, \sigma^2)$ would be fitted and estimates of β_0 , β_1 , and σ^2 obtained. Then the effects of the control and noise factors on each of β_0 , β_1 , and σ^2 would be assessed.

Case Study

The experiment used to investigate the injection-molding system described in the introduction (Example 2) will be used to illustrate the performance measure modeling and the response function modeling approaches. This article will just summarize these analyses – the interested reader is referred to Miller and Wu [4] for a more detailed description. One of the key objectives of this experiment was described by DeMates [2] as “. . . to identify the control factors and optimum settings that would reduce part weight variation over a range of molded part weights”.

Recall that for this example, high injection pressure was designated as the signal factor and the amount (weight) of material delivered to the mold used as the response. Eight levels were used for

4 Signal–Response Systems

Table 1 Control and noise factors for injection-molding case study

Control factor	Levels		Compound noise factor	Levels	
	+1	−1		+1	−1
A: Injection speed	2	0	Melt index	22	18
B: Clamp time (s)	44	49	Percent regrind (%)	0	5
C: High injection time (s)	6.3	6.8	Operator	Experienced	New
D: Low injection time (s)	17	20	Resin moisture	Low	High
E: Clamp pressure (psi)	1900	1700	Day	First	Second
F: Water cooling (F)	70	80			
G: Low injection pressure (psi)	650	550			

the signal factor which corresponded to varying the high injection pressure from 650 to 1000 psi in increments of 50 psi. A preliminary study indicated that the signal–response relationship could be adequately approximated by a quadratic polynomial with constant variance. The experiment investigated seven two-level control factors that represented modifications to the injection-molding system. These factors and the levels used are listed in Table 1. It was thought the system could also be affected by changes in raw materials, operators and environmental conditions. A single two-level compound noise factor (which involved setting the levels of five individual factors) was used to introduce variation from all three of these sources – the factors included in the compound noise factor are also listed in Table 1.

As this experiment was designed with a performance measure modeling approach in mind, an inner array was used to vary the control factors and an outer array was used to set the signal and noise factors. A 2^{7-4} fractional factorial design was used for the inner array and a full factorial (8×2) design was used for the outer array.

The experiment was run over 2 days where the compound noise factor was set to its low level on the first day and to its high level on the second day. On each day all combinations of the control factors in the inner array were tested. For each combination, the signal factor was varied over eight levels and four observations were taken at each level. From the description in DeMates [2] it appears that these four replicates were taken one after another without resetting the signal level between. Thus the variability between replicates should reflect the part-to-part variability in the process. Arguably, it may have been better to spread out replicates over time so that longer-term variation would have also been captured.

Performance Measure Modeling Analysis

The first step is to identify a suitable performance measure. If the primary goal is to reduce variability in part weight for a specific set of target weights then the obvious choice is the variance of part weights with the caveat that the full range of target values must be attainable.

Thus, for each run in the inner array, the variance is estimated using the set of observations obtained by varying the signal and noise factors over the outer array. Notice that there are 64 runs in each set: 32 on day 1 ($N = -1$), which consist of four replicates at eight levels of the signal factor and a similar set of 32 on day 2 ($N = +1$). For each inner-array run a quadratic regression model,

$$Y = \beta_0 + \beta_1 M + \beta_2 M^2 + \epsilon$$

where

$$\epsilon \sim N(0, \sigma^2) \quad (2)$$

is fitted and $\hat{\sigma}^2$ is used as the estimated performance measure – see Table 2. These values are treated as the response for the eight-run control factor design. In actual fact, $\log(\hat{\sigma}^2)$ is analyzed because it better meets the regression model assumptions. The estimated effects for the seven factors are 1.07 (A), 0.67 (B), −0.39 (C), −0.64 (D), 0.14 (E), −0.19 (F), and −0.16 (G). A **half-normal plot** (not given) of these effects indicates that factor A and possibly factors B and D warrant further consideration. From the signs of the effects the settings $A = -1$, $B = -1$, and $D = +1$ are recommended to minimize σ^2 . Finally, it would need to be confirmed that entire set of target weights can be attained for this combination of factor settings.

Table 2 Inner and outer arrays for performance measure modeling analysis

Run	Control factors							$\hat{\sigma}^2$	Run	M	N
	A	B	C	D	E	F	G				
1	1	1	1	1	1	1	1	8.39	1	600	-1
2	1	1	1	-1	-1	-1	-1	19.70	2	650	-1
3	1	-1	-1	1	1	-1	-1	9.06	$\vdots$	$\vdots$	$\vdots$
4	1	-1	-1	-1	-1	1	1	10.40	8	1000	-1
5	-1	1	-1	1	-1	1	-1	4.34	9	600	1
6	-1	1	-1	-1	1	-1	1	9.75	10	650	1
7	-1	-1	1	1	-1	-1	1	1.54	$\vdots$	$\vdots$	$\vdots$
8	-1	-1	1	-1	1	1	-1	3.27	16	1000	1
Inner array									Outer array		

Response Function Modeling Analysis

For this approach a fitted model of the signal-response relationship is obtained for each combination of the control and signal factors (16 in total). Orthogonal polynomials will be used in the quadratic model as to simplify the interpretation of the results. The model can be written as follows:

$$E(Y) = \beta_0 + \beta_1 P_1(M) + \beta_2 P_2(M) + \epsilon, \epsilon \sim N(0, \sigma^2) \quad (3)$$

where the values $P_1(M)$ and $P_2(M)$ are given in Table 3. For this model, β_0 is the mean of Y , β_1 is the linear component of the signal-response relationship, and β_2 is the quadratic departure from linearity. In this case σ^2 represents the part-to-part variability in the process observed when the control and noise factors are fixed. For each of the 16 runs in the combined array estimates are obtained for β_0 , β_1 , β_2 , and σ^2 . For this example, Miller and Wu [4] argued that just the pure error estimate of σ^2 (the estimate based on replicate observations) should be used since residual plots revealed some unusual patterns that indicated an estimate based on **lack of fit** may not be reliable. Table 4 contains the estimated parameters for each run in the control/noise array ($\hat{\sigma}_p^2$ is the pure error estimate of σ^2).

Table 3 Levels for orthogonal polynomials

	Signal factor levels							
	650	700	750	800	850	900	950	1000
$P1(M)$	-7	-5	-3	-1	1	3	5	7
$P2(M)$	7	1	-3	-5	-5	-3	1	7

Each of the estimated parameters is modeled as a function of the control and noise factors. A total of 15 effects are calculated corresponding to the main effects of the control and noise factors and the seven control $\times$ noise two-factor interactions. Half-normal plots were used to evaluate which effects should be included in the models for the parameters – the half-normal plots and model selection details are omitted here but can be found in Miller and Wu [4]. The fitted models for the parameters are

$$\begin{aligned} \beta_0 &= 666.4 + 1.2A - 1.8C + 1.4E - 1.0F \\ &\quad + 1.8G + 1.1N \\ \beta_1 &= 4.79 + 0.16C \\ \beta_2 &= 1.33 + 0.03B - 0.04D - 0.05E - 0.04N \\ &\quad + 0.01A \times N - 0.03F \times N - 0.02G \times N \\ \log(\sigma_p^2) &= 0.12 + 1.10A + 0.22B - 0.21C \\ &\quad + 0.40N + 0.28E \times N + 0.04E \end{aligned} \quad (4)$$

Variability in the response can be decreased in two ways: (a) decrease the impact of the noise factor on the signal-response relationship and (b) decrease the part-to-part variability. First consider decreasing the impact of the noise factor. The fitted models indicate that N affects β_0 and β_2 but not β_1 . The model for β_0 indicates that the signal-response curve is shifted up when N is at the +1 level relative to -1 level. As no control $\times$ noise interaction shows up in this model, no method of reducing this shift has been found. The model for β_2 indicates that the signal-response relationship exhibits more curvature when N is at the -1 level than at the +1 level. As the model also includes terms for the $A \times N$, $F \times N$, and $G \times N$

Table 4 Control/noise array for response function modeling analysis

Run	Control and noise factors								Parameters			
	<i>A</i>	<i>B</i>	<i>C</i>	<i>D</i>	<i>E</i>	<i>F</i>	<i>G</i>	<i>N</i>	$\hat{\beta}_0$	$\hat{\beta}_1$	$\hat{\beta}_2$	$\hat{\sigma}_p^2$
1	1	1	1	1	1	1	1	1	666.5	5.02	1.16	7.78
2	1	1	1	−1	−1	−1	−1	1	664.2	5.12	1.44	4.45
3	1	−1	−1	1	1	−1	−1	1	668.2	4.98	1.22	4.99
4	1	−1	−1	−1	−1	1	1	1	668.4	4.76	1.25	3.53
5	−1	1	−1	1	−1	1	−1	1	666.3	4.66	1.35	0.67
6	−1	1	−1	−1	1	−1	1	1	674.4	4.32	1.32	1.00
7	−1	−1	1	1	−1	−1	1	1	666.6	4.92	1.31	0.21
8	−1	−1	1	−1	1	1	−1	1	664.9	4.90	1.25	0.75
9	1	1	1	1	1	1	1	−1	665.0	4.98	1.33	1.20
10	1	1	1	−1	−1	−1	−1	−1	660.0	4.69	1.48	3.20
11	1	−1	−1	1	1	−1	−1	−1	665.2	4.86	1.26	2.70
12	1	−1	−1	−1	−1	1	1	−1	664.2	4.55	1.54	2.64
13	−1	1	−1	1	−1	1	−1	−1	664.2	4.46	1.39	0.56
14	−1	1	−1	−1	1	−1	1	−1	674.1	4.33	1.36	0.30
15	−1	−1	1	1	−1	−1	1	−1	666.1	4.91	1.30	0.18
16	−1	−1	1	−1	1	1	−1	−1	663.6	5.02	1.29	0.12

interactions, it should be possible to reduce the impact of N by judiciously selecting levels for A , F , and G . Now consider using the $\log(\sigma_p^2)$ model to reduce part-to-part variability. This model clearly indicates that A has the biggest impact and should be set to the -1 level and that factor B should be set to the $+1$ level and C to the -1 level. The model also includes terms for the compound noise factor N and the $E \times N$ interaction. It indicates that part-to-part variation is higher for N at the $+1$ than the -1 level. Setting E to -1 will make part-to-part variation more consistent across levels of N (it does this by decreasing the variance for the $+1$ level but increasing the variance for the -1 level).

Notice that the response function modeling approach has given a much clearer picture of how the control factors affect the nature of the signal–response relationship and the impact of the compound noise factor than the performance measure modeling procedure did. Performance measure modeling only provides an overall assessment of how control factors affect the performance of the system. No information is gained on how specific control factors affect the shape of the signal–response relationship or interact with noise factors. This type of information can be very valuable in suggesting directions for future research.

Concluding Remarks

There is considerable diversity of opinion with regards to methodology that should be used to analyze experiments involving signal–response systems. General discussions of signal–response systems including detailed treatments of both the performance measure modeling and the response function modeling approaches can be found in Miller and Wu [4] and Wu and Hamada [7]. Taguchi’s signal-to-noise ratio approach is still very popular and there are numerous books available that explain this methodology. Two examples are Taguchi *et al.* [8] and Wu and Wu [9]. Lunani *et al.* [10] suggest a method of generalizing Taguchi’s dynamic signal-to-noise ratio approach to make it more flexible. Tsui [11, 12] discusses and develops the response modeling approach. Joseph and Wu [13] propose a methodology that extends Taguchi’s signal-to-noise ratio approach and a methodology that extends response modeling. Miller [14] discusses the response function modeling approach and suggests some additional refinements. Nair *et al.* [15] adapt some functional response techniques to the analysis of signal–response systems. Their approach is similar to response function modeling in that the signal–response relationship is treated as the response but it allows systems to be investigated where no simple parametric model that describes the

signal-response relationship can be identified. Lesperance and Park [16] propose using joint **generalized linear models** for the analysis of signal-response systems adapting the procedure proposed by Nelder and Lee [17].

References

- [1] Alashe, W., Barnes, E., Bath, K., Jacobson, K. & O'Keefe, M. (1994). Engine idle quality robustness using Taguchi methods, *Proceedings of the Symposium on Taguchi Methods*, American Supplier Institute Press, Dearborn, pp. 111–121.
- [2] DeMates, J. (1990). Dynamic analysis of injection molding using Taguchi methods, *Eighth Symposium on Taguchi Methods*, American Supplier Institute Press, Dearborn, pp. 313–331.
- [3] McCaskey, S. & Tsui, K.-L. (1997). Analysis of dynamic robust parameter design experiments, *International Journal of Production Research* **35**, 1561–1574.
- [4] Miller, A. & Wu, C.F.J. (1996). Parameter design for signal-response systems: a different look at Taguchi's dynamic parameter design, *Statistical Science* **11**, 122–136.
- [5] Mandel, J. (1964). *The Statistical Analysis of Experimental Data*, Interscience, New York, pp. 365–369.
- [6] Taguchi, G. (1987). *System of Experimental Design*, Unipub/Kraus International Publications, White Plains.
- [7] Wu, C.F.J. & Hamada, M. (2000). *Experiments: Planning, Analysis and Parameter Design Optimization*, John Wiley & Sons, New York, Chapter 11.
- [8] Taguchi, G., Chowdhury, S. & Taguchi, S. (2000). *Robust Engineering*, McGraw-Hill, New York.
- [9] Wu, Y. & Wu, A. (2000). *Robust Design*, ASME, New York.
- [10] Lunani, M., Nair, V. & Wasserman, G. (1997). Graphical methods for robust parameter design with dynamic characteristics, *Journal of Quality Technology* **29**, 327–338.
- [11] Tsui, K.-L. (1998). Alternatives of Taguchi's approach for dynamic robust design problems, *International Journal of Reliability, Quality and Safety Engineering* **5**, 115–131.
- [12] Tsui, K.-L. (1999). Modeling and analysis of dynamic robust design experiments, *IIE Transactions* **31**, 1113–1122.
- [13] Joseph, V.R. & Wu, C.F.J. (2002). Robust parameter design of multiple-target systems, *Technometrics* **44**, 338–346.
- [14] Miller, A. (2002). Analysis of parameter design experiments for signal-response systems, *Journal of Quality Technology* **34**, 139–151.
- [15] Nair, V., Taam, W. & Ye, K. (2002). Analysis of functional responses from robust design studies, *Journal of Quality Technology* **34**, 355–370.
- [16] Lesperance, M. & Park, S.M. (2003). GLMs for the analysis of robust designs with dynamic characteristics, *Journal of Quality Technology* **35**, 253–263.
- [17] Nelder, J. & Lee, Y. (1991). Generalized linear models for the analysis of Taguchi-type experiments, *Applied Stochastic Models and Data Analysis* **7**, 107–120.

Related Articles

Factorial Experiments; Performance Measures for Robust Design; Product Array Designs; Robust Design; Signal-to-Noise Ratios for Robust Design.

ARDEN MILLER

Signal-to-Noise Ratios for Robust Design

Taguchi's Parameter Design for a Fixed Target

Robust design (RD; *see* **Robust Design**) is a systematic methodology that applies statistical experimental design to improve the design of products and processes. By making product and process performance insensitive (robust) to hard-to-control disturbances (noise), RD simultaneously improves product quality, the manufacturing process, and reliability. The original RD method developed by the Japanese quality consultant, Genichi Taguchi ([1]; is called *parameter design*. Taguchi's 1980 introduction of robust parameter design to several major American industries resulted in significant quality improvements in product and process design (see [2]). Since then, a great deal of research on RD has improved related statistical techniques and clarified underlying principles. Representative examples include Shoemaker *et al.* [3] and Tsui [4, 5].

The most widely applied of Taguchi's robust parameter design methods is static RD, in which the process response has a single, unchanging target value. All of Taguchi's robust parameter design methods assume the response is affected by easily manipulated variables, that is, control factors, as well as hard-to-control variables, that is, noise factors. The aim of static RD is to choose a vector of control variable settings that simultaneously reduces the variability of the process response and sets the response to the desired target value.

The key ideas behind static RD are nicely illustrated by a tile manufacturing example described in Taguchi and Wu [1] and Phadke [6]. The Ina Tile Company was experiencing large variation in its tile dimensions, resulting in unwanted scrap and expense. The engineers first determined that the variability in the response of interest, that is, tile dimensions, was caused by uneven heating of the tiles in the kiln. Tiles in the center of the stack in the kiln were heated to a lesser degree than those placed on the outside of the stack, hence causing the variation in dimension. A redesign of the kiln would solve the problem, but it would be an expensive solution.

Instead, the engineers identified several process factors that could be varied inexpensively. They performed carefully planned experiments to determine which combination of process factors would make tile dimensions robust against the uneven distribution of heat within the kiln. The engineers found that by increasing the lime content of the clay from 1 to 5%, the variability in tile dimensions would be greatly decreased.

The engineers discovered that an easily adjusted control variable, that is, lime content, made tile dimensions robust against a difficult to adjust noise factor, namely the uneven distribution of heat within the kiln. Adjusting the lime content allowed the tile dimensions to meet target specifications with reduced variation.

As the Ina Tile example shows, the goal in parameter design is to improve quality by achieving a desired target level coincident with low variability. More formally, the objective is to minimize the quadratic loss function popularized by Taguchi,

$$L(Y, t) = k(Y - t)^2 \quad (1)$$

where Y is the value of the process response, t the desired target value for Y , and k is a constant called the *quality loss coefficient*, equal to A_0/Δ_0^2 in which Δ_0 is the customer's tolerance around target t , and A_0 represents the loss in dollars or other units when the tolerance Δ_0 is exceeded. The average or expected loss may be decomposed as follows:

$$\begin{aligned} R(\mathbf{C}) &= kE_{\mathbf{N},\epsilon}(Y - t)^2 \\ &= k[\text{Var}_{\mathbf{N},\epsilon}(Y) + (E_{\mathbf{N},\epsilon}(Y) - t)^2] \end{aligned} \quad (2)$$

where $\mathbf{C}$ is a vector of design variables that can be partitioned into two mutually exclusive parts corresponding to "nonadjustment" factors, which do not affect the response mean, that is, $\mathbf{d} = (d_1, \dots, d_p)$, and "adjustment" factors, which can be used to adjust the response mean to the target, that is, $\mathbf{a} = (a_1, \dots, a_q)$. The goal is to find the factor settings in $\mathbf{d}$ and $\mathbf{a}$ that will minimize $R(\mathbf{C})$. Note that this partitioning of design variables is a theoretical ideal, and mutually exclusive groupings may not exist in practice. $\text{Var}_{\mathbf{N},\epsilon}(Y)$ and $E_{\mathbf{N},\epsilon}(Y)$ are the mean and variance of the process response over the random sources of variability represented by $\mathbf{N}$ and ϵ . $\mathbf{N}$ is a vector of uncontrollable "noise variables" (*see* **Factorial Experiments**) and ϵ represents purely random

variation. This decomposition of average quadratic loss allows us to frame RD as the minimization of the sum of process variance and the square of the deviation of process mean from the target.

Taguchi [7] introduced a family of performance measures called *signal-to-noise ratios* (SNRs) whose specific form depends on the desired response outcome. Instead of directly minimizing quadratic loss, Taguchi uses these SNR in his optimization procedure. Here we present three of the most common SNR for the static, or fixed target, problem that was illustrated previously with the Ina Tile example. The case where the response has a fixed nonzero target is called the *nominal-the-best* (NTB) case. Likewise the cases where the response has a unique smaller-the-better target or larger-the-better target are respectively called the *STB* and *LTB* cases. For these three cases Taguchi defined the SNR as follows:

$$SNR_i = \begin{cases} 10 \log_{10} \left[\frac{\bar{Y}_i^2}{S_{Y_i}^2} \right] & \text{for the NTB case} \\ -10 \log_{10} \left[\frac{1}{n} \sum_{j=1}^n Y_{ij}^2 \right] & \text{for the STB case} \\ -10 \log_{10} \left[\frac{1}{n} \sum_{j=1}^n \left(\frac{1}{Y_{ij}} \right)^2 \right] & \text{for the LTB case} \end{cases} \quad (3)$$

where SNR_i is the SNR for the i th combination of design factor settings in the experiment, $\bar{Y}_i$ and $S_{Y_i}^2$ are respectively the sample mean and variance of the response estimated from the n observations, that is, $Y_{i1}, \dots, Y_{in}$, at the i th design factor setting in the experiment.

To accomplish the objective of minimal expected squared-error loss for the NTB case, Taguchi proposes the following two-step optimization procedure:

1. Calculate and model the SNRs and find the values of the nonadjustment factor settings, that is, $\mathbf{d}$, which maximize the SNR.
2. Shift mean response to the target by changing the values of the adjustment factor(s) $\mathbf{a}$.

For the STB and LTB cases, Taguchi recommends directly searching for the values of the design vector

$\mathbf{C}$ which maximize the respective SNR. Alternatives to Taguchi's optimization approaches for these cases are proposed by Nair and Pregibon [8], León *et al.* [9], Box [10], and Tsui and Li [11].

Taguchi's Parameter Design for Multiple Targets

The term *dynamic robust design* (DRD), also known as a *signal-response system* (see **Signal-Response Systems**), signifies that there are multiple targets of interest for the process response. In static RD there is a single, fixed target value and the variables affecting process response are considered as either design or noise variables. In DRD there exists a third type of variable, the signal variable M , whose magnitude directly affects the mean value, and thus the target, of the output response.

Consider, for example, a push-pull cable actuator system which consists of a flexible release cable assembly, similar to one used for an automobile's remote mirror adjuster, or hood and trunk releases. This example was examined in McCaskey and Tsui [12] and Byrne and Quinlan [13]. The response of interest is output displacement measured in millimeters, which is affected by control and noise factors, as in the static case. In addition however, the cable system is affected by a user-influenced signal variable M , the input displacement, which affects the target value of the output response.

Note that for the static case of Ina Tile the targets of the tile dimensions were always fixed at specific values. A mirror adjuster moves the mirror to different positions based on the drivers' needs. These different mirror locations desired by the driver represent different target values to be achieved by adjustment of the control variables while maintaining low variability in response.

In order to optimize the design of a cable actuator, a study was planned in which 11 control factors (labeled $A-K$) were varied at two levels each using an OA12 orthogonal array (see **Orthogonal Arrays**). This 12-run, highly fractionated Taguchi design is comparable to the $N = 12$ Plackett-Burman design (see **Plackett-Burman Designs**), in which each two-factor interaction is partially aliased with nine main effects. This design is useful for efficiently screening a large number of variables when two-factor interactions are not anticipated.

The 11 control factors varied in the study included liner material modulus, braid wire diameter, liner wall thickness, braid density, coating material modulus, coating thickness, cable diameter, and cable material modulus. In addition to the OA12 inner array, two combinations of noise factors and three settings of the signal variable were incorporated into the outer array of the design. The signal variable settings corresponded to the specific mirror locations drivers were likely to use most. The noise factors were manufacturing characteristics that are very difficult to control, specifically number of cycles, assembly length, routing, and operating load. Three experimental replicates for each combination of control, noise, and signal settings were performed, resulting in 216 ($12 \times 2 \times 3 \times 3$) total observations. The experimental design and data are listed in Table 1. Note that the data in this presentation has been modified from the original experimental results found in McCaskey and Tsui [12] and Byrne and Quinlan [13] to more clearly illustrate Taguchi's optimization procedure for the DRD situation.

In analyzing the cable actuator data, we can begin

by specifying the loss function. A common choice of dynamic loss function is the quadratic loss function popularized by Taguchi,

$$L(Y, t(M)) = k(Y - t(M))^2 \quad (4)$$

where k is the constant described in equation (1) and $t(M)$ indicates that the target level is a function of the signal variable M . This loss function provides a good approximation to many realistic loss functions. It follows that the average loss becomes

$$\begin{aligned} R(C) &= kE_M E_{N,\epsilon}(Y - t(M))^2 \\ &= kE_M [\text{Var}_{N,\epsilon}(Y) + (E_{N,\epsilon}(Y) - t(M))^2] \end{aligned} \quad (5)$$

where E_M is the expected value over the range of the user-influenced signal variable M and N and ϵ are sources of variability contributed by the uncontrollable noise variables and pure random error.

The SNR for the cable actuator study must also incorporate the effect of the signal variable. Taguchi

Table 1 Design and data of the push–pull cable actuator experiment – simulated data

Run	Control factor ABCDEFGHIJK	Noise level	Signal factor levels								
			$M = 8$			$M = 16$			$M = 24$		
1	11111111111	N_1	4.03	4.28	4.37	8.41	8.56	9.09	12.75	12.92	12.75
		N_2	4.56	4.91	4.32	8.78	8.65	8.81	12.61	13.10	12.64
2	11111222222	N_1	5.32	5.90	5.70	11.92	11.70	11.65	17.14	17.59	17.24
		N_2	5.99	5.42	6.52	11.69	11.22	11.53	16.39	16.83	16.87
3	11222111222	N_1	4.14	4.03	3.61	7.32	6.71	7.74	11.41	10.70	10.14
		N_2	3.75	3.82	4.25	7.15	7.19	6.97	10.94	10.61	10.63
4	12122122112	N_1	4.76	5.53	5.41	9.68	10.78	10.47	16.19	16.13	14.66
		N_2	4.60	4.00	3.76	8.44	10.13	9.94	12.51	14.65	16.08
5	12212212121	N_1	4.60	6.23	5.65	11.08	10.09	11.39	17.42	16.14	17.56
		N_2	5.37	4.06	6.44	11.46	10.09	12.09	17.47	16.46	17.01
6	12221221211	N_1	6.00	4.32	4.54	8.43	8.66	8.68	12.50	11.46	12.35
		N_2	2.62	3.91	4.50	8.24	7.32	7.39	12.46	12.09	13.59
7	21221122121	N_1	4.59	5.17	4.73	9.78	10.91	10.58	15.58	14.48	14.84
		N_2	4.49	5.14	4.67	9.53	9.22	9.69	15.00	14.21	15.34
8	21212221112	N_1	3.92	3.75	3.79	8.09	8.42	8.18	12.06	12.25	12.01
		N_2	4.66	4.50	3.91	7.93	8.23	7.82	11.85	12.89	12.17
9	21122212211	N_1	5.76	5.33	4.65	10.38	10.58	9.66	15.77	14.87	15.35
		N_2	6.11	4.75	4.98	9.60	9.96	10.40	15.37	15.60	14.86
10	22211112212	N_1	5.18	6.02	5.64	11.35	11.59	10.36	17.72	17.01	16.42
		N_2	4.42	4.95	6.34	12.06	9.02	11.37	15.36	16.76	15.51
11	22121211122	N_1	3.30	3.83	3.90	7.85	8.23	6.94	12.01	13.34	12.38
		N_2	3.36	6.08	4.74	7.15	8.51	8.50	11.41	12.57	11.49
12	22112121221	N_1	2.65	3.46	3.22	8.73	6.74	7.05	11.10	10.95	11.77
		N_2	3.57	4.85	3.42	7.99	8.27	6.98	11.59	11.42	12.32

4 Signal-to-Noise Ratios for Robust Design

suggested the following SNR as a performance measure for the dynamic case:

$$SNR_i = 10 \log_{10} \frac{\hat{\beta}_i^2}{S_i^2} \quad (6)$$

where $\hat{\beta}_i$ is the regression coefficient from fitting the responses Y_{ij} to the signal M , that is,

$$\hat{Y}_{ij} = \hat{\beta}_i M \quad (7)$$

where β_i can be interpreted as the slope of the linear model between Y_{ij} and M with a zero intercept for the i th design factor setting and S_i^2 is the sample error variance of the i th fitted regression model. As in the static case, Taguchi recommends a two-step optimization procedure:

1. Calculate and model the SNR of equation (6) and find the nonadjustment factor settings, that is, $\mathbf{d}$, which maximize it.
2. Shift the mean slope $\hat{\beta}$ to the desired sensitivity level using the adjustment factor(s) $\mathbf{a}$.

Alternatives to Taguchi's performance measures and optimization procedures for DRD are proposed in Tsui [14] and Joseph and Wu [15, 16].

The upper half of Table 2 shows SNR responses for each level of the 11 control factors. From visual inspection of these results and consultation with the engineers designing the cable actuator system, control factors B and D , respectively corresponding to braid wire diameter and braid density, were chosen as the set of nonadjustment factors $\mathbf{d}$.

Next, a separate analysis was performed to determine the adjustment factors $\mathbf{a}$ that affect the signal-response slope. For this example the engineers desired a 1:1 relationship between the input

displacement signal and the output displacement response, corresponding to a slope of 1. Taguchi suggests calculating an SNR for the estimated slopes, $SNR_{\beta i}$:

$$SNR_{\beta i} = 10 \log_{10} \hat{\beta}_i^2 \quad (8)$$

The $SNR_{\beta i}$ responses for each level of the 11 control factors are shown in the lower half of Table 2. We see that factor H , that is, the cable material modulus, can be used to adjust the slope to its theoretically ideal value which in this case is 1. The design engineers verified that this was a reasonable choice to use as an adjustment variable.

Using the information in Table 2 to apply Taguchi's two-step optimization procedure we have the following conclusions:

1. Set the factors B and D to levels $B = 1$, $D = 1$ to maximize the SNR.
2. Assuming that it can be adjusted over a sufficiently wide range, use factor H to adjust the signal-response slope to its theoretically ideal value.

The factor settings determined from the two-step procedure will then, according to Taguchi, minimize the expected quadratic loss for the dynamic problem.

Strengths and Limitations of Taguchi's SNR

The concept of the SNR originated in electrical engineering to quantify a circuit's ability to discern a meaningful signal from the wide frequency spectrum of which it is a single component. The concept of comparing the strength of a signal to its companion noise is best demonstrated by the NTB case of

Table 2 SNR and SNR_{β} responses by control factor levels for the cable actuator experiment – simulated data

Factor	A	B	C	D	E	F	G	H	I	J	K
SNR											
Level 1	0.83	3.80	0.66	1.74	0.75	0.16	0.71	0.26	0.68	0.65	0.69
Level 2	0.06	−2.91	0.22	−0.85	0.14	0.72	−0.18	0.63	0.21	0.24	0.20
Difference	0.77	6.71	0.44	2.59	0.60	0.56	0.53	0.37	0.47	0.41	0.49
SNR_{β}											
Level 1	−4.70	−4.86	−4.79	−4.50	−4.57	−5.05	−4.75	−6.05	−4.76	−4.73	−4.79
Level 2	−4.94	−4.77	−4.85	−5.13	−5.06	−4.58	−4.88	−3.58	−4.87	−4.91	−4.85
Difference	0.24	0.09	0.06	0.63	0.50	0.47	0.13	2.47	0.11	0.18	0.06

equation (3) which is the logarithm of the ratio of the sample response mean (i.e., signal) squared to the sample response variance (i.e., noise). The SNR is the most controversial part of Taguchi's methods because it summarizes the quality information from the entire experiment with a single performance measure. Other methods examine separate measures of location and dispersion or transformations thereof.

In RD, the performance measure tells how a process is functioning. The advantages of the SNR advocated by Taguchi are twofold. Foremost, it has great appeal to an audience who are not statistically sophisticated as a single measure that encompasses all experimental information. Secondly, when combined with the simplicity of the two-step procedure, it has functioned as a powerful incentive for engineers and scientists to employ designed experiments to improve the quality of their products. Bridging the worlds of engineering and applied statistics is generally acknowledged as Taguchi's greatest contribution.

Conversely the statistical community has been troubled by the unexplained connection between maximizing the SNR and achieving the stated goal of minimal quadratic loss. Recall that the two-step procedure consisted of maximizing the SNR and then using adjustment factors which have negligible effect on the SNR to bring the process mean to target. Taguchi did not rigorously prescribe how this procedure guaranteed minimal quadratic loss.

The relationship between minimization of expected quadratic loss and Taguchi's SNR was articulated by León *et al.* [9]. They defined a function called the *performance measure independent of adjustment* (PerMIA; see **Performance Measures for Robust Design**) which justified the use of a two-step optimization procedure. They showed that Taguchi's SNR for the NTB case is a PerMIA when an adjustment factor exists and σ_Y^2 is proportional to μ_Y . When Taguchi's SNR complies with the properties of a PerMIA, his two-step procedure minimizes the expected value of squared-error loss.

León *et al.* [9] also emphasized two major advantages of Taguchi's two-step procedure:

- it reduces the dimension of the original optimization problem;
- it does not require reoptimization of the nonadjustment factors (**d**) for future changes of the target value. The adjustment factors (**a**) can be used to accommodate future changes in target.

Since the work of León *et al.* [9] the SNR has been extensively reviewed by the statistical community. Box [10] recommends modeling the process response Y with appropriate transformation as needed to find design factor settings that minimize **mean squared error**. He points out that even when σ_Y^2 is proportional to μ_Y , the analysis of $\log(Y)$ is more efficient than the SNR. A more recent comparison of the SNR with alternative performance measures is presented in Bérubé and Wu [17], and Wu and Hamada [18] provide a variety of approaches for related situations.

An important paper that compared use of Taguchi's SNR to process response was that of Nair and Pregibon [8] who advocate a three-phase process for RD consisting of exploration, modeling, and optimization. Their performance measure is process response which they assume is a function of the design and noise variables, that is, **C** and **N** referred to in equation (2). They reduce the RD problem to minimization of VAR_Y subject to the constraint that E_Y is held as close as possible to the target level t .

Other excellent discussions of the SNR and Taguchi's related parameter design techniques include Kackar [19], Nair [20], and Tsui [21]. In the last 10 years most research on RD has focused on DRD with seminal work including that of Miller and Wu [22], Lunani *et al.* [23], and McCaskey and Tsui [12]. More recent contributions include those of Joseph and Wu [15, 16].

References

- [1] Taguchi, G. & Wu, Y. (1980). *Introduction to Off-Line Quality Control*, Central Japan Quality Control Association, Nagoya.
- [2] Phadke, M.S., Kackar, R.N., Speeney, D.V. & Grieco, M.J. (1983). Introduction to quality engineering: designing quality into products and processes, *The Bell System Technical Journal* **1**(5), 1273–1309.
- [3] Shoemaker, A.C., Tsui, K.-L. & Wu, C.F.J. (1991). Economical experimentation methods for robust parameter design, *Technometrics* **33**, 415–428.
- [4] Tsui, K.-L. (1996). A multi-step analysis procedure for robust design, *Statistica Sinica* **6**, 631–648.
- [5] Tsui, K.-L. (1999a). Robust design optimization with multiple characteristics, *International Journal of Production Research* **37**, 433–445.
- [6] Phadke, M.S. (1989). *Quality Engineering Using Robust Design*, Prentice-Hall.
- [7] Taguchi, G. (1986). *Introduction to Quality Engineering: Designing Quality into Products and Processes*, Kraus International Publications, White Plains.

- [8] Nair, V.N. & Pregibon, D. (1986). A data analysis strategy for quality engineering experiments, *AT&T Technical Journal* **65**, 73–84.
- [9] León, R.V., Shoemaker, A.C. & Kackar, R.N. (1987). Performance measure independent of adjustment: an explanation and extension of Taguchi's signal to noise ratio, *Technometrics* **29**, 253–285.
- [10] Box, G.E.P. (1988). Signal to noise ratios, performance criteria and transformations, *Technometrics* **30**, 1–31.
- [11] Tsui, K.-L. & Li, A. (1994). Analysis of smaller and larger the better robust design experiments, *International Journal of Industrial Engineering* **1**, 193–202.
- [12] McCaskey, S.D. & Tsui, K.-L. (1997). Analysis of dynamic robust design experiments, *International Journal of Production Research* **35**, 1561–1574.
- [13] Byrne, D. & Quinlan, J. (1993). Robust function for attaining high reliability at low cost, in *IEEE Proceedings of the Annual Reliability and Maintainability Symposium*, Atlanta, GA, USA pp. 183–191.
- [14] Tsui, K.-L. (1999b). Response model analysis of dynamic robust design experiments, *IIE Transactions on Quality and Reliability* **31**, 1113–1122.
- [15] Joseph, V.R. & Wu, C.F.J. (2002a). Performance measures in dynamic parameter design, *Journal of Japanese Quality Engineering society* **10**, 82–86.
- [16] Joseph, V.R. & Wu, C.F.J. (2002b). Robust parameter design of multiple target systems, *Technometrics* **44**, 338–346.
- [17] Bérubé, J. & Wu, C.F.J. (1998). Signal-to-noise ratio and related measures parameter design optimization, Technical Report 321, University of Michigan.
- [18] Wu, C.F.J. & Hamada, M. (2000). *Experiments: Planning, Analysis and Parameter Design Optimization*, John Wiley & Sons, New York.
- [19] Kackar, R.N. (1985). Off-line quality control, parameter design, and the Taguchi method, *Journal of Quality Technology* **17**, 176–200.
- [20] Nair, V.N. (1992). Taguchi's parameter design: a panel discussion, *Technometrics* **34**, 127–161.
- [21] Tsui, K.-L. (1992). An overview of Taguchi method and newly developed statistical methods for robust design, *IIE Transactions on Quality and Reliability* **24**(5), 44–57.
- [22] Miller, A. & Wu, C.F.J. (1996). Parameter design for signal-response systems: a different look at Taguchi's dynamic parameter design, *Statistical Science* **11**, 122–136.
- [23] Lunani, M., Nair, V.N. & Wasserman, G.S. (1997). Graphical methods for robust design with dynamic characteristics, *Journal of Quality Technology* **29**, 327–338.

Related Articles

Orthogonal Arrays; Performance Measures for Robust Design; Plackett–Burman Designs; Robust Design; Signal–Response Systems.

TERRENCE E. MURPHY, KWOK-LEUNG TSUI
AND MARY MCSHANE-VAUGHN

Tolerance Design

Introduction

The three main stages in a product/process design are system design, parameter design, and tolerance design [1]. In system design the basic architecture of the system is selected; in parameter design the nominal values of the parameters in the system are selected; and in tolerance design the allowable variations of the parameters around their nominal values are selected. This chapter discusses tolerance design. Tightening the tolerances increases the manufacturing cost, whereas widening the tolerances decreases the quality. Thus designing tolerances is about balancing cost with quality.

The quality of a product is measured with respect to customer requirements. Thus tolerance design activities begin with acquiring customer specifications and translating them to product specifications. A product may be an assembly of several components and each component may have several parameters. The objective in product design is to find the specifications of each parameter so that the desired product quality can be achieved. Each component itself may be manufactured in several stages of a manufacturing process and each stage may contain several parameters. Therefore the objective in process design is to find the specifications of each process parameter so that the component specifications can be achieved. In this chapter we discuss only the designing component and the process parameter tolerances and not the conversion of customer tolerances into product tolerances.

Let y be the output quality characteristic and $x_1, x_2, \dots, x_p$ be the input parameters (or factors). In reality, the input parameters are not constants and have variations around their nominal values. For example, product/component parameters vary because of the manufacturing variations; similarly manufacturing process parameters vary because they cannot be controlled precisely over time. Assume that (maybe after a suitable transformation) x_i follows a **normal distribution** with mean μ_i and variance σ_i^2 . We can write $x_i = \mu_i + \delta_i$, where $\delta_i \sim \mathcal{N}(0, \sigma_i^2)$. For some parameters, the designer can choose a value for μ_i , in which case δ_i is called an *internal noise* factor. If the designer has no control over the nominal value, then x_i is called an *external noise* factor.

Let the input–output relationship be $y = f(x_1, \dots, x_p)$. Because the x 's vary, y also varies. The variations in y can be large and not acceptable. One approach to reduce the variations in y is to reduce the variations in the x 's. This can be done by tightening the tolerances, but it inadvertently increases the manufacturing cost. A more cost-effective approach is to use robust parameter design (*see Robust Design*). Here, the nominal values of the x 's are chosen such that the variations of y are as small as possible. Note that in robust parameter design, no effort is made to reduce the variability of the x 's, thus it does not increase the cost of manufacturing. However, even after the application of robust parameter design, the variations can still be large and tightening the tolerances becomes inevitable. But this clearly shows that tolerance design should be performed only after parameter design, not the other way. We will explain this using a simple example.

Example

Consider an example discussed in [2] on the design of a cyclone (see also [3]). A cyclone is used for separating solid mass and gaseous mass. The characteristic of output quality, the critical parameter of particles y , is related to seven input parameters by

$$y = 174.42 \left(\frac{x_1}{x_5} \right) \left(\frac{x_3}{x_2 - x_1} \right)^{0.85} \times \left(\frac{1 - 2.62[1 - 0.36(x_4/x_2)^{-0.56}]^{3/2}(x_4/x_2)^{1.16}}{x_6 x_7} \right)^{1/2} \quad (1)$$

The existing nominal values of the factors and their tolerances are given in Table 1. Here the tolerances are expressed as a percentage of the nominal values. For example, the specification of x_6 is 16 ± 4 . The objective is to find the optimal nominal values and tolerances such that the quality is maximized while minimizing the manufacturing cost.

The quality can be measured using a quality loss function [1] (see also [4]). In this example, the response is a nominal-the-best type characteristic and thus the following quadratic loss function can be used

$$L(y) = k(y - \tau)^2 \quad (2)$$

where τ is the target for the characteristic and k is a cost coefficient. Here $\tau = 1.5\text{ m}$ and $k =$

2 Tolerance Design

Table 1 Existing values of the factors and tolerances

Factors	x_1	x_2	x_3	x_4	x_5	x_6	x_7
Nominal values	0.10 m	0.30 m	0.10 m	0.10 m	1.50	16.00 m s ⁻¹	0.75 m
Tolerances (%)	25	25	25	25	25	25	25

$1/9 \times 10^5$ yen m⁻². The determination of k is not trivial, see [2, 3] for details. Although an accurate **estimation** of k is not important for parameter design, it is important for the tolerance design. The average quality loss can be obtained as

$$Q = E[L(Y)] = kE[f(\mu_1 + \delta_1, \dots, \mu_p + \delta_p) - \tau]^2 \quad (3)$$

where the expectation is taken with respect to the distribution of the noise factors.

To evaluate the quality loss, we need to know the distribution of the noise factors. Because no information is available about σ_i in the cyclone design example, we will indirectly estimate it by assuming that the tolerance is equal to three times the standard deviation. Let t_i denote the tolerance of factor i . Thus, assume that $\sigma_i = t_i/3$. In general, it is difficult to find a closed-form expression for Q . One simple method to estimate Q is to use a **Monte Carlo** approximation. Generate n independent samples for each δ_i from $\mathcal{N}(0, \sigma_i^2)$. Then

$$\hat{Q} = \frac{k}{n} \sum_{j=1}^n [f(\mu_1 + \delta_{1j}, \dots, \mu_p + \delta_{pj}) - \tau]^2 \quad (4)$$

Using $n = 10\,000$, we obtain $Q = ¥2096$ per product for the existing design.

In the robust parameter design, the nominal values of the control factors are chosen to minimize the sensitivity of the response to the noise factors. This can be achieved by minimizing the average quality loss with respect to $\mu_1, \dots, \mu_p$. Because there is no explicit expression for Q in terms of μ_i 's, this cannot be easily done. However, $\hat{Q}$ can be easily minimized with respect to μ_i 's. It is known that when n is large, the solution obtained by minimizing $\hat{Q}$ is close to the solution obtained by minimizing Q . This method, known as *sample average approximation*, is widely used in stochastic programming [5].

Let $\mu = (\mu_1, \dots, \mu_p)'$. Assume that the nominal values can be changed by at most 25% of the existing values given in Table 1. Minimizing $\hat{Q}$ subject to the

above constraints on μ_i , we obtain

$$\mu^* = (0.075, 0.375, 0.125, 0.125, 1.319, 14.183, 0.655)' \quad (5)$$

The average quality loss at the above setting is ¥756 per product, which is almost one-third of that of the existing design. Suppose that this is still large and we want to reduce it further; We can now consider tolerance design.

Tolerance Design

Taguchi's Approach

The quality loss can be reduced by tightening the tolerances. Tighter tolerances can be achieved by using better-quality raw materials, or by instituting better control devices, or by rejecting/reworking products falling outside the specifications. All of these increase the manufacturing cost. Therefore the decision of tightening tolerances should be made carefully so that the desired quality improvement can be achieved without unduly increasing the cost.

The manufacturing costs per product for three classes of tolerances are given in Table 2. We need to decide which tolerances should be tightened? By looking at Table 2, we may think that the tolerances of x_1 , x_6 , and x_7 should be tightened, because these

Table 2 Tolerance costs and model parameters

Factors	Costs			Model parameters		
	$\pm 25\%$	$\pm 12.5\%$	$\pm 2.5\%$	A	B	γ
x_1	10	25	100	-17.09	11.22	0.6357
x_2	20	50	200	-34.18	22.45	0.6357
x_3	20	50	200	-34.18	22.45	0.6357
x_4	50	100	500	0	12.5	1
x_5	50	200	1000	-196.06	94.97	0.6867
x_6	10	25	100	-17.09	11.22	0.6357
x_7	10	25	100	-17.09	11.22	0.6357

are least expensive. But if they do not have much effect on the response, then we may not get the required improvement in the quality. Thus, to take a judicious decision on which tolerances should be tightened, we need to know the contribution of each factor to quality. Taguchi [1] proposed to find this using **orthogonal array** experiments.

An 18-run orthogonal array [1] is given in Table 3, which can accommodate the seven factors. The three levels of each factor x_i can be chosen as a representative sample from $\mathcal{N}(\mu_i^*, \sigma_i^2)$. If we take the three levels as

$$\mu_i^* - \sqrt{\frac{3}{2}}\sigma_i, \mu_i^*, \text{ and } \mu_i^* + \sqrt{\frac{3}{2}}\sigma_i \quad (6)$$

Table 3 18-run orthogonal array and data

Run	x_1	x_2	x_3	x_4	x_5	x_6	x_7	y
1	1	1	1	1	1	1	1	1.60
2	1	2	2	2	2	2	2	1.26
3	1	3	3	3	3	3	3	1.02
4	2	1	1	2	2	3	3	1.29
5	2	2	2	3	3	1	1	1.41
6	2	3	3	1	1	2	2	1.63
7	3	1	2	1	3	2	3	1.57
8	3	2	3	2	1	3	1	1.96
9	3	3	1	3	2	1	2	1.40
10	1	1	3	3	2	2	1	1.52
11	1	2	1	1	3	3	2	1.02
12	1	3	2	2	1	1	3	1.30
13	2	1	2	3	1	3	2	1.61
14	2	2	3	1	2	1	3	1.60
15	2	3	1	2	3	2	1	1.15
16	3	1	3	2	3	1	2	1.84
17	3	2	1	3	1	2	3	1.52
18	3	3	2	1	2	3	1	1.52

Table 4 Analysis of Variance

Factors	df	SS	Contribution%
x_1	2	0.369	33.9
x_2	2	0.165	15.2
x_3	2	0.210	19.3
x_4	2	0.018	1.6
x_5	2	0.220	20.2
x_6	2	0.046	4.3
x_7	2	0.060	5.5
Error	3	0.001	0.0
Total	17	1.089	100

then the mean and variance of these three values will be exactly μ_i^* and σ_i^2 . The function in equation (1) is evaluated for the 18 runs and is given in Table 3. The analysis of variance is shown in Table 4. The contribution of each factor is calculated as the ratio of its **sum of squares** to the total sum of squares. We can see that x_1 is the most important factor followed by x_5 , x_3 , and x_2 . Because tightening the tolerance of x_5 is very costly, it makes sense to focus on the tolerances of x_1 , x_2 , and x_3 . The decision about how much to tighten can be determined by comparing the reduction in quality loss with the increase in the manufacturing costs. This requires additional calculations, which will not be presented here. Instead, we present in the next section a more general and flexible method based on Monte Carlo simulations. However, we should point out that Taguchi's approach is much superior to the traditional tolerancing methods, such as worst-case and root sum-of-squares methods. See [1, 6] for details.

Monte Carlo Approach

Let $\delta_i = \sigma_i e_i$. Thus, we have $x_i = \mu_i + \sigma_i e_i$, where $e_i \sim \mathcal{N}(0, 1)$. Again assume that $t_i = 3\sigma_i$. We should explain the implications of this assumption. Tolerance represents the allowable variation not the actual variation. The actual variation in x_i could be more than t_i , in which case we assume that additional effort will be put into the process for reducing the variation so that $3\sigma_i = t_i$. On the other hand, if the actual variation is less than t_i , then by assuming $t_i = 3\sigma_i$ we underestimate the quality. Thus under this assumption, we are considering the worst-case scenario in terms of quality.

As in the case of the cyclone design example, sometimes it is more appropriate to express the tolerances as a percentage of their nominal values. Define $a_i = t_i/\mu_i$. Then $\sigma_i = \mu_i a_i/3$. Thus the average quality loss after parameter design can be computed as

$$\begin{aligned} Q &= E[L(Y)] = E[f(\mu_1^* + \sigma_1 e_1, \dots, \mu_p^* + \sigma_p e_p) - \tau]^2 \\ &= E\left[f\left(\mu_1^* + \frac{\mu_1^* a_1}{3} e_1, \dots, \mu_p^* + \frac{\mu_p^* a_p}{3} e_p\right) - \tau\right]^2 \end{aligned} \quad (7)$$

We also need to obtain an expression for the cost of changing tolerances. Let $c_i(a_i)$ be the cost of

4 Tolerance Design

manufacturing a product with tolerance for the i th factor equal to $t_i = a_i \mu_i$. Thus the total cost of manufacturing is

$$C = \sum_{i=1}^p c_i(a_i) \quad (8)$$

For the cyclone design the costs for each factor are given for three different tolerances. It is more convenient to use a continuous function for the costs. Many cost–tolerance models are given in [7]. We use the following model

$$c_i(a_i) = A_i + \frac{B_i}{a_i^{\gamma_i}} \quad (9)$$

The three parameters A_i , B_i , and γ_i can be deduced from the cost data in Table 2 and are also reported in the same table. The actual estimation of the tolerance costs can be tedious; see [8] for another example.

In [2, 3] the cost of the existing design is adjusted to zero. To be consistent with their approach, we reparameterize the cost function as

$$c_i(a_i) = \bar{A}_i + B_i \left(\frac{1}{a_i^{\gamma_i}} - \frac{1}{\bar{a}_i^{\gamma_i}} \right) \quad (10)$$

where $\bar{a}_i$ denotes the tolerance for the existing design and $\bar{A}_i$ the corresponding cost. We then let $\bar{A}_i = 0$.

Our objective is to find the optimal tolerances $a_1, \dots, a_p$ that minimize $Q + C$. For optimization we use the Monte Carlo approach described before. Thus the problem is to

$$\begin{aligned} & \min_{a_1, \dots, a_p} \frac{k}{n} \\ & \times \sum_{j=1}^n \left[f \left(\mu_1^* + \frac{\mu_1^* a_1}{3} e_{1j}, \dots, \mu_p^* + \frac{\mu_p^* a_p}{3} e_{pj} \right) - \tau \right]^2 \\ & + \sum_{i=1}^p c_i(a_i) \end{aligned}$$

The optimal solution is given by $\mathbf{a}^* = (0.073, 0.125, 0.122, 0.25, 0.197, 0.138, 0.145)'$. As expected, the largest decrease in the tolerance is for x_1 followed by x_2 and x_3 . No change is made for the tolerance of x_4 because its contribution to quality is negligibly small. At the optimal setting, the quality loss is obtained as ¥251 per product, which is one-third of the quality loss obtained after parameter design. However, the

cost of tolerancing is increased by ¥160 per product. The total cost of ¥411 per product is still much lower than the total cost after parameter design.

We note that the optimal solution obtained here is much better than that given in [2, 3]. This is because we have used a continuous cost–tolerance model, which gives more flexibility than using some discrete choices for the tolerances.

An important assumption in carrying out the Monte Carlo approach is that the function $f(x_1, \dots, x_p)$ is easy to evaluate. Although this is true for the cyclone design example, it need not be the case always. For example, the function could be implicitly defined through a finite element model, which can be quite time consuming to evaluate. In such cases, the function should be approximated by a metamodel using techniques, such as kriging (see **Computer Experiments**). Because the metamodel is easy to evaluate, the Monte Carlo approach can be efficiently implemented. When analytic or finite element models are not available, physical experiments should be carried out to estimate the function. In physical experiments, it is difficult to incorporate all of the external noise factors, and thus an error term should be included in the model: $y = f(x_1, \dots, x_p) + \epsilon$, where $\epsilon \sim \mathcal{N}(0, \sigma^2)$. Note that σ^2 can also depend on $x_1, \dots, x_p$. This will change the evaluation of the average quality loss, but everything else in the procedure remains the same.

The estimation of Q can be improved by using better sampling techniques, such as **Latin hypercube** sampling. Even better space filling designs are now available (see **Computer Experiments; Latin Hypercube Designs**).

Integrated Parameter and Tolerance Design

The parameter design solution was obtained by fixing the tolerances at the existing values. The optimal tolerances are then obtained by fixing the nominal values at the parameter design solution. Clearly, the parameter design solution that we obtained is not optimal anymore because the tolerances are different. We can go back and perform the parameter design optimization again with the new tolerances. This iterative process can be continued until convergence [9]. A more straightforward approach is to optimize the total cost simultaneously with respect to both nominal

Table 5 Comparison of different designs

Design	Nominal and tolerance values							Mean	Variance	Q	C	$Q + C$
Existing	0.10 0.25	0.30 0.25	0.10 0.25	0.10 0.25	1.50 0.25	16 0.25	0.75 0.25	1.76	0.120	2096	0	2096
Parameter	0.075 0.25	0.375 0.25	0.125 0.25	0.125 0.25	1.319 0.25	14.2 0.25	0.655 0.25	1.45	0.066	756	0	756
Tolerance	0.075 0.073	0.375 0.125	0.125 0.122	0.125 0.25	1.319 0.197	14.2 0.138	0.655 0.145	1.44	0.019	251	160	411
Integrated	0.075 0.071	0.375 0.120	0.125 0.120	0.125 0.25	1.176 0.189	16 0.131	0.685 0.137	1.49	0.019	209	177	386

and tolerance values [3]. Thus the problem is to

$$\min_{\mu_1, \dots, \mu_p, a_1, \dots, a_p} \frac{k}{n} + \sum_{j=1}^n \left[f\left(\mu_1 + \frac{\mu_1 a_1}{3} e_{1j}, \dots, \mu_p + \frac{\mu_p a_p}{3} e_{pj}\right) - \tau \right]^2 + \sum_{i=1}^p c_i(a_i) \quad (11)$$

We note that some combinations of μ_i 's and a_i 's can generate infeasible solutions. For example, in the cyclone design x_2/x_4 should lie within [1.2, 6.2], otherwise negative values will be produced inside the square root leading to infeasible solutions. Therefore special care should be taken during simulations to avoid such values.

The optimal solution is obtained as $\mu^* = (0.075, 0.375, 0.125, 0.125, 1.176, 16, 0.685)'$ and $a^* = (0.071, 0.120, 0.120, 0.25, 0.189, 0.131, 0.137)'$. The total cost is ¥386 per product, which is lower than what we obtained after the two-stage parameter and tolerance designs. Thus, as expected, the integrated approach performs better than the two-stage approach. For comparison, the results are summarized in Table 5.

Acknowledgments

The research was supported by the U.S. National Science Foundation grant CMMI-0620259.

References

- [1] Taguchi, G. (1986). *Introduction to Quality Engineering*, Unipub Quality Resources, White Plains.
- [2] Mori, T. (1985). *Case Studies in Experimental Design*, Management System Research Laboratory, Tokyo (in Japanese).
- [3] Li, W. & Wu, C.F.J. (1999). An integrated method of parameter design and tolerance design, *Quality Engineering* **11**, 417–425.
- [4] Joseph, V.R. (2004). Quality loss functions for nonnegative variables and their applications, *Journal of Quality Technology* **36**, 129–138.
- [5] Ruszczynski, A. & Shapiro, A. (2003). *Stochastic Programming, Handbook in Operations Research and Management Science*, Elsevier.
- [6] Creveling, C.M. (1997). *Tolerance Design: A Handbook for Developing Optimal Specifications*, Addison-Wesley, Reading.
- [7] Chase, K.W., Greenwood, W.H., Loosli, B.G. & Hauglund, L.F. (1990). Least cost tolerance allocation for mechanical assemblies with automated process selection, *Manufacturing Review* **3**, 49–59.
- [8] Chase, K.W. (1999). Tolerance allocation methods for designers, ADCATS Report NO. 99-6, Department of Mechanical Engineering, Brigham Young University, Provo, Utah.
- [9] Chan, L.K. & Xiao, P.H. (1995). Combined robust design, *Quality Engineering* **8**, 47–56.

Related Articles

Analysis of Variance; Computer Experiments; Latin Hypercube Designs; Orthogonal Arrays Robust Design.

V. ROSHAN JOSEPH

Bergman–Hynén Method for Dispersion Effects

Introduction

During the last decades, **robust design** experimentation has become an important part of industrial statistics, see, for example, Wu and Hamada [1], Myers and Montgomery [2], Robinson *et al.* [3] and Myers *et al.* [4]. Dispersion analysis is of specific interest in robust design experimentation, where the objective is to find design solutions that make products with minimal dispersion (*see* **Dispersion Effects**).

Finding **dispersion effects** from unreplicated two-level **fractional factorial** designs has attracted much interest since the seminal paper by Box and Meyer [5]. The Box–Meyer method is based on a test statistic containing the ratio of the residual **sum of squares** connected with the high and low levels of each factor (*see* **Box–Meyer Method for Dispersion Effects**). The drawback of this method is that the sums of squares in the numerator and denominator are correlated as they are calculated from the same location model. Consequently, Box–Meyer’s method does not contain any formal statistical test to identify active dispersion effects.

Recently, a large number of papers on the identification of dispersion effects have given interesting generalizations, complements, and alternatives to the methods suggested by Box and Meyer [5]; see, for example, Engel and Huele [6], Bergman and Hynén [7], Nelder and Lee [8], Blomkvist *et al.* [9], Holm and Wiklander [10], Pan [11], Arvidsson *et al.* [12], Brenneman and Nair [13], McGrath and Lin [14, 15], McGrath [16], and Arvidsson *et al.* [17]. The purpose of this paper is to elaborate on the usefulness and limitations of the method for dispersion analysis suggested in Bergman and Hynén [7]. The next section presents the model and notation used throughout. The section titled “Identification of Dispersion Effects” presents two different ways of representing the Bergman–Hynén method, whereas aspects of modelling are discussed in the section following it. The discussion in the last section elaborates on the usefulness of the method.

Model and Notation

Let $\mathbf{y}$ denote an $n \times 1$ response vector from a 2^{k-p} fractional factorial experiment and $\mathbf{X} = [\mathbf{x}_0 | \dots | \mathbf{x}_{n-1}]$ the corresponding saturated $n \times n$ design matrix. Here, $\mathbf{x}_0 = \mathbf{1}$ is an intercept column of +1’s, and $\mathbf{x}_1, \dots, \mathbf{x}_{n-1}$ are orthogonal columns of +1’s and –1’s corresponding to the high and low levels of the factors and interactions. We assume a linear location model $\mathcal{L}$ with $k_{\mathcal{L}}$ parameters:

$$\mathbf{y} = \mathbf{X}_{\mathcal{L}}\boldsymbol{\beta} + \boldsymbol{\varepsilon} \quad (1)$$

where $\mathbf{X}_{\mathcal{L}}$ is an $n \times k_{\mathcal{L}}$ orthogonal design matrix containing the intercept column together with $k_{\mathcal{L}} - 1$ columns in $\mathbf{X}$. Here $\boldsymbol{\beta}$ denotes a vector containing the $k_{\mathcal{L}}$ regression coefficients and $\boldsymbol{\varepsilon}^T = (\varepsilon_1, \varepsilon_2, \dots, \varepsilon_n)$ an error vector of independent and normally distributed deviations with means zero. In addition to the location model we assume a dispersion model such that the variances of $\boldsymbol{\varepsilon}$ are $\sigma_j^2 = \mathbf{Z}_j\boldsymbol{\gamma}$, $j = 1, \dots, n$. Here, $\mathbf{Z}$ is an $n \times k_{\mathcal{D}}$ orthogonal matrix containing an intercept column together with $k_{\mathcal{D}} - 1$ columns in $\mathbf{X}$ related to factors that influence the dispersion. Further, $\mathbf{Z}_j$ are the corresponding row vectors, and $\boldsymbol{\gamma}$ is a $k_{\mathcal{D}} \times 1$ vector including dispersion coefficients. If **estimation** of the parameters in $\boldsymbol{\gamma}$ is focused, it is necessary to add the constraint $\sum_{i=1}^{k_{\mathcal{D}}-1} |\gamma_i| < \gamma_0$ on the dispersion coefficients.

Identification of Dispersion Effects

The Bergman–Hynén method for dispersion effect identification can be represented in a number of different ways. In this section, the residual-based representation is reviewed to point out similarities and differences with the Box–Meyer method. Further, a representation based on half contrasts is reviewed.

The Extended Location Model Representation

Let $\mathcal{L}_i^E$ denote an expanded location model that in addition to all factors and interactions in $\mathcal{L}$ contains all interactions between factor i and the factors and interactions in $\mathcal{L}$. Further, let $\tilde{r}_i$ denote the **residuals** corresponding to this expanded location model. The test statistic of the Bergman–Hynén method can then be written as follows:

$$D_i^{\text{BH}} = \frac{\sum \tilde{r}_{i+}^2}{\sum \tilde{r}_{i-}^2} \quad (2)$$

where $\tilde{r}_{i+}$ and $\tilde{r}_{i-}$ are the residuals of the expanded location model associated with the high and low levels of i , respectively. The test statistic of factor i can be calculated if $n > 2(k_L - k_i)$, where k_i is the number of pairs in $\mathcal{L}$ whose interaction is i (see, e.g., [13]). The main difference between D_i^{BH} and the Box–Meyer test statistic is that an expanded location model is used for calculating the residuals of factor i . The advantage of this innovation is that the residuals associated with the high and low levels of i are uncorrelated provided that no dispersion effects exist. This is an important property as it increases the test efficiency.

The Half-contrast Representation

Let $\mathbf{X}_{\mathcal{L}_i^E}$ be the design matrix corresponding to $\mathcal{L}_i^E$ and let $\mathbf{X}_{\tilde{\mathcal{L}}_i^E}$ contain the remaining columns of the saturated design matrix $\mathbf{X}$. Further, let $\mathbf{x}_i$ denote the column in $\mathbf{X}$ corresponding to factor i and $\text{diag}(\mathbf{x}_i)$ the corresponding diagonal matrix. Then

$$\begin{aligned} \mathbf{P}_{i+} &= \frac{1}{2}(\mathbf{I}_n + \text{diag}(\mathbf{x}_i)) \text{ and} \\ \mathbf{P}_{i-} &= \frac{1}{2}(\mathbf{I}_n - \text{diag}(\mathbf{x}_i)) \end{aligned} \quad (3)$$

are diagonal matrices partitioning $\mathbf{X}$ into two matrices $\mathbf{P}_{i+}\mathbf{X}$ and $\mathbf{P}_{i-}\mathbf{X}$ corresponding to high and low levels of i , respectively.

Now, let the m th column in $\frac{1}{\sqrt{n}}\mathbf{P}_{i+}\mathbf{X}_{\tilde{\mathcal{L}}_i^E}$ be denoted $\mathbf{V}_{i+,m}$ and define the contrast $v_{i+,m} = \mathbf{V}_{i+,m}^T \mathbf{y}$; further, define $v_{i-,m}$ similarly. We now have

$$D_i^{\text{BH}} = \frac{\sum_{m \in \tilde{\mathcal{L}}} v_{i+,m}^2}{\sum_{m \in \tilde{\mathcal{L}}} v_{i-,m}^2} \quad (4)$$

As shown in Bergman and Hynén [7], $v_{i+,m}$ and $v_{i-,m}$ correspond to $n/2 - (k_L - k_i)$ pairs of independent contrasts not involved in the location model. Thus, their expected values are zero and their variances are $V(v_{i+,m}) = \sigma_{i+}^2/2$ and $V(v_{i-,m}) = \sigma_{i-}^2/2$, respectively; note that $\sigma_{i+}^2 = 2n^{-1} \sum_{x_{j,i}=+1} \sigma_j^2$, where $\mathbf{x}_{j,i} = +1$ indicates summation over the high levels of $\mathbf{x}_i$, and σ_{i-}^2 is defined correspondingly. Thus, under the assumption that experiments contain no dispersion effects $D_i^{\text{BH}} \sim F(n/2 - (k_L - k_i), n/2 - (k_L - k_i))$ and formal **hypothesis testing** can be used to investigate whether i has a dispersion effect or not.

The equivalence of the two representations of Bergman–Hynén’s method follows immediately from the result within linear models that the sum of squared residuals is equal to the sum of squares of contrasts not included in the model, see Rao [18].

Modeling Aspects

A common practice is to assume log-linear dispersion models. In these models, the expected value is described as a linear function of the factor levels and the variance is modeled as a log-linear function of the factor levels. This is a convenient model because it makes estimated variances necessarily positive. However, this argument is not very strong when estimates are derived merely to provide support for the identification of factors affecting dispersion. In these cases, differences in the magnitude of the estimates are a more important aspect than the actual values of the estimates. Furthermore, the physical interpretation of the log-linear model is not clear. Arvidsson *et al.* [17] argue that it is unnatural to use different link functions for the expected value and the variance in the robust design framework. The rationale of this argument is that dispersion effects are assumed to be interactions between design factors and noise factors not systematically varied in the experiment. Therefore, the application of log-linear dispersion models might be questioned in robust design experimentation.

It should be noted that multiplicative models could be assumed in many engineering situations. However, these models are easily transformed into additive models by taking logarithms of the raw data before fitting any models. Of course, additional types of scale transformations might also be necessary in order to achieve an additive model. The point is that we want to analyze the overdispersion in the model, that is, dispersion not attributable to a scale transformation. (For related discussions, see León *et al.* [19], Box [20], and **Performance Measures for Robust Design**.) In summary, our proposal for modeling robust design experiments, possibly after a scale transformation, is a linear location model $\mathcal{L}$. This location model is complemented with a linear variance model interpreted as interactions between design factors and unknown or nonobserved noise factors.

Brenneman and Nair [13] showed that the appropriateness of different methods for dispersion effect

identification depends on the assumed variance model. McGrath and Lin [14] and Brenneman and Nair [13] showed that the Bergman–Hynén test statistic is biased under the log-linear variance model if multiple dispersion effects exist. However, this is not the case when an additive variance model is assumed.

Discussion

The assumption of a correctly identified location model is a common denominator in all methods for identification of dispersion effects in unreplicated fractional factorials. If this assumption is inaccurate, **aliasing** between location and dispersion effects may complicate the identification. In fact, McGrath and Lin [15] showed that two unidentified location effects create a spurious dispersion effect in their interaction column. McGrath [16] proposed two procedures for determining whether a dispersion effect is real or whether it actually stems from two unidentified location effects. Moreover, Bisgaard and Pinho [21] discussed possible follow-up strategies to investigate whether identified dispersion effects are actually due to aberrant observations or faulty execution of some of the experimental runs.

One aspect of the identification procedure suggested by Bergman and Hynén [7], not addressed in this paper, is the situation that appears when multiple dispersion effects exist. In those cases it is easily seen that the half contrasts in the numerator and denominator of the test statistic are correlated. Thus, the test statistics will not be F distributed. One way to address this problem is to adopt the idea of McGrath and Lin [14] to partition the experiment with respect to the levels of all factors expected to influence dispersion. This extended partitioning of the experiment is done to obtain uncorrelated test statistics, thus making a subsequent testing or normal plotting strategy more adequate. A problem that may occur when pursuing this strategy is that there will be few **degrees of freedom** for dispersion effect testing. However, a continuation of these ideas must be postponed to future research.

Acknowledgments

The authors would like to thank the Swedish Foundation for Strategic Research/Gothenburg Mathematical Modelling Centre (GMMC) for financial support.

References

- [1] Wu, C.F.J. & Hamada, M. (2000). *Experiments – Planning, Analysis and Parameter Design Optimization*, John Wiley & Sons, New York.
- [2] Myers, R.H. & Montgomery, D.C. (2002). *Response Surface Methodology – Process and Product Optimization Using Designed Experiments*, John Wiley & Sons, New York.
- [3] Robinson, T.J., Borror, C.M. & Myers, R.H. (2003). Robust parameter design: a review, *Quality and Reliability Engineering International* **20**(1), 81–101.
- [4] Myers, R.H., Montgomery, D.C., Vining, G.G., Borror, C.M. & Kowalski, S.M. (2004). Response surface methodology: a retrospective and literature survey, *Journal of Quality Technology* **36**(1), 53–77.
- [5] Box, G.E.P. & Meyer, R.D. (1986). Dispersion effects from fractional designs, *Technometrics* **28**(1), 19–27.
- [6] Engel, J. & Huele, A.F. (1996). A generalized linear modeling approach to robust design, *Technometrics* **38**(4), 365–373.
- [7] Bergman, B. & Hynén, A. (1997). Dispersion effects from unreplicated designs in the 2^{k-p} series, *Technometrics* **39**(2), 191–198.
- [8] Nelder, J.A. & Lee, Y. (1998). Joint modeling of mean and dispersion, *Technometrics* **40**, 168–175.
- [9] Blomkvist, O., Hynén, A. & Bergman, B. (1997). A method to identify dispersion effects from unreplicated multilevel experiments, *Quality and Reliability Engineering International* **13**, 127–138.
- [10] Holm, S. & Wiklander, K. (1999). Simultaneous estimation of location and dispersion in two-level fractional factorial designs, *Journal of Applied Statistics* **26**(2), 235–242.
- [11] Pan, G. (1999). The impact of unidentified location effects on dispersion-effects identification from unreplicated factorial designs, *Technometrics* **41**(4), 313–326.
- [12] Arvidsson, M., Kammerlind, P., Hynén, A. & Bergman, B. (2001). Identification of factors influencing dispersion in split-plot experiments, *Journal of Applied Statistics* **28**(3), 269–283.
- [13] Brenneman, W.A. & Nair, V.N. (2001). Methods for identifying dispersion effects in unreplicated factorial experiments: a critical analysis and proposed strategies, *Technometrics* **43**(4), 388–405.
- [14] McGrath, R.N. & Lin, D.K.J. (2001). Testing multiple dispersion effects in unreplicated fractional factorial designs, *Technometrics* **43**(4), 406–414.
- [15] McGrath, R.N. & Lin, D.K.J. (2001). Confounding of location and dispersion effects in unreplicated fractional factorials, *Journal of Quality Technology* **33**(2), 129–139.
- [16] McGrath, R.N. (2003). Separating location and dispersion effects in unreplicated fractional factorials, *Journal of Quality Technology* **35**(3), 306–316.

4 Bergman–Hynén Method for Dispersion Effects

- [17] Arvidsson, M., Gremyr, I. & Bergman, B. (2006). Interpretation of dispersion effects in a robust design context, *Journal of Applied Statistics* **33**(6), 623–627.
- [18] Rao, C.R. (1973). *Linear Statistical Inference and its Applications*, John Wiley & Sons, New York.
- [19] León, R.V., Shoemaker, A.C. & Kacker, R.N. (1987). Performance measures independent of adjustment, **29**(3), *Technometrics* 253–285.
- [20] Box, G.E.P. (1988). Signal-to-noise ratios, performance criteria, and transformations, *Technometrics* **30**(1), 1–40.
- [21] Bisgaard, S. & Pinho, A. (2003). Follow-up experiments to verify dispersion effects: Taguchi’s welding experiment, *Quality Engineering* **16**(2), 335–343.

MARTIN ARVIDSSON AND BO BERGMAN

Box–Meyer Method for Dispersion Effects

Introduction

During the process of new product development, we run engineering experiments to optimize both product design and manufacturing processes. An example of developing the tablet for a new drug product illustrates the concept. Using **factorial** and **response surface** designs, we vary the levels of inert ingredients in the formulation, as well as process variables such as roller compaction speed and pressure. The results of the experiments allow us to set the values of these variables in the product/process design, in order to optimize features such as product shelf life. Typical designs employed in earlier stages are 2^{k-p} factorial or fractional factorial designs (*see Fractional Factorial Designs*), which are followed by response surface designs (*see Response Surface Methodology*) in the search for optimum settings.

One aspect of product quality we are often concerned with is minimizing variability of finished product. In the 1980s Taguchi [1] and others showed us that the product design has a major influence on this variability. Taguchi also showed how the same engineering experiments used to optimize product features could also be used to reduce variability, or dispersion. As an alternative to the Taguchi **signal-to-noise** approach, the concepts of location effects and **dispersion effects** were developed [2] to separate the ideas of factors affecting mean performance of the product and factors affecting the variability of the product, respectively. Specifically, the Box–Meyer method was introduced for identifying dispersion effects from unreplicated and fractionally replicated designs.

Method

The concept of dispersion effects and dispersion modeling can be applied to designs with any number of factors and levels. (For general theory *see Generalized Linear Models*.) The Box–Meyer method was devised for two-level designs such as 2^{k-p} (*see Fractional Factorial Designs*) and Plackett–Burman designs (*see Plackett–Burman Designs*), where k

is the number of experimental factors and n the number of runs. For these designs, we typically have $n - 1$ orthogonal columns of +1's and -1's, labeled $\mathbf{x}_1, \dots, \mathbf{x}_{n-1}$. We use k of the columns to vary the experimental factors in the design. Contrasts calculated from these k columns estimate main effects (with possible confounding with interactions). The remaining unused columns estimate interactions among two or more factors (*see Aliasing in Fractional Designs*).

After we run our two-level fractional designed experiment, location effects are identified in the usual way by any number of methods, for example, half-normal **probability plot** of effects (*see Half-Normal Plot*) or Lenth's method (*see Lenth's Method for the Analysis of Unreplicated Experiments*). Location effects can include both main effects and interactions. Once active location effects are identified, we fit a linear regression model to the data, incorporating those effects. The statistical model is as follows:

$$\mathbf{Y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon} \quad (1)$$

Here $\mathbf{Y}$ is the vector of response variable measurements, $\mathbf{X}$ is the matrix corresponding to active location effects and includes a column for the intercept, $\boldsymbol{\beta}$ is the vector of regression coefficients, and $\boldsymbol{\varepsilon}$ is the vector of errors assumed to be normally distributed with zero mean. We estimate regression coefficients by ordinary least squares (*see Least-Squares Estimation*):

$$\hat{\boldsymbol{\beta}} = (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{Y} \quad (2)$$

The residuals (*see Residuals*), $\mathbf{e}$, from this fit are then used to identify and estimate the dispersion effects.

$$\mathbf{e} = \mathbf{Y} - \mathbf{X}\hat{\boldsymbol{\beta}} \quad (3)$$

For each of the $n - 1$ orthogonal columns $\mathbf{x}_i$ in the original design, the estimated dispersion effect is defined as

$$F_i = \frac{\sum_{j:\mathbf{x}_i=+1} e_j^2}{\sum_{j:\mathbf{x}_i=-1} e_j^2} \quad (4)$$

This is the sum of squared residuals for those runs with $\mathbf{x}_i$ at its high level (+1), divided by the sum of squared residuals for those runs with $\mathbf{x}_i$ at its low level (-1). There will be k dispersion main effects, and $n - 1 - k$ interaction dispersion effects.

2 Box–Meyer Method for Dispersion Effects

Dispersion effects F_i either much greater than 1 or much less than 1 indicate potentially significant effects.

It is important to note that, in general, under the null hypothesis of no dispersion effect, one cannot assume that F_i necessarily follows an F distribution. To identify those estimated dispersion effects that may be important, the logarithms of the estimated dispersion effects can be plotted on a **half-normal plot** to identify those that stand out from the others.

Example

A nine-factor two-level study of welding comes from [1]. The design and responses are shown in Table 1. Factors B and C have significant location effects. After calculating residuals from the regression model with those location effects, dispersion effects were calculated from the residuals. Table 2 shows in detail how the dispersion effect for factor C is calculated. A half-normal plot of the natural logarithm of all the

Table 1 Design and responses for welding example [1].^(a)

Run	A	B	C	D	E	F	G	H	J	Y
1	–	–	–	–	–	–	–	–	–	43.7
2	–	–	+	+	+	+	–	–	+	40.2
3	–	+	+	–	–	–	–	+	–	42.4
4	–	+	–	+	+	+	–	+	+	44.7
5	–	+	+	–	–	+	+	–	+	42.4
6	–	+	–	+	+	–	+	–	–	45.9
7	–	–	–	–	–	+	+	+	+	42.2
8	–	–	+	+	+	–	+	+	–	40.6
9	+	+	+	–	+	–	–	–	+	42.4
10	+	+	–	+	–	+	–	–	–	45.5
11	+	–	–	–	+	–	–	+	+	43.6
12	+	–	+	+	–	+	–	+	–	40.6
13	+	–	–	–	+	+	+	–	–	44.0
14	+	–	+	+	–	–	+	–	+	40.2
15	+	+	+	–	+	+	+	+	–	42.5
16	+	+	–	+	–	–	+	+	+	46.5

^a © ASI Press A : kind of welding rods; B : period of drying; C : welded material; D : thickness; E : angle; F : opening; G : current; H : welding method; J : preheating; Y : tensile strength

Table 2 Levels for factor C , fitted values for Y , and residuals for welding example [1]. The five largest residuals in absolute value occur for $C = -1$, illustrating the large dispersion effect of this factor: $F_C = (3.66875)/(0.19875) = 18.46$

Run	C	Fitted Y	Residuals	Squared residuals	
1	–1	43.4375	0.2625	0.0689	Sum = 3.66875
4	–1	45.5875	–0.8875	0.7877	
6	–1	45.5875	0.3125	0.0977	
7	–1	43.4375	–1.2375	1.5314	
10	–1	45.5875	–0.0875	0.0077	
11	–1	43.4375	0.1625	0.0264	
13	–1	43.4375	0.5625	0.3164	Sum = 0.19875
16	–1	45.5875	0.9125	0.8327	
2	1	40.3375	–0.1375	0.0189	
3	1	42.4875	–0.0875	0.0077	
5	1	42.4875	–0.0875	0.0077	
8	1	40.3375	0.2625	0.0689	
9	1	42.4875	–0.0875	0.0077	
12	1	40.3375	0.2625	0.0689	
14	1	40.3375	–0.1375	0.0189	
15	1	42.4875	0.0125	0.0002	

estimated dispersion effects is shown in Figure 1. The estimates are given in Table 3. One could conclude a possible dispersion effect of factor *C* (welded material). In other words, variability in tensile strength of the weld appears to be higher for one of the welded materials than for the other.

Strengths of the Technique

Ordinarily, experiments to study dispersion would require repeated runs. Much larger experiments would be needed. The Box–Meyer technique can be applied to two-level unreplicated factorial and fractional factorial designs. So a major advancement of

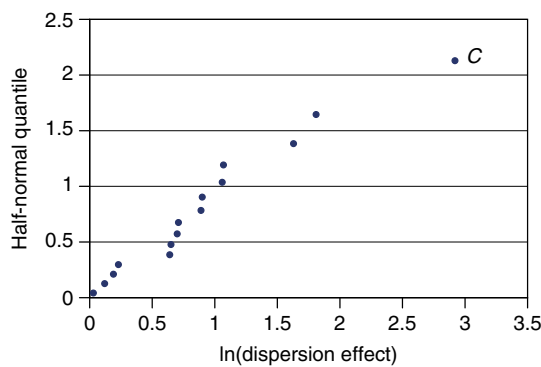


Figure 1 Half-normal plot of natural logarithm of estimated dispersion effects

Table 3 Estimated dispersion effects for welding example [1]

Column	Effect	Aliases	Dispersion effect (F_i)	$ \ln F_i $
1	<i>A</i>	<i>DE, FJ</i>	2.024	0.71
2	<i>B</i>	<i>CD</i>	1.212	0.19
3	<i>C</i>	<i>BD, HJ</i>	18.459	2.92
4	<i>D</i>	<i>AE, BC, FG</i>	1.034	0.03
5	<i>E</i>	<i>AD, GJ</i>	1.922	0.65
6	<i>F</i>	<i>AJ, DG</i>	0.412	0.89
7	<i>G</i>	<i>DF, EJ</i>	0.346	1.06
8	<i>H</i>	<i>CJ</i>	0.164	1.81
9	<i>J</i>	<i>AF, CH, EG</i>	0.197	1.63
10	<i>AB</i>	<i>CE, GH</i>	0.525	0.64
11	<i>AC</i>	<i>BE, FH</i>	0.498	0.70
12	<i>AG</i>	<i>BH, DJ, EF</i>	0.855	0.12
13	<i>AH</i>	<i>BG, CF</i>	2.449	0.90
14	<i>BF</i>	<i>CG, EH</i>	2.921	1.07
15	<i>BJ</i>	<i>DH</i>	0.792	0.23

Box and Meyer was just recognizing that dispersion effects could be estimated from unreplicated designs.

The Box–Meyer method is very simple to understand, very easy to apply, and intuitive. While there are other methods that have been developed for estimating dispersion effects in two-level designs (*see Bergman–Hynén Method for Dispersion Effects; Harvey’s Method for Dispersion Effects*), they are more complicated and typically the estimates are similar.

Weaknesses of the Technique

Application of the method is limited to two-level orthogonal designs, though extensions to more than two levels are straightforward. Modeling dispersion effects in more complex situations, for example, when important covariates must be included in the modeling, requires other more advanced methods such as generalized linear models (*see Generalized Linear Models*). Also, making inferences about dispersion from an unreplicated design is a risky proposition. Conservative scientists may find this an uncomfortable situation. Thus dispersion effects discovered in this manner would more reliably be verified in a confirmatory experiment that included repeat runs. Alternatively, analyzing replicated experiments for dispersion effects is described in [3].

As described in the original article [2], there is an **aliasing** relationship between dispersion effects and location effects in unreplicated designs. One can easily show that estimated dispersion effects are a simple quadratic function of the estimated location effects assumed to be inert (and not included in the location model). This is not a major problem, but real location effects assumed to be inert could falsely dampen or inflate the size of estimated dispersion effects [4].

The Box–Meyer technique does not include a tractable null distribution to judge the statistical significance of potential dispersion effects. The Bergman–Hynén dispersion effect method [5] is a similar method in which the null distribution can be shown to be the usual *F* distribution (under a null hypothesis of no dispersion effects). This allows a statistical test of the null hypothesis of no effect, and construction of an *F*-based **confidence interval** for the dispersion effect.

4 Box–Meyer Method for Dispersion Effects

When a substantial number of effects are included in the location model (1) above, the **degrees of freedom** for estimating dispersion effects are reduced. One should exercise even more caution in interpreting the dispersion effects in these situations.

References

- [1] Taguchi, G. & Wu, Y. (1980). *Introduction to Off-Line Quality Control*, Central Japan Quality Control Association, Nagoya.
- [2] Box, G.E.P. & Meyer, R.D. (1986). Dispersion effects from fractional designs, *Technometrics* **28**, 19–27.
- [3] Nair, V.N. & Pregibon, D. (1988). Analyzing dispersion effects from replicated factorial experiments, *Technometrics* **30**, 247–257.
- [4] Pan, G. (1999). The impact of unidentified location effects on dispersion effects identification from unreplicated factorial designs, *Technometrics* **41**, 313–326.
- [5] Bergman, B. & Hynen, A. (1997). Dispersion effects from unreplicated designs in the 2^{k-p} series, *Technometrics* **39**, 191–198.

R. DANIEL MEYER

Generalized Linear Models

Overview of the Linear Regression Model

“All models are wrong; some models are useful.” This quote from the eminent statistician George Box attempts to put modeling into the proper perspective. While no mathematical model can describe nature exactly, some models do a better job than others. Indeed, we often state assumptions that were made in order to estimate the parameters of our model or to perform some inference knowing in our heart that they do not hold. We rely on the procedure being robust to these assumptions, or assume that they are met asymptotically (i.e., never in practice but close for large data sets). We sometimes start with a model and if there appears to be a disagreement between the model and the data, we change the data through some transformation, known more softly as *reexpression*.

The standard linear regression model is

$$y_i = \beta_0 + \beta_1 x_{i1} + \cdots + \beta_k x_{ik} + \epsilon_i \quad (1)$$

where $E(\epsilon_i) = 0$, $\text{Var}(\epsilon_i) = \sigma^2$, and $\text{Cov}(\epsilon_i, \epsilon_j) = 0$ for $i \neq j$. For notational purposes, we will abbreviate the term linear regression model as LM in this exposition.

In matrix notation, $\mathbf{Y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\epsilon}$ where

$$\begin{aligned} \mathbf{Y} &= \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix}, \mathbf{X} = (\mathbf{1}, \mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_k) \\ &= \begin{pmatrix} 1 & x_{11} & x_{12} & \cdots & x_{1k} \\ 1 & x_{21} & x_{22} & \cdots & x_{2k} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 1 & x_{n1} & x_{n2} & \cdots & x_{nk} \end{pmatrix}, \\ \boldsymbol{\beta} &= \begin{pmatrix} \beta_0 \\ \beta_1 \\ \vdots \\ \beta_k \end{pmatrix}, \boldsymbol{\epsilon} = \begin{pmatrix} \epsilon_1 \\ \epsilon_2 \\ \vdots \\ \epsilon_n \end{pmatrix} \end{aligned} \quad (2)$$

So $E(\mathbf{Y}) = \mathbf{X}\boldsymbol{\beta} = \boldsymbol{\mu}$ and $E(y_i) = x_i'\boldsymbol{\beta} = \mu_i$. Here, $x_i'\boldsymbol{\beta}$ is called the *linear predictor*. In the ordinary least-squares (OLS) approach, we find our vector of estimates, $\mathbf{b}$, of $\boldsymbol{\beta}$ by minimizing the sum of the squared **residuals**, $Q = \sum_{i=1}^n e_i^2 = \sum_{i=1}^n (y_i - x_i'\boldsymbol{\beta})^2$ which, in matrix notation is $Q = (\mathbf{Y} - \mathbf{X}\boldsymbol{\beta})'(\mathbf{Y} - \mathbf{X}\boldsymbol{\beta})$.

The minimum is found by differentiating Q with respect to $\boldsymbol{\beta}$ and equating to 0: $\partial Q / \partial \boldsymbol{\beta} = -2\mathbf{X}'\mathbf{Y} + 2\mathbf{X}'\mathbf{X}\boldsymbol{\beta} = \mathbf{0}$. Note that this may also be expressed as $\mathbf{X}'(\mathbf{Y} - \mathbf{X}\boldsymbol{\beta}) = \mathbf{0}$ or

$$\mathbf{X}'(\mathbf{Y} - \boldsymbol{\mu}) = \mathbf{0} \quad (3)$$

Solving this set of linear equations results in

$$\mathbf{b} = (\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{Y} \quad (4)$$

In order to perform inference about the parameters, or to calculate interval estimates of predictions, a normality assumption is often made. If we assume the error terms follow a multivariate **normal distribution** with common variance and no **correlation**, then we have $\boldsymbol{\epsilon} \sim N_n(\mathbf{0}, \sigma^2 \mathbf{I}_n)$ where $\mathbf{I}_n$ is the $n \times n$ identity matrix. For convenience, we will drop the subscript n when the dimension is obvious.

Another popular method for finding $\mathbf{b}$ is maximum-likelihood **estimation** (MLE). Under the normality assumption described above, we can write the likelihood function as

$$\begin{aligned} l(\mathbf{Y}, \boldsymbol{\beta}, \sigma^2) &= (2\pi\sigma^2)^{-n/2} \exp(-(\mathbf{Y} - \mathbf{X}\boldsymbol{\beta})'(\mathbf{Y} - \mathbf{X}\boldsymbol{\beta}) \\ &\quad \times (2\sigma^2)^{-1}) \end{aligned} \quad (5)$$

The maximum-likelihood estimators of $\boldsymbol{\beta}$ and σ^2 of course maximize the likelihood, or equivalently, the natural logarithm of the likelihood commonly called *log-likelihood* or $\log(l)$. (We use the designation *log* to represent the natural logarithm throughout.) The derivative of $\log(l)$ is known as the *score function* or simply the *score*. Differentiating $\log(l)$ and equating to 0 we have $\partial \log(l) / \partial \boldsymbol{\beta} = (2\sigma^2)^{-1} (-2\mathbf{X}'\mathbf{Y} + 2\mathbf{X}'\mathbf{X}\mathbf{b}) = \mathbf{0}$ which simplifies to $\partial \log(l) / \partial \boldsymbol{\beta} = \mathbf{X}'(\mathbf{Y} - \boldsymbol{\mu}) = \mathbf{0}$. Thus, the maximum-likelihood estimators are the same as those obtained by **least squares**.

If the normality assumption does not hold, we often transform the data, $\mathbf{y}$ to $g(\mathbf{y}) = \mathbf{y}'$ so that the errors more closely resemble a normal distribution. Box-Cox transformations (see **Box and Cox Transformation** and Box and Cox [1]) comprise a family of power transformations that are used frequently in practice.

Introduction to Generalized Linear Models

In this section, we first define the *generalized linear model* (GLM) and then discuss inference and

2 Generalized Linear Models

diagnostic procedures. For a full treatment of GLMs, see the classic text of McCullagh and Nelder [2]. For practitioners in the quality and reliability fields, see Myers *et al.* [3] for a more applied introduction.

The idea behind GLMs is to use a distribution that more naturally models the response instead of “forcing” normality. The distributions that can be employed in GLMs are members of the exponential family. Many common distributions are members of this family including, in addition to the normal, the binomial, Poisson, exponential, and gamma distributions. For these distributions, the probability density function (pdf) may be written as

$$f(y; \theta, \phi) = \exp(a(\phi)^{-1}(y\theta - b(\theta)) + c(y, \phi)) \quad (6)$$

It can be shown that for any random variable Y with a density of the above form,

$$E(Y) = \mu = b'(\theta) \quad (7)$$

and

$$\text{Var}(Y) = b''(\theta)a(\phi) = \frac{d\mu}{d\theta}a(\phi) \quad (8)$$

An advantage of GLMs is that, for example, data that seem to be modeled well by the binomial distribution can be analyzed using the binomial itself instead of transforming to the normal distribution. In GLMs, we use a transformation of μ_i , not of y_i . In other words, we let $\eta_i = g(\mu_i)$ where $g(\cdot)$, called the *link function*, is monotone and differentiable. The function $\eta_i = g(\mu_i) = \theta_i$ is known as the *canonical link*. Using the Poisson distribution as an example we have

$$f(y; \lambda) = \lambda^y \exp(\lambda)(y!)^{-1} = \exp((y \log(\lambda) - \lambda) \times (1 - y!)^{-1}) \quad (9)$$

so $\theta = \log(\lambda)$, $b(\theta) = -\lambda = \exp(-\theta)$, $\phi = 1$, $a(\phi) = 1$, and $c(y, \phi) = y!$. The canonical link is then $\eta_i = g(\mu_i) = \log(\lambda_i)$, that is, the log link. While other links may be (and are) used, selection of the canonical link simplifies the mathematics. If a noncanonical link is used, the algorithm used to find our estimates is more complicated, but that need not concern us here. For a discussion of using noncanonical links, see [3, pp. 165–166].

So we have $\eta_i = g(\mu_i) = \mathbf{x}_i' \boldsymbol{\beta}$, that is, that some function of the mean is related to the experimental

factors through a regression model. As in linear regression, an important goal is to estimate the parameters, $\boldsymbol{\beta}$ and we will use the maximum-likelihood approach to accomplish this. By restricting our attention to the exponential family, we note that $\log(l) = \sum_{i=1}^n a(\phi)^{-1}(y_i \theta_i - b(\theta_i)) + c(y, \phi)$ and using equation (7) we have $\partial \log(l) / \partial \theta_i = \sum_{i=1}^n a(\phi)^{-1}(y_i - b'(\theta_i)) = \sum_{i=1}^n a(\phi)^{-1}(y_i - \mu_i)$. If the canonical link is used, $\eta_i = \theta_i = \mathbf{x}_i' \boldsymbol{\beta}$ and $\partial \theta_i / \partial \boldsymbol{\beta} = \mathbf{x}_i'$. So, using the chain rule, we have $\partial \log(l) / \partial \boldsymbol{\beta} = (\partial \log(l) / \partial \theta) \times (\partial \theta / \partial \boldsymbol{\beta}) = \sum_{i=1}^n a(\phi)^{-1}(y_i - \mu_i) \mathbf{x}_i'$. Finally, equating these to 0 provides $\partial \log(l) / \partial \boldsymbol{\beta} = \sum_{i=1}^n (y_i - \mu_i) \mathbf{x}_i' = 0$, or, in matrix notation

$$\frac{\partial \log(l)}{\partial \boldsymbol{\beta}} = \mathbf{X}'(\mathbf{Y} - \boldsymbol{\mu}) = \mathbf{0} \quad (10)$$

which are the score equations. Although equation (10) appears to be identical to equation (3), there is a difference in $\boldsymbol{\mu}$. For the linear regression model (LM), we had $\mu_i = \mathbf{x}_i' \boldsymbol{\beta}$ which is linear in $\boldsymbol{\beta}$. This not generally true for GLMs. Referring back to our Poisson model in equation (9) with the log link function, we have $\log(\mu_i) = \mathbf{x}_i' \boldsymbol{\beta}$ or $\mu_i = \exp(\mathbf{x}_i' \boldsymbol{\beta})$ which is clearly nonlinear in $\boldsymbol{\beta}$. In general, there is no closed form solution to these score equations. Accordingly, we must use methods of nonlinear regression such as iteratively reweighted least squares. These methods will provide us with the GLM estimator $\mathbf{b}$ which can be shown to be asymptotically unbiased. To find the asymptotic **covariance** matrix of $\mathbf{b}$, we recall from equation (8) that $\text{Var}(Y_i) = \sigma_i^2 = (d\mu_i/d\theta_i)a(\phi)$ so that assuming zero correlation among responses we have $\text{Cov}(\mathbf{Y}) = \text{diag}(\sigma_i^2) = \mathbf{V}$. It can be shown that the asymptotic covariance matrix for $\mathbf{b}$ is

$$\text{Cov}(\mathbf{b})_{\text{can}} = (\mathbf{X}'\mathbf{V}\mathbf{X})^{-1} (a(\phi))^2 \quad (11)$$

Recall for the normal distribution we have $\mu_i = \theta_i$ and $a(\phi) = \sigma^2$. Therefore, $\mathbf{V} = \sigma^2 \mathbf{I}$ and we get the well-known result

$$\text{Cov}(\mathbf{b})_{lm} = (\mathbf{X}'\mathbf{X})^{-1} \sigma^2 \quad (12)$$

So the common LM assuming independent normally distributed errors can be viewed as a GLM using the normal distribution and an identity link, $\eta_i = \mu_i$. If a noncanonical link is used, the computations become more complex and the asymptotic covariance matrix becomes

$$\text{Cov}(\mathbf{b})_{\text{non}} = (\mathbf{X}'\Delta\mathbf{V}\Delta\mathbf{X})^{-1} \sigma^2 (a(\phi))^2 \quad (13)$$

where $\mathbf{\Delta}$ is a diagonal matrix with i th element $\partial\theta_i/\partial\eta_i$.

Given data Y_i that are independent with $E(Y_i) = \mu_i$, the steps for fitting a GLM are outlined here.

1. The distribution of Y_i is specified as a member of the exponential family as previously defined.
2. The linear predictor is $\mathbf{x}_i'\boldsymbol{\beta}$.
3. Some function of the mean equals the linear predictor, $\eta_i = g(\mu_i) = \mathbf{x}_i'\boldsymbol{\beta}$. The function $g(\mu_i)$ is called the *link function* as it “links” the mean to the linear predictor.
4. To find $\mathbf{b}$, the vector of estimators of $\boldsymbol{\beta}$, we must use iteratively reweighted least squares, or specialized GLM software.

Inference in Generalized Linear Models

In the general linear model, parameter inference is relatively straightforward due to the normality assumption. Indeed, the OLS and maximum-likelihood estimator of β_j is b_j with **standard error** $\text{SE}\{b_j\}$ and

$$t_j = b_j(\text{SE}\{b_j\})^{-1} \sim t_{(n-p)} \quad (14)$$

where $\text{SE}\{b_j\}$ is the j th diagonal term of $\text{Cov}(\mathbf{b})_{lm}$ and p is the number of parameters estimated. **Confidence intervals** for and hypothesis tests about β_j are constructed based on this distribution and the inference is exact under normality. Of course the normal distribution is merely a model of the real-life behavior and is therefore an approximation. Similar results are generally found when the data do not depart much from normality. In GLMs, we rely on an asymptotic approximation known as *Wald inference*, that is,

$$\chi_j^2 = (b_j(\text{SE}\{b_j\})^{-1})^2 \sim \chi_{(1)}^2 \quad (15)$$

where $\text{SE}\{b_j\}$ is the j th diagonal term of $\text{Cov}(\mathbf{b})_{can}$ or $\text{Cov}(\mathbf{b})_{non}$ depending on the category of link used. Note the similarity between equations (14) and (15). As $n - p$ increases without bound, $t_{(n-p)}$ approaches $N(0, 1)$, the square of which is $\chi_{(1)}^2$. For GLMs, the approximation may be considered rough for small values of $n - p$ where p is the number of parameters estimated.

Diagnostics in GLMs

In the linear model, $SSE = \sum_{i=1}^n (y_i - \hat{\mu}_i)^2 = \sum e_i^2$ is a measure of the variation unexplained by the

model. If a saturated model (one in which n parameters are estimated) is fit, then $SSE = 0$. So SSE is the difference between variation left from the fitted model and the saturated model. An analogous term in GLMs is called the *deviance*, which is a function of likelihoods. A saturated model will provide the maximum possible likelihood for the observed data. Any model with fewer parameters will have a smaller likelihood. The *deviance* is defined as

$$D = -2 [\log(l_{\text{fit}}) - \log(l_{\text{sat}})] \quad (16)$$

The subscripts “fit” and “sat” in equation (16) refer to the fitted and saturated models respectively. Since each $\log(l)$ is a sum of n components, we can write $D = \sum_{i=1}^n D_i$ where $D_i \geq 0$ and define the deviance residuals as $d_i = \text{sgn}(e_i)\sqrt{D_i}$ so that the deviance and traditional residuals have the same sign. These d_i may be used for diagnostic purposes as the e_i are used in LM. For example, a normal **probability plot** of e_i is often used to check the normality assumption in LMs. Significant departures from a line on this plot may suggest a transformation is warranted. Similarly, the d_i from a GLM may be plotted on a normal probability plot and significant departures from a line suggest consideration of another link function.

Generalized Linear Models in 2^{k-p} Experiments

GLMs have been employed successfully in many different areas such as the biological and social sciences. Often they are used in **observational studies** as an alternative to OLS estimation in multiple regression. In these studies, the predictor variables, the X 's, are not controlled but merely observed. Also, it is often known that the normal distribution provides a poor model for the response. However, in the engineering and physical sciences, experiments are often performed where the predictors are under the experimenter's control and it is often assumed that the normal distribution provides an adequate model. If there appears to be **skewness** or nonconstant variance, the data may be transformed, perhaps using a Box–Cox transformation (see [1]). GLMs may be applied in either of these settings as well as others. In this section, we will focus on the use of GLMs in two-level fractional factorial (2^{k-p}) designs (see **Fractional Factorial Designs**). We will discuss three uses of GLMs in such experiments: with attribute

4 Generalized Linear Models

data such as defects or defective counts, for identifying factors that affect the mean of a continuous response, and finally for estimating both mean (location) and variance (dispersion) effects for continuous responses.

Attribute Data

In many industrial experiments, it is prohibitively expensive, or even impossible, to use a quantitative variable or variables as quality characteristics. We may be forced into recording the number of defects on tested units or just the number of defective units (those having one or more defects). If m_i units are inspected in run i of an experiment, then the only possible values for the number of defective units are $0, 1, \dots, m_i - 1, m_i$. If we assume the probability of a unit being defective within an experimental run i is a constant, p_i , and that the units are independent, then the binomial distribution is a reasonable model. If y_i is the number of defective units in experimental run i and we let $\hat{p}_i = y_i m_i^{-1}$, then $0 \leq \hat{p}_i \leq 1$. Under a binomial model, $E(Y_i) = m_i p_i$ and $\text{Var}(Y_i) = m_i p_i (1 - p_i)$. Clearly the variance is a function of the mean. Assuming a normal distribution for values that are bounded can lead to predicted values that are not feasible, for example, greater than 1. In order to use an LM, some authors suggest using the transformation $\hat{p}' = \arcsin(\sqrt{\hat{p}})$. This is known as a *variance stabilizing transformation*. The data would be analyzed on this scale and then transformed back to the original for interpretation.

If the number of defects is the response, then we know the possible values for y_i are $0, 1, \dots$. If the defects are assumed to occur independently, then the Poisson distribution may provide a reasonable model. Using equation (9) for the Poisson distribution

we have $E(Y) = \lambda$ and $\text{Var}(Y) = \lambda$. Again, the variance is a function of the mean and the LM can produce undesirable values less than 0. To use an LM, the variance stabilizing transformation $y' = \sqrt{y}$ is often employed. As an alternative to these transformations, we can make use of GLMs, directly modeling the response with either the binomial or the Poisson distribution with an appropriate link function.

We now provide an example wherein each unit within a run is classified as acceptable or defective. The experiment is a 2^3 design with factors A , B , and C coded as -1 and $+1$ for their low and high levels respectively. In run i of the experiment, we observe y_i , the number of acceptable units out of the $m_i = m = 100$ units tested. Note that since m_i is a constant, yield may be described as number or percentage or fraction ($\hat{p}_i = y_i m_i^{-1}$) of acceptable units. We use three models: (a) LM on the observed data (LM1), (b) LM using the transformation $y' = \arcsin(\sqrt{\hat{p}_i})$ (LM2), and (c) a GLM with the binomial distribution and a logit link function (GLM). Letting $E(Y_i) = m_i p_i$ and $y_i m_i^{-1} = \hat{p}_i$, the logit link is $\log(p_i(1 - p_i)^{-1})$. The first four columns of Table 1 display the design for our example and the response y_i recorded as yield. Table 2 contains the parameter estimates as well as common inference output.

Analysis of the three models shows that factors A and B are active and that factor C most likely is not. AB is highly significant for the analysis on the original scale and is mildly significant on the transformed scale. For these models, we include A , B , and AB as active. For the GLM, AB had a p value of 0.66 and so only the main effects A and B were included. The parameter estimates in Table 2 were used to find the fitted values (reported as fraction acceptable) for each model. For LM1, the fitted values are simply obtained

Table 1 2^3 design with response and fitted values for each of three models

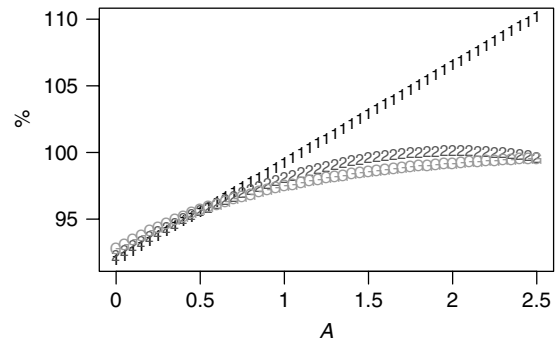
A	B	C	Yield	Original	Transformed	GLM
-1	-1	-1	30	0.320	0.325	0.320
1	-1	-1	81	0.810	0.805	0.810
-1	1	-1	72	0.740	0.735	0.740
1	1	-1	93	0.955	0.960	0.959
-1	-1	1	34	0.320	0.325	0.320
1	-1	1	81	0.810	0.805	0.810
-1	1	1	76	0.740	0.735	0.740
1	1	1	98	0.955	0.960	0.959

Table 2 Output for (a) linear model untransformed, (b) linear model transformed, and (c) generalized linear model, logit link, and binomial distribution

(a)				
Predictor	Coefficient	SE coefficient	T	P
Constant	70.6250	0.9437	74.84	0.000
A	17.6250	0.9437	18.68	0.000
B	14.1250	0.94377	14.97	0.000
AB	-6.8750	0.9437	-7.28	0.002
(b)				
Predictor	Coefficient	SE coefficient	T	P
Constant	1.03071	0.01757	58.65	0.000
A	0.21216	0.01757	12.07	0.000
B	0.17028	0.01757	9.69	0.001
AB	-0.04718	0.01757	-2.68	0.055
(c)				
Predictor	Coefficient	SE coefficient	T	P
Intercept	1.2202	0.1022	142.57	0.000
A	1.0756	0.1005	114.44	0.000
B	0.8766	0.0965	82.43	0.000

by substituting the appropriate values of A and B and dividing by 100. The fitted values for LM2 were obtained as $\hat{p}_i = \sin^2(1.03071 + 0.212126A + 0.17028B - 0.04718AB)$ and for the GLM, $\hat{p}_i = e^{1.2202 + 1.0756A + 0.8766B} / (1 + e^{1.2202 + 1.0756A + 0.8766B})^{-1}$. These values from the three models occupy the last three columns of Table 1. Clearly, all three models give similar fitted values. All models suggest higher yields should occur at values of A and B greater than those studied, that is, outside the range of $(A, B) \in [-1, 1]^2$.

So what if we look outside the range of values studied? While extrapolation should always be done with caution, we will extrapolate here to see what happens to our predicted values beyond the range of data. Since LM1 and LM2 contain the AB interaction, we can find the values of A and B that give the maximum predicted yield. These are $(A, B) = (2.05, 2.56)$ and $(4.50, 3.61)$ for LM1 and LM2 respectively. While not extrapolating that far, Figure 1 shows $100\hat{p}_i$ as a function of A with $B = 1.5$ (for illustrative purposes) for each of the three models. Note that LM1 increases without bound (for fixed B), quickly surpassing the highest possible value of 100%. The transformation in LM2 avoids

**Figure 1** Estimated yield (in proportion) for three models. 1: LM on original data, 2: LM on arcsin transformed data, and G: GLM with binomial distribution and logit link

this problem, achieving its maximum of 100% near $A = 2$ for $B = 1.5$. The GLM, fit without an interaction, approaches 1 asymptotically. So the three models give similar results within the range of data studied, but differ substantially outside this range. While we cannot tell which provides the best fit between LM2 and GLM, the GLM model has the advantage of simplicity in that it fits the data well without including the interaction term.

Location Effects

As opposed to producing attribute (count) data, many industrial experiments have quantitative quality characteristics such as resistivity, voltage, diameter, and so on. While a normal distribution may provide a reasonable model for many such characteristics, others may be known to have a skewed distribution. For example, a simple model for measuring lifetimes in the reliability field is the exponential distribution. In other applications, tool wear often leads to skewed distributions of quality characteristics which can be modeled by the gamma distribution, of which the exponential is a special case.

Myers *et al.* [3, p. 176] provide an example of a 2^4 design to study the effects of four factors labeled X_1 , X_2 , X_3 , and X_4 on the response, resistivity. The design matrix is given in columns 1–4 of Table 3 with labels in the 2^4 row and the resistivity values in the y_1 column. They assumed, as we do, that the three and four factor interactions are negligible and are pooled into an error term for inference purposes. As resistivity is known to have a positively skewed distribution, we analyze the data using a log transformation on y , $y' = \log(y)$ with an LM. Table 4(a) shows edited output from SAS. As an alternative, we fit a GLM with a gamma distribution

and a log link and display edited output from SAS Proc Genmod in Table 4(b). The Appendix shows the SAS code (version 9.1) needed to perform these analyses. Note the similarities in the code.

We note that in part (a) of the table, t statistics from equation (14) are used for testing and in part (b) χ^2 statistics from equation (15) are used. Also, the parameter estimates are very similar, but the LM estimates have standard errors of 0.0206 compared to the asymptotic standard error of 0.0115 for the GLM estimates, a 44% reduction. This difference in standard errors results in drawing different conclusions from the two models. In the LM, X_1 , X_2 , and X_3 are clearly active and X_2X_3 and X_3X_4 show mild significance. However, in the GLM, X_1 , X_2 , X_3 , X_4 , X_1X_3 , X_2X_3 , and X_3X_4 are all highly significant. The addition of X_4 and noting that X_3 also interacts with X_1 was missed by the LM. Regardless of which model is used, we may choose to remove the nonsignificant terms and refit the model. After doing so, of course, we should use diagnostic procedures such as residual plots to assess each model. For the GLM, the deviance residuals (d_i)s may be used for this purpose. As previously mentioned, plots of the deviance residuals may display **outliers** or suggest investigation of a different link function. In the interest of brevity, we do not show these plots.

Table 3 Design matrix and response for resistivity experiment (2^4 and y_1) and welding experiment (2^{9-5} and y_2)

[illegible]

Table 4 Resistivity example (a) linear model with log transformation and (b) generalized linear model, gamma distribution with log link

(a)				
Parameter	Estimate	Standard error	t	$Pr > t $
Intercept	5.4131	0.0206	262.93	< 0.0001
x_1	0.0614	0.0206	2.98	0.0307
x_2	-0.1499	0.0206	7.28	0.0008
x_3	0.0892	0.0206	4.33	0.0075
x_4	-0.0281	0.0206	-1.36	0.2309
$x_1 \times x_2$	0.0049	0.0206	0.24	0.8230
$x_1 \times x_3$	-0.0391	0.0206	-1.90	0.1163
$x_1 \times x_4$	0.0019	0.0206	0.09	0.9289
$x_2 \times x_3$	-0.0446	0.0206	-2.16	0.0827
$x_2 \times x_4$	-0.0110	0.0206	-0.54	0.6154
$x_3 \times x_4$	-0.0458	0.0206	-2.22	0.0769
(b)				
Parameter	Estimate	Standard error	χ^2	$Pr > \chi$
Intercept	5.4141	0.0115	221557	< 0.0001
x_1	0.0617	0.0115	28.70	< 0.0001
x_2	-0.1496	0.0115	169.01	< 0.0001
x_3	0.0900	0.0115	61.24	< 0.0001
x_4	-0.0280	0.0115	5.91	0.0151
$x_1 \times x_2$	0.0050	0.0115	0.19	0.6640
$x_1 \times x_3$	-0.0386	0.0115	11.27	0.0008
$x_1 \times x_4$	0.0020	0.0115	0.03	0.8644
$x_2 \times x_3$	-0.0441	0.0115	14.70	< 0.0001
$x_2 \times x_4$	-0.0110	0.0115	0.91	0.3402
$x_3 \times x_4$	-0.0456	0.0115	15.71	< 0.0001

Location and Dispersion Effects

Experimenters have been interested in studying process means (location) for many years. Taguchi *et al.* [9] generated much interest in developing products and processes that are robust to noise. A **robust design** is one in which we observe little variation at the target value of the process. While there are many different approaches to accomplish this (see **Robust Design**), the general principle is to find factors that affect the variation in the response (**dispersion effects**) and set them at levels to minimize the variance. We then find factors that affect the mean of the response (location effects), but do not impact the variance, and set these to adjust to the desired target value of the response. The mechanics of doing so depend on whether the experiment is replicated or not.

As before, let $\mathbf{X}$ be a matrix whose first column consists of 1's and the rest be potential location effect columns. In a similar manner, let $\mathbf{Z}$ be a matrix whose

first column consists of 1's and the rest be potential dispersion effect columns. Often the columns in $\mathbf{Z}$ are a subset of those in $\mathbf{X}$ but this is not necessary. The generic model we will follow is given in equation (17).

$$Y_i = \mu_i + \epsilon_i, g(\mu_i) = \mathbf{x}_i' \boldsymbol{\beta}, \sigma_i^2 = e^{z_i' \boldsymbol{\gamma}} \quad (17)$$

where $E(Y_i) = \mu_i$, $E(\epsilon_i) = 0$, and $\text{Var}(\epsilon_i) = \sigma_i^2$. $\boldsymbol{\beta}$ is the collection of location effects and $\boldsymbol{\gamma}$ represents the dispersion effects. We must determine which columns are to be included in the location model and which are to be included in the dispersion model. Often, we want to initially consider all location effects that can be estimated. In this case, we may use a half-normal probability plot of the location effect estimates (see **Half-Normal Plot**) or Lenth's method (see **Lenth's Method for the Analysis of Unreplicated Experiments**). The next two sections describe the different approaches for replicated and unreplicated experiments.

Location and Dispersion Effects in Replicated Experiments. Suppose that in the previously discussed resistivity example, each of the 16 factor combinations (treatments) in Table 3 was repeated m times, resulting in m observations for each row. This is called a *replicated experiment* in that each combination is repeated (replicated) multiple times. Accordingly, in run i of the experiment, we can calculate the sample mean $\bar{y}_i = \sum_{j=1}^m y_{ij} m^{-1}$ and variance $s_i^2 = (m-1)^{-1} \sum_{j=1}^m (y_{ij} - \bar{y})^2$. In the LM framework one may choose to model s_i or $\log(s_i^2)$ using the square root or log transformation respectively for investigating dispersion effects. Alternatively, we can model the s_i^2 's using a GLM with an appropriate distribution and link function and identify and estimate dispersion effects. These dispersion effects may then be used in the estimation of the location effects.

For example, assume the response is normally distributed with $\mathbf{Y} \sim \mathbf{N}(\boldsymbol{\mu}, \mathbf{V})$ where $\mathbf{V} = \text{diag}(\sigma_i^2)$. Then $(m-1)s_i^2\sigma_i^{-2}$ has a gamma distribution. We could use a GLM with a gamma distribution and, say, a log link to analyze the sample variances. If all columns in $\mathbf{X}$ are candidates, we can use Lenth's method (see **Lenth's Method for the Analysis of Unreplicated Experiments**) to determine which are considered active. We then save the fitted values, $\hat{\sigma}_i^2$ from the variance model and form the matrix $\hat{\mathbf{V}} = \text{diag}(\hat{\sigma}_i^2)$. Lastly, we can perform weighted least-squares (WLS) analysis to estimate the location effects. In WLS, we estimate the location effects as

$$\mathbf{b}_w = (\mathbf{X}'\hat{\mathbf{V}}^{-1}\mathbf{X})^{-1}\mathbf{X}'\hat{\mathbf{V}}^{-1}\mathbf{Y} \quad (18)$$

(See Kutner *et al.* [4] for details on WLS estimation.) The implementation of this two-step procedure is relatively straightforward so we next discuss location and dispersion effect estimation with unreplicated data.

Location and Dispersion Effects in Unreplicated Experiments. It is much more difficult to study dispersion effects in unreplicated experiments. Generally, we need to *identify* location and dispersion effects before using GLM models for estimation and testing. Again, location effect identification techniques described in **Half-Normal Plot** and **Lenth's Method for the Analysis of Unreplicated Experiments** may be used. For information on dispersion effect identification in unreplicated designs under the LM framework, see **Bergman-Hynén Method for**

Dispersion Effects; Box-Meyer Method for Dispersion Effects; Harvey's Method for Dispersion Effects, and McGrath and Lin [5] for examples. We will outline a GLM approach that shares much in common with the work of Lee and Nelder [6], Engel and Huele [7], and Pan and Taam [8] as outlined in [3, pp. 294–297].

With replication, we first estimated the variances and then used these as weights in the procedure for estimating the means. However, without replication we have no estimate of variance at each run and therefore have no initial weights. So we first fit the model for the means without weights and store the residuals. The squared residuals then become the response for the variance model. The fitted values of the variance model are stored and their reciprocals become the weights for estimating the means. Residuals are again calculated and the process is reiterated until the parameters converge.

We list the specific steps in the procedure when we assume a normal model with identity link for the location model and a gamma distribution with a log link for the dispersion model. These steps are easily generalized to other distributions and links. This method is called the *double GLM approach*.

1. Obtain the initial location effect estimates, $\mathbf{b}^{(0)}$ by using OLS.
2. Calculate the residuals (e_i) from step 1, or later, step 4.
3. Letting e_i^2 be the responses for a GLM with a gamma distribution and a log link, find the dispersion effect estimates $\boldsymbol{\gamma}^{(0)}$ and the estimated variances, $\hat{\sigma}_i^{2(0)}$.
4. Perform WLS estimation of the means by substituting $\hat{\mathbf{V}}_0 = \text{diag}(\hat{\sigma}_i^2)$ for $\mathbf{V}$ in (18). Call the new location effect estimates $\mathbf{b}^{(1)}$.
5. Repeat steps 2–4 until the parameters converge.

The above procedure is based on MLE. Unfortunately, the variance estimates from this method are biased as no adjustment is made for the loss of **degrees of freedom** from estimating the mean model. These underestimations lead to **bias** in the standard errors for the location effect estimates. A method that makes an adjustment for this loss of degrees of freedom is called *restricted maximum-likelihood estimation* (REML). From a practical point of view, the only change to the MLE procedure above is that instead of using e_i^2 as the response in step 3, we

use $e_i^2 + h_{ii}\hat{\sigma}_i^2$ where h_{ii} is the i th diagonal term of $\mathbf{H} = \hat{\mathbf{V}}^{-1/2}\mathbf{X}(\mathbf{X}'\hat{\mathbf{V}}^{-1}\mathbf{X})^{-1}\mathbf{X}'\hat{\mathbf{V}}^{-1/2}$ (see [7] and [[3], pp. 296–297]).

As an example of using GLMs to estimate location and dispersion effects in unreplicated designs, we study a welding example previously discussed by many (see [8] and references therein). Nine factors are studied in a 2^{9-5} design as shown in Table 3 with the response of weld strength in the y_2 column. In row 2^{9-5} of Table 3, we use the factor labeling of [8]. From the half-normal probability plot of the 15 location effect estimates (not shown), it is clear that the largest location effects by far are B and C . Assuming these constitute the location model, the Bergman and Hynén method finds significant dispersion effects H , J , and C . We run the double GLM REML procedure described above and obtain the results given in Table 5.

It is clear from this output that the main effects of B and C are significant, as expected from the **half-normal plot**. However, while the Bergman and Hynén method found the three dispersion effects C , H , and J to be active, the double GLM method indicates that H is not active. This dataset exemplifies the difficulty of estimating location and dispersion effects in unreplicated experiments. Indeed, different authors suggest different models for this example (see [6, 7], and [8]). Even different algorithms for finding the REML estimates may give varying results. (The output in Table 5 was obtained by running the dglm script in S-plus 6.1.) While more research is needed in this area, it is clear that GLMs are useful for attacking this difficult problem.

Summary

GLMs can be viewed as an extension of standard regression models where the distributional assumptions are broadened from the normal distribution to members of the exponential family. This allows, for example, count data to be modeled by a more natural distribution such as the binomial or Poisson instead of transforming the data to achieve near normality. The same can be accomplished for continuous data by using, for example, the gamma distribution.

In the interest of brevity, we have left out many details of generalized linear modeling. We encourage those interested in the more theoretical aspects to consult [2]. For more details on applications in the quality and engineering areas, see [3].

While GLMs have become somewhat common in the biological sciences, for example, they are slowly being recognized as useful models in engineering and quality and reliability studies. We hope that our discussion of uses in fractional **factorial** designs will encourage further use.

Appendix

SAS Code for the Welding Example

```
proc glm ;
  model logy = x1 x2 x3 x4 x1*x2 x1*x3
    x1*x4 x2*x3 x2*x4 x3*x4;
  output out=new p=yhat r=resid ;
run;

proc genmod data=gamma;
  model y = x1 x2 x3 x4 x1*x3 x2*x3
    x3*x4/dist=gamma link=log;
```

Table 5 Location and dispersion effect estimation example

Mean coefficients			
	Value	Standard error	<i>t</i> value
(Intercept)	43.824199	0.10391390	421.73566
B	1.862278	0.04806503	38.74496
C	3.237478	0.10408022	31.10560
Dispersion coefficients			
	Value	Standard error	<i>t</i> value
(Intercept)	−3.15889105	0.9481690	−3.33156965
C	2.73527516	0.9627868	2.84099779
H	−0.08581873	0.9627868	−0.08913576
J	3.33236105	0.9627868	3.46116198


```
*dataset out2 contains fitted values,
residuals, and deviance residuals;
output out=out2
      p=pred
      reschi=res
      resdev=dev;
run;
```

References

- [1] Box, G.E.P. & Cox, D.R. (1964). An analysis of transformations, *Journal of the Royal Statistical Society, Series B* **26**, 211–243.
- [2] McCullagh, P. & Nelder, J.A. (1989). *Generalized Linear Models*, 2nd Edition, Chapman & Hall, London, Great Britain.
- [3] Myers, R.H., Montgomery, D.C. & Vining G.G. (2002). *Generalized Linear Models with Applications in Engineering and the Sciences*, John Wiley & Sons, New York.
- [4] Kutner, M.H., Nachtsheim, C.J. & Neter, J. (2004). *Applied Linear Regression Models*, McGraw-Hill/Irwin, New York, pp. 421–431.
- [5] McGrath, R.N. & Lin, D.K.J. (2001). Testing multiple dispersion effects in unreplicated fractional designs, *Technometrics* **43**, 406–414.
- [6] Lee, Y. & Nelder, J.A. (1998). Generalized linear models for the analysis of quality-improvement experiments, *Canadian Journal of Statistics* **26**, 95–105.
- [7] Engel, J. & Huele, A.F. (1996). A generalized linear modeling approach to robust design, *Technometrics* **38**, 363–373.
- [8] Pan, G. & Taam, W. (2002). On generalized linear model method for detecting dispersion effects in unreplicated factorial designs, *Journal of Statistical Computation and Simulation* **72**, 431–450.
- [9] Taguchi, G., Chowdhury, S. & Wu, Y. (2004). *Taguchi's Quality Engineering Handbook*, John Wiley & Sons, New York.

RICHARD N. MCGRATH

Harvey's Method for Dispersion Effects

Introduction

Estimation of location and **dispersion effects** is typically done through replicated experiments (see [1–3]), but due to the cost of such experiments there has been much interest in developing methods for identification and estimation of these effects in unreplicated experiments. A similar problem for identifying just location effects in unreplicated designs has been studied extensively and a review of the various methods is summarized in [4]. Several methods have been proposed for studying dispersion effects from unreplicated (fractional) **factorial** designs [5–8].

The problem of identifying location and dispersion effects from unreplicated 2^{k-p} designs is difficult and relies on the affect sparsity principle given that the number of runs is $N = 2^{k-p}$ while the potential number of parameters is $2N$. However, even under the effect sparsity principle, Pan [9] and McGrath and Lin [10] discuss the problem of misidentification of the location and the impact on inference for dispersion effects. In addition, under the assumption that the location model is correctly identified, the presence of multiple dispersion effects can greatly influence the dispersion analysis (see [8]). It is also shown in [8] that all noniterative methods exhibit biases that depend on the fitted location model and active dispersion effects (parameters other than the one being estimated) and therefore can be large. This **bias** can lead to spurious identification of effects as well as underestimation or overestimation of real effects.

This article takes a look at the two-step approach proposed in [8] which incorporates a noniterative identification step followed by an iterative final estimation step. The noniterative step is a modification of Harvey's method (see [11]) and the iterative step is restricted maximum likelihood (REML). This two-step approach eliminates some of the bias present in the noniterative methods.

Notation, Model, and Methods

This section provides the notation that will be used throughout along with the assumed dispersion model

and a description of the two-step modified Harvey (MH) and REML estimation scheme. Throughout we consider only two-level (fractional) factorial designs along with the key assumption that the true location model $\mathcal{L}$ is correctly identified during the initial location-modeling stage. The misidentification of the location model can lead to confounding of location and dispersion effects, as discussed in [9, 10].

Notation

The following notation will be used throughout:

1. $X = [X_0 | \cdots | X_m]$ denotes an $N \times (m+1)$ design matrix corresponding to a 2^{k-p} fractional design with $N = 2^{k-p}$.
2. X_0 is the first column of X and is a vector of $+1$ s; for $j \geq 1$, X_j is the vector of $+1$ s and -1 s corresponding to the high and low levels of factor j .
3. x_i denotes the i th row of X .
4. $S(j+) = \{i : X_{i,j} = +1\}$, where $X_{i,j}$ is the i th element of X_j . Similarly, $S(j-) = \{i : X_{i,j} = -1\}$. In addition, $S(j_1+, j_2+) = \{i : X_{i,j_1} = +1, X_{i,j_2} = +1\}$; similarly for $S(j_1+, j_2-)$, $S(j_1-, j_2+)$, and $S(j_1-, j_2-)$. This notation can be generalized to $S(j_1\pm, \dots, j_m\pm) = \{i : X_{i,j_1} = \pm 1, \dots, X_{i,j_m} = \pm 1\}$.
5. For $j_1, j_2, j_3 \in \{1, 2, \dots, N\}$, we write $j_1 = j_2 \circ j_3$ if X_{j_1} can be obtained by elementwise multiplication of X_{j_2} and X_{j_3} . For example, the interaction term $AB = A \circ B$.
6. Let $\mathcal{L} = \{0\} \cup \{l_1, \dots, l_q\}$ represent the set of active location effects and let $\mathcal{D} = \{0\} \cup \{d_1, \dots, d_r\}$ represent the set of active dispersion effects. Letters such as A , B , and so forth will be used interchangeably with l_j and d_j . Both 0 and I denote the intercept.
7. A set $\mathcal{S}$ is closed if it is closed under the $\circ$ operation. The closure of a set $\mathcal{S}$ is denoted by $\overline{\mathcal{S}}$. Given a set $\mathcal{S}$, define the set $\mathcal{S}_{-j} = \mathcal{S} - \{j\}$.
8. $(r_1, \dots, r_N)'$ denotes the vector of **residuals** from a fitted location model.

Models

We begin with the model for the replicated case. Suppose the data, possibly after a suitable transformation, follow the location-scale model

$$Y_{ir} = \mu_i + \sigma_i \varepsilon_{ir}, \quad i = 1, \dots, N, \quad r = 1, \dots, R \quad (1)$$

2 Harvey's Method for Dispersion Effects

where R represents the number of replications of each design point and ε_{ir} are independent and identically distributed (i.i.d.) with mean 0 and variance 1. A structural model relates the mean, μ_i , to the location effects, β_j , and the variances, σ_i^2 , to the dispersion effects, ϕ_j , through link functions $g(\cdot)$ and $h(\cdot)$, $g(\mu) = X\beta$, $h(\sigma^2) = X\phi$. The linear link function is the most common one for the mean and for the variance we will consider the multiplicative or log-linear model

$$\sigma_i^2 = \exp(x_i' \phi), i = 1, \dots, N \quad (2)$$

A linear dispersion model is also considered in [8] where they provide modeling strategies for that case. The log-linear dispersion model is used most frequently since it does not require constraints on the parameters to ensure the estimates of σ_i^2 are positive.

With replicated data, the location and dispersion analysis can be based on the mean $\bar{Y}_i$ and sample variance S_i^2 of the replications at each design run. For dispersion effects, one can fit an appropriate model to the sample variances and estimate the dispersion effects ϕ_j . A linear analysis of the sample variances is discussed in [12, 13] for random-effects and variance-components models. The idea of doing the log-linear analysis of sample variances goes back to Bartlett and Kendall (see [1]). See [3] for a comparison of the properties of this estimation method with other methods including the maximum likelihood (ML).

We will now look at the case of interest with no replications, that is model (1) with $R = 1$. The method that we will focus on is the two-step approach proposed in [8] which incorporates a noniterative identification step followed by an iterative final estimation step. The noniterative step is a modification of Harvey's method [11] and the iterative step is REML.

Methods

Harvey's Method. With a log-linear dispersion model the most natural extension of Bartlett and Kendall's (see [1]) idea applied to the unreplicated case is to do a log-linear analysis of the squared residuals instead of the sample variances. That is, we fit a log-linear model

$$\log r_i^2 = x_i' \phi + u_i \quad (3)$$

and estimate the dispersion effects ϕ_j through **least squares**. Harvey [11] studied this method in the

general regression context. See also [14] for results about modeling functions of residuals for doing variance estimation. We refer to this method as Harvey's method (H method).

For a 2^{k-p} design, the estimate of ϕ_j from the H method can be written as

$$\begin{aligned} D_j^H &= \frac{1}{N} \left(\sum_{S(j+)} \log r_i^2 - \sum_{S(j-)} \log r_i^2 \right) \\ &= \log \left(\frac{\prod_{S(j+)} r_i^2}{\prod_{S(j-)} r_i^2} \right) \end{aligned} \quad (4)$$

It is shown in [3] that for the replicated case, where the sample variances are used instead of the squared residuals, D_j^H in equation (4) is the ML under a fully saturated dispersion model and is recommended in the initial model-identification stage. It is pointed out in [8] that this method, along with the methods of [5–7] are all biased in the context of multiple active dispersion effects. The amount of bias depends on the fitted location model along with the active dispersion effects. However, unlike the other methods, the size of the bias of the H method decreases as the run size increases.

Modified Harvey Method. A modification to the H method is proposed in [8] to decrease the bias of the H method. This modified method (MH method) is obtained by applying the H method to the modified residuals from an expanded location model discussed in [7]. The idea is that we first estimate a location model through least squares and identify the location model $\mathcal{L}$. For each dispersion effect j , let $\mathcal{L}_j^E$ be the expanded location model, which is the union of $\mathcal{L}$, $\{j\}$, and all interactions between j and the elements of $\mathcal{L}$. For example, if $\mathcal{L} = \{I, A, B\}$ and $j = C$, then $\mathcal{L}_j^E = \{I, A, B, C, AC, BC\}$. If we let $\{\tilde{r}_i, i = 1, \dots, N\}$ denote the residuals from the expanded location model, then the MH estimator is

$$D_j^{MH} = \frac{1}{N} \left(\sum_{S(j+)} \log \tilde{r}_i^2 - \sum_{S(j-)} \log \tilde{r}_i^2 \right) \quad (5)$$

Iterative Methods. Once a tentative model for location and dispersion effects has been identified,

one can use iterative methods to refine the model and estimate the parameters more efficiently. One approach is to use the full ML for location and dispersion effects under normal theory (e.g., see [11]) or under **generalized linear models** (see [15–17]). Another approach is to use a log-linear model to estimate the dispersion effects (as in the H or MH method) and then iterate, using iteratively reweighted least squares (IRLS), to refit the location model at each step. Both approaches are discussed in [18], but they do not discuss model-selection issues. Lee and Nelder [19] developed REML for dispersion analysis that takes into account the loss of **degrees of freedom** from location model estimation. Iterative methods are considered later in the context of the proposed modeling strategy from [8] where they show that the bias is reduced even further through the use of REML at the final estimation stage.

Properties of the H and MH Methods

Properties of the H Method

For replicated studies, it is shown in [3] that the original Bartlett–Kendall method is unbiased. However, in unreplicated 2^{k-p} experiments the corresponding H method is biased, and the bias is related to the structure of the active dispersion effects as well as the fitted location model (see [8]). The main result is for the case when the location model is closed.

Proposition 1. Suppose $\mathcal{L} = \overline{\mathcal{L}}$ with $\dim(\mathcal{L}) \leq N/2$. Then D_j^H is unbiased if $j \notin \overline{\mathcal{D} - \mathcal{L}}$.

Let us take a look at some special cases of the proposition: (a) If $\mathcal{D} \subset \mathcal{L}$, then the estimators of all the active dispersion effects are unbiased. (b) Let $\mathcal{L} = \{I, A\}$ and $\mathcal{D} = \{I, d\}$. If $d = A$, then D_j^H is unbiased for all j . If $d \neq A$, then D_j^H is biased for $j = d$ and unbiased for others.

An example will also illustrate the point of the bias present in the H method. Consider the case where $\mathcal{L} = \{I, A\}$ and we want to determine if factor A has a dispersion effect. If $\mathcal{D} = \{I, A\}$, then D_A^H is unbiased. But if $\mathcal{D} = \{I, A, B, AB = C\}$ then D_A^H is biased. Figure 1(a) shows the distribution of D_A^H when $\phi_A = 1, \phi_B = -1.67$, and $\phi_{AB} = 1.67$. Note that the distribution of D_A^H is centered at 0 even though factor A has an active dispersion effect. So in the presence of active dispersion effects for B and AB (or its aliased **main effect** C), one can miss an active dispersion effect for A . Similarly, if $\phi_A = 0$ instead of $\phi_A = 1$, then the distribution of D_A^H would be centered at -1 , and a spurious dispersion effect for A may occur when B and AB have active dispersion effects.

Properties of the MH Method

The MH method removes some of the bias present in the H method. For example, when $\mathcal{L} = \{I\}$ and $\mathcal{D} = \{I, A\}$ we know from Proposition 1 that D_A^H is biased. However, by purposefully fitting

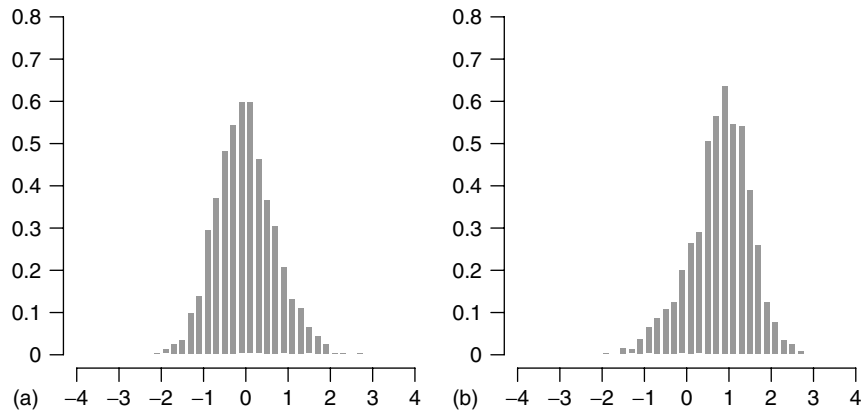


Figure 1 Distribution of noniterative and iterative estimators of ϕ_A . From the example where $\phi_A = 1, \phi_B = -1.67$, and $\phi_{AB} = 1.67, N = 2^2$: (a) distribution of D_A^H (D_A^{MH}); (b) distribution of the ML estimate of ϕ_A [© American Statistical Association, 2001.]

the expanded location model $\mathcal{L}_A^E = \{I, A\}$, then by Proposition 1, D_A^H is unbiased. This is the motivation for the MH method. Now the MH method is not free from bias altogether since it can coincide with the H method in some cases.

By Proposition 1, if $\mathcal{L}_A^E$ is closed with $\dim(\mathcal{L}_k^E) \leq N/2$ and $\mathcal{D} \subset \mathcal{L}_j^E$, then D_j^{MH} is unbiased. The following characterization of other situations where D_j^{MH} is unbiased is also given in [8].

Proposition 2. Let $\mathcal{D}$ be the set of active dispersion effects. If $j \notin \overline{\mathcal{D}}_{-j}$, then D_j^{MH} is unbiased.

Brenneman and Nair [8] also show that the bias diminishes as $N \rightarrow \infty$ for both the H and MH methods, so the methods are not structurally flawed. They also obtain a lower bound for the variance of D_j^{MH} under normality as $V(D_j^{\text{MH}}) \geq V(\log \chi_1^2)/N \approx 4.93/N$. This can then be used to obtain liberal tests for active effects in the model-selection strategy.

Properties of Iterative Method (ML and REML)

Through simulation studies, it is shown in [8] that the bias present in the H and MH methods can be reduced through ML and REML. As a simple example, we can revisit the case when $\mathcal{L} = \{I, A\}$, $\mathcal{D} = \{I, A, B, AB = C\}$ and we want to determine if factor A has a dispersion effect. Recall that Figure 1(a) shows the distribution of D_A^H ($= D_A^{\text{MH}}$, in this case) when $\phi_A = 1$, $\phi_B = -1.67$, and $\phi_{AB} = 1.67$. Since the center of the distribution of D_A^H is at 0, we would in most instances miss this effect. However, if we keep A in the dispersion model along with B and AB then we see in Figure 1(b) that the distribution of $\hat{\phi}_A$ from ML (similarly for REML) is centered at 1. Thus by placing interactions of identified dispersion effects from the MH method in the initial dispersion model and then performing an iterative estimation method on this model, one can correct for the bias of the MH method.

Modeling Strategy

The H method would be the most natural approach under the assumption of a log-linear model for dispersion. However, since one first has to estimate and remove location effects, this can lead to biased dispersion estimates. Part of the bias problem is due to the lack of balance in the location model in relation to the

dispersion effect being estimated. The lack of balance is addressed by the MH method which expands the location model so that the dispersion effect of interest, as well as all of its interactions with active location effects, are estimated. This induces the balance in the estimated residuals, as originally observed in [7], and reduces the underestimation phenomenon. A second problem is the biased estimation of ϕ_A when dispersion effects of B and AB are active. This can lead to: (a) the spurious identification of inert effect A or (b) missing the active effect A . In order to take care of these problems, the proposed strategy in [8] is to initially choose a possibly larger model for dispersion effects $\mathcal{D}^{\text{init}}$ and then use an iterative method, REML, to eliminate spurious effects or pick up active effects. Their strategy is outlined below:

1. Fit a location model using ordinary least squares (OLS) and identify the location model $\mathcal{L}$.
2. If $\mathcal{L}$ contains a closed set $\mathcal{L}'$ with $\dim(\mathcal{L}') \geq N/2$, then identification of dispersion effects is *a priori* difficult and additional experimentation is necessary.
3. If not, then use the MH method to estimate dispersion effects and use an appropriate method to determine the active effects.
4. For a dispersion effect j that is identified as inert, include j in $\mathcal{D}^{\text{init}}$ (the initially selected dispersion model) if there are active effects d_1 and d_2 such that $j = d_1 \circ d_2$.
5. Using $\mathcal{L}$ and $\mathcal{D}^{\text{init}}$ as the identified location and dispersion models, estimate the effects for these models by REML. Discard any dispersion effects that are now identified as inert and refit the submodel using REML.

One can use half-normal plots or more formal methods found in [4] to detect active location effects in the initial model-selection stage. For more formal inference in the final stage, one can rely on parametric **bootstrap** techniques.

Example

This section contains a simulated data set from [8] that illustrates the modeling strategy. Consider the data in Table 1 generated from a 2^4 design in which $\mathcal{L} = \{I, A, B\}$ and $\mathcal{D} = \{I, A, C\}$. The true values of the parameters are $\beta_A = 6$, $\beta_B = 5$, $\phi_A = 2$, and $\phi_C = 1.5$. The half-normal plots for

Table 1 Simulated 2^4 design dataset for which models $\mathcal{L} = \{I, A, B\}$ and $\mathcal{D}^{\text{init}} = \{I, A, C, AC\}$; $\beta_A = 6$, $\beta_B = 5$, $\phi_A = 2$, and $\phi_C = 1.5^{(a)}$

A	B	C	D	y
1	1	1	1	62.30
1	1	1	-1	45.42
1	1	-1	1	49.24
1	1	-1	-1	53.65
1	-1	1	1	47.75
1	-1	1	-1	25.78
1	-1	-1	1	39.05
1	-1	-1	-1	42.25
-1	1	1	1	38.59
-1	1	1	-1	39.05
-1	1	-1	1	38.93
-1	1	-1	-1	39.27
-1	-1	1	1	28.83
-1	-1	1	-1	30.00
-1	-1	-1	1	28.85
-1	-1	-1	-1	28.66

^(a) © American Statistical Association, 2001.

identifying location and dispersion effects are given in Figure 2. From the plot of location effects, it is clear that the only active location effects are A and B with estimates of 5.83 and 5.96, respectively. The plot of dispersion effects using the MH method shows that A , C , and AC are initially identified as being active with estimates of 2.78, 2.34, and 1.60, respectively. The MH method's approximate critical

value for $\alpha = 0.05$ is $1.96 \times \sqrt{4.93/16} = 1.09$, so $\mathcal{D}^{\text{init}} = \{I, A, C, AC\}$. Note that if only factors A and C would have been identified as active, we would still add their interaction, factor AC , to the initially identified dispersion model prior to REML.

The next step is to use REML estimation with the identified models $\mathcal{L} = \{I, A, B\}$ and $\mathcal{D}^{\text{init}} = \{I, A, C, AC\}$. The estimates are $\hat{\beta}_A = 6.04$, $\hat{\beta}_B = 5.14$ for location effects and $\hat{\phi}_A = 2.60$, $\hat{\phi}_C = 1.38$, and $\hat{\phi}_{AC} = 0.16$ for dispersion effects. The estimates for dispersion factors A and C did not change much but the estimate for factor AC changed dramatically from 1.60 to 0.16. This indicates that AC may have been picked up spuriously by the MH method during the initial modeling stage.

For more formal inference about AC , one can use parametric bootstrap under normal error structure. The resulting 95% bootstrap **confidence intervals** are: factor $A - (1.69, 3.71)$, factor $C - (-0.01, 2.52)$, factor $AC - (-1.00, 1.55)$. This provides evidence that AC was spuriously picked up by the MH method. After AC is removed, the REML estimates are practically identical to the initial estimates while the final 95% bootstrap confidence intervals are factor $A - (1.71, 3.66)$ and factor $C - (0.47, 2.52)$. The final conclusion is that factors A and C are significant dispersion effects while factor AC is not.

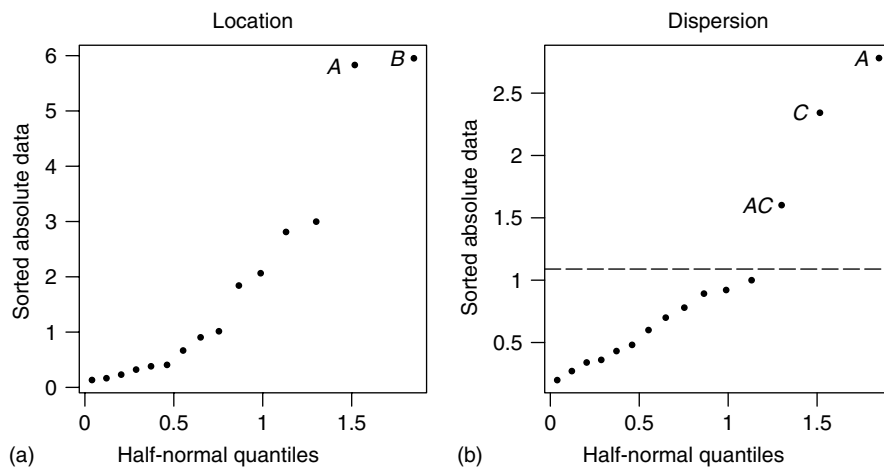


Figure 2 Half-normal plots of location and dispersion effects for simulated dataset in which $\mathcal{L} = \{I, A, B\}$ and $\mathcal{D}^{\text{init}} = \{I, A, C, AC\}$; $\beta_A = 6$, $\beta_B = 5$, $\phi_A = 2$, and $\phi_C = 1.5$. The dashed line in the plot of dispersion-effect estimates is the approximate critical value for the MH method [© American Statistical Association, 2001.]

Weaknesses of Technique

There are two main weaknesses to the strategy of identification and estimation of dispersion effects presented here. The first weakness is that the MH method requires fitting a new location model for each dispersion-effect estimate. For example, in the $2^4 = 16$ -run experiment in the previous section, we had to fit 15 different location models in order to get the 15 MH estimates. This can be very time consuming. The second weakness is that there is a problem when one of the residuals turns out to be identically 0, where $\log(0)$ is undefined, or close enough to 0 for the computer to have trouble taking the logarithm. An *ad hoc* fix to this problem is provided in [20] where they propose replacing $r_i = 0$ with $0.5 \times \min\{|r_j| : j \text{ with } r_j \neq 0\}$. While this provides a way to get MH estimates, the properties of this type of substitution have not been studied. And finally, a third weakness, related to the second weakness, is when a residual that is exactly 0 is computed as a very small "machine zero" number. In this case the program would not have trouble computing the logarithm but the user would not be aware that a problem has occurred.

Conclusion

All noniterative methods (e.g., [5–7]) exhibit biases that depend on the fitted location model and active dispersion effects (parameters other than the one being estimated) and therefore can be large. This bias can lead to spurious identification of effects as well as underestimation or overestimation of real effects. The two-step approach proposed in [8], which incorporates the noniterative MH method identification step followed by an iterative REML final estimation step, eliminates some of the bias present in the noniterative methods.

References

- [1] Bartlett, M.S. & Kendall, I.A. (1946). The statistical analysis of variance-heterogeneity and the logarithmic transformation, *Journal of the Royal Statistical Society. Series B* **8**, 128–138.
- [2] Box, G.E.P. (1988). Signal-to-noise ratios, performance criteria, and transformations, *Technometrics* **30**, 1–40.
- [3] Nair, V.N. & Pregibon, D. (1988). Analyzing dispersion effects from replicated factorial experiments, *Technometrics* **30**, 247–257.
- [4] Hamada, M. & Balakrishnan, N. (1998). Analyzing unreplicated factorial experiments: a review with some new proposals, *Statistica Sinica* **8**, 1–41.
- [5] Box, G.E.P. & Meyer, R.D. (1986). Dispersion effects from fractional designs, *Technometrics* **28**, 19–27.
- [6] Wang, P.C. (1989). Tests for dispersion effects from orthogonal arrays, *Computational Statistics and Data Analysis* **8**, 109–117.
- [7] Bergman, B. & Hyñen, A. (1997). Dispersion effects from unreplicated designs in the 2^{k-p} series, *Technometrics* **39**, 191–198.
- [8] Brenneman, W.A. & Nair, V.N. (2001). Methods for identifying dispersion effects in unreplicated factorial experiments, *Technometrics* **43**, 388–405.
- [9] Pan, G. (1999). The impact of unidentified location effects on dispersion-effects identification from unreplicated factorial designs, *Technometrics* **41**, 313–326.
- [10] McGrath, P. & Lin, K.J. (2001). Confounding of location and dispersion effects in unreplicated fractional factorials, *Journal of Quality Technology* **33**, 129–139.
- [11] Harvey, A.C. (1976). Estimating regression models with multiplicative heteroscedasticity, *Econometrica* **44**, 461–465.
- [12] Rao, C.R. (1970). Estimation of heteroscedastic variances in linear models, *Journal of the American Statistical Association* **65**, 161–172.
- [13] Froehlich, B.R. (1973). Some estimators for a random coefficient regression model, *Journal of the American Statistical Association* **68**, 329–335.
- [14] Davidian, M. & Carroll, R.J. (1987). Variance function estimation, *Journal of the American Statistical Association* **82**, 1079–1091.
- [15] McCullagh, P. & Nelder, J.A. (1989). *Generalized Linear Models*, 2nd Edition, Chapman & Hall, London.
- [16] Nelder, J.A. & Lee, Y. (1991). Generalized linear models for the analysis of taguchi-type experiments, *Applied Stochastic Models and Data Analysis* **7**, 107–120.
- [17] Nelder, J.A. & Lee, Y. (1998). Joint modeling of mean and dispersion, *Technometrics* **40**, 168–175.
- [18] Engel, J. & Huele, A.F. (1996). A generalized linear modeling approach to robust design, *Technometrics* **38**, 365–373.
- [19] Lee, Y. & Nelder, J.A. (1998). Generalized linear models for analysis of quality-improvement experiments, *The Canadian Journal of Statistics* **26**, 95–105.
- [20] Ferrer, A.J. & Romero, R. (1993). Small samples estimation of dispersion effects from unreplicated data, *Communications in Statistics, Part B: Simulation and Computation* **22**, 975–995.

Related Articles

Bergman–Hynén Method for Dispersion Effects; Box–Meyer Method for Dispersion Effects; Factorial Experiments; Fractional Factorial Designs;

Generalized Linear Models; Half-Normal Plot; Lenth's Method for the Analysis of Unreplicated Experiments.

WILLIAM A. BRENNEMAN

Ridge Analysis in Experimental Design

Fitting a Response Surface to Data

Suppose an experimenter wishes to examine several predictor variables $x_1, x_2, \dots, x_k$ in order to determine their effects, individual and combined, on a response variable y . A standard way of proceeding is to decide on n (say) specific combinations of levels $(x_{1u}, x_{2u}, \dots, x_{ku})$, $u = 1, 2, \dots, n$, at each of which the value of the response y , y_u for the u th experiment, will be measured. These n experiments, or *runs* as they are usually called, form the *experimental design* or simply the *design*. Further analysis of the experimental results arising from performing the design can consist of fitting a model, often of polynomial form, to the data, using the least-squares method of regression analysis, and then assessing what the computations reveal. The fitted model could be first order (planar), for example, $\hat{y} = b_0 + b_1x_1 + b_2x_2 + \dots + b_kx_k$, where each of the x 's enters only to first order (power 1); or could be second order (quadratic) of the form (1), where all possible second-order terms (power 2, squares and cross-products) are included.

$$\begin{aligned} \hat{y} = & b_0 + b_1x_1 + b_2x_2 + \dots + b_kx_k \\ & + b_{11}x_1^2 + b_{22}x_2^2 + \dots + b_{kk}x_k^2 + b_{12}x_1x_2 \\ & + b_{13}x_1x_3 + \dots + b_{k-1,k}x_{k-1}x_k \end{aligned} \quad (1)$$

The *hat* (caret) over the y indicates the *predicted value* of y . The b 's are numbers, estimates obtained via the least-squares method, of the *true but unknown* regression coefficients. We assume here that appropriate statistical tests have been carried out to determine that the chosen fitted model is *adequate* to express the main features of the response data and does not suffer from *lack of fit*.

If a second-order surface seems useful and explains the data well, further analysis could take several forms such as viewing two- or three-dimensional surface or sectional plots, performing canonical reduction (a technique that determines the principal features of the surface), or performing *ridge analysis*, the topic of this article. Because ridge analysis does not require one to be able to view the fitted regression surface as a whole, it may be applied even when it is difficult to draw or visualize the surface, for

example, in four or more dimensions of the predictor variable space. The starting point can be *any* selected point in the experimental space; we shall designate it by $\mathbf{f} = (f_1, f_2, \dots, f_k)'$, and call it the *focus*. In many cases, the focus might simply be the coded origin $\mathbf{f} = \mathbf{0}$, the center of the experimental region. Higher, lower, or intermediate local optima values of the response can be specified as desirable for the ridge analysis. The calculation of the coordinates of a ridge trace path also helps one to assess the typically complex interplay between the input variables as the response improves. Use of this tool can supplement, or replace, the canonical reduction of a second-order surface. We explain below how the various paths can be distinguished.

Ridge Analysis for Examining Second-Order Fitted Models, Unrestricted Case, General Focus $\mathbf{f}$

Suppose, in the x space where we have fitted a second-order response function, we (mentally) construct a sphere of radius R about the selected focus $\mathbf{f}$. Somewhere on that sphere there will be a point with the highest response value. As the radius R is changed, all such highest response points will mark out a locus, or path, moving outward from the focus. Similarly, there will also be a path of lowest response. If we extend this idea to points on the spheres where the response is stationary, that is, a minimum or maximum *locally but NOT globally*, more paths will come to light. In fact, there are always $2k$ such paths. Usually, but not always, only the paths that provide the maximum response and the minimum response are likely to be interesting to the experimenter. For that reason, these two paths tend to receive the most emphasis. However, secondary paths could well be of interest in certain practical situations; for example, they might provide good, though not optimum, response values at lower cost or with fewer impurities. As we shall see, we shall always know which of the $2k$ paths we are on.

The basic ridge analysis method is as follows. The fitted second-order surface of equation (1) can be written in matrix form as

$$\hat{y} = b_0 + \mathbf{x}'\mathbf{b} + \mathbf{x}'\mathbf{B}\mathbf{x} \quad (2)$$

2 Ridge Analysis in Experimental Design

where $\mathbf{x}' = (x_1, x_2, \dots, x_k)$, $\mathbf{b}' = (b_1, b_2, \dots, b_k)$, and where

$$\mathbf{B} = \begin{bmatrix} b_{11} & \frac{1}{2}b_{12} & \dots & \frac{1}{2}b_{1k} \\ \dots & b_{22} & \dots & \frac{1}{2}b_{2k} \\ \dots & \dots & \dots & \dots \\ \text{sym} & \dots & \dots & b_{kk} \end{bmatrix} \quad (3)$$

is a symmetric matrix containing all second-order coefficients. Note the off-diagonal halves; it is easy to forget them! We seek the stationary (or “turning”) values of equation (2) subject to being on a sphere, radius R , centered at the focus $\mathbf{f} = (f_1, f_2, \dots, f_k)'$, with equation

$$(\mathbf{x} - \mathbf{f})'(\mathbf{x} - \mathbf{f}) = R^2 \quad (4)$$

We consider the Lagrangian function

$$G = b_0 + \mathbf{x}'\mathbf{b} + \mathbf{x}'\mathbf{B}\mathbf{x} - \lambda\{(\mathbf{x} - \mathbf{f})'(\mathbf{x} - \mathbf{f}) - R^2\} \quad (5)$$

where λ is a Lagrangian multiplier. (For the method of “Lagrange’s undetermined multipliers”, see [1], pp. 231–233.) Differentiation with respect to $\mathbf{x}'$ leads to

$$\frac{\partial G}{\partial \mathbf{x}'} = \mathbf{b} + 2\mathbf{B}\mathbf{x} - 2\lambda(\mathbf{x} - \mathbf{f}) \quad (6)$$

and setting equation (6) equal to a zero vector implies that

$$2(\mathbf{B} - \lambda\mathbf{I})\mathbf{x} = -\mathbf{b} - 2\lambda\mathbf{f} \quad (7)$$

For many given values of λ (the specific choices will be discussed below) we can write a solution for $\mathbf{x}$ as

$$\mathbf{x} = -\frac{1}{2}(\mathbf{B} - \lambda\mathbf{I})^{-1}(\mathbf{b} + 2\lambda\mathbf{f}) \quad (8)$$

This leads us into the following solution sequence:

1. Choose a value of λ appropriate for the desired path (to be explained below).
2. Obtain $\mathbf{x}$ from equation (8).
3. Using that $\mathbf{x}$
 - (a) Evaluate $R^2 = (\mathbf{x} - \mathbf{f})'(\mathbf{x} - \mathbf{f})$ and so obtain R .
 - (b) Evaluate $\hat{y} = b_0 + \mathbf{x}'\mathbf{b} + \mathbf{x}'\mathbf{B}\mathbf{x}$.
 - (c) Tabulate $\mathbf{x}$, R , and $\hat{y}$.

Then the point $\mathbf{x}$ will be on the desired path of stationary values and will lie on a sphere of radius R centered at $\mathbf{f}$.

Distinguishing between the Various Paths

There are, theoretically, $2k$ possible stationary paths. How do we know, when we pick a λ value, on which of these paths we have determined a point? The key lies in the k eigenvalues of $\mathbf{B}$, namely, the values that result from solving $\det(\mathbf{B} - \lambda\mathbf{I}) = 0$, a routine computer calculation. These eigenvalues of $\mathbf{B}$ will be denoted here by $\mu_1, \mu_2, \dots, \mu_k$. In using this notation, we shall assume that they are ordered from most negative to most positive, that is, in ascending order with regard to sign. Words like *lower* and *higher*, or *up* and *down*, are always used in that sense here.

The theory in [2] tells us that, if we select values of λ from $+\infty$ downwards as far as the highest eigenvalue μ_k , we shall always be on the “maximum $\hat{y}$ ” path. Similarly, values of λ from $-\infty$ upwards to the lowest eigenvalue μ_1 put us on the “minimum $\hat{y}$ ” path. Intermediate paths lie in the ranges of λ between the eigenvalues of $\mathbf{B}$. When λ is $+\infty$ or $-\infty$, the value of R is zero and we are at the focus $\mathbf{f}$. As we move λ back down from $+\infty$ to the largest eigenvalue μ_k of $\mathbf{B}$, R rises to ∞ . As we move λ forward from $-\infty$ to the smallest eigenvalue μ_1 of $\mathbf{B}$, R again rises to ∞ . In between each successive pair of eigenvalues, R falls from ∞ to some value above zero and rises again to ∞ , each “loop” corresponding to *two* paths of stationary values of $\hat{y}$. This means that there are $2k$ ridge paths in all, some of which may not appear at the low R values that are usually of most practical interest. So,

1. choosing $\lambda > \mu_k$ provides a locus of maximum $\hat{y}$ as R changes;
2. choosing $\lambda < \mu_1$ provides a locus of minimum $\hat{y}$ as R changes; and
3. choosing $\mu_1 < \lambda < \mu_k$ gives intermediate stationary values.

Note that we do not actually need the eigenvalues of $\mathbf{B}$ to obtain the paths, but only to distinguish between paths. For the loci of the maximum $\hat{y}$ and the minimum $\hat{y}$ which are the most common applications, the eigenvalues are really not needed, because choosing λ values decreasing from ∞ gives the path of maximum $\hat{y}$ while using values increasing from $-\infty$ gives the path of minimum $\hat{y}$. However, obtaining the actual eigenvalues is recommended, simply to avoid the possibility of mistakenly getting onto an unwanted path. When $\lambda = 0$, equation (8)

becomes $\mathbf{x} = -\frac{1}{2}\mathbf{B}^{-1}\mathbf{b}$, which defines the overall stationary point of the second-order surface (e.g., the center of the system if the contours are all ellipsoidal ones). Note that this stationary point could be on *any* of the paths, depending on where zero fits into the listing of the eigenvalues of $\mathbf{B}$.

Applying the Above Results to a First-Order Model as a Special Case

The first-order model $\hat{y} = b_0 + b_1x_1 + b_2x_2 + \cdots + b_qx_k$ occurs when $\mathbf{B} = \mathbf{0}$. The “eigenvalues of $\mathbf{B}$ ” are thus all zero, and there are only two segments for λ , $-\infty$ to 0, and 0 to $+\infty$. The solution for $\mathbf{x}$ reduces to

$$\mathbf{x} = \mathbf{f} + (2\lambda)^{-1}\mathbf{b} \quad (9)$$

Choosing λ in the range $[0, \infty]$, that is, positive, gives the straight-line steepest-ascent direction, and choosing λ negative in the range $[-\infty, 0]$ gives the straight-line steepest-descent direction. Note that $\mathbf{x} - \mathbf{f}$ is proportional to $\mathbf{b}$, as it should be.

Behavior of the Ridge Paths in General

The $\hat{y}$ values on a ridge path will always exhibit one of the following four behaviors. The response will

1. decrease monotonically, or
2. increase monotonically, or
3. pass through a maximum and then decrease monotonically, or
4. pass through a minimum and then increase monotonically.

In cases 3 and 4, the path changes slope as it passes through the stationary point of the second-order **response surface** at the value $\lambda = 0$. The path that contains the stationary point is determined by the position of $\lambda = 0$ in relation to the eigenvalues of $\mathbf{B}$, as mentioned above.

Worked Example for a Three-Predictor Response Surface

A specific example of a second-order fit was discussed elsewhere in this encyclopedia (*see* **Central**

Composite Designs). The surface was fitted to 16 observations and took the form

$$\begin{aligned} \hat{y} = & 57.31 + 1.50x_1 - 2.13x_2 + 1.81x_3 - 4.69x_1^2 \\ & - 6.27x_2^2 - 5.21x_3^2 - 7.13x_1x_2 \\ & - 3.27x_1x_3 - 2.73x_2x_3 \end{aligned} \quad (10)$$

A **blocking** coefficient ($1.29x_B$) also appeared in the original equation, but it has been omitted here because it plays no role in the ridge analysis to follow. In terms of the notation used earlier, $b_0 = 57.31$,

$$\begin{aligned} \mathbf{b} &= \begin{bmatrix} 1.50 \\ -2.13 \\ 1.81 \end{bmatrix}, \text{ and} \\ \mathbf{B} &= \begin{bmatrix} -4.69 & -3.565 & -1.635 \\ -3.565 & -6.27 & -1.365 \\ -1.635 & -1.365 & -5.21 \end{bmatrix} \end{aligned} \quad (11)$$

The eigenvalues of $\mathbf{B}$ are -10.04 , -4.37 , and -1.77 , all negative. Thus, to follow the path of maximum response for example, we substitute values of λ between $+\infty$ and -1.77 in equation (8). (To approximate ∞ in computations, we use $\lambda = 1000$ or some other large number. This provides a check on the calculations since the solution should be very close to $\mathbf{x} = \mathbf{f}$, with R very close to zero.) Selected λ values for our example are shown in Table 1 together with the corresponding results of applying the solution sequence below equation (8), using a Minitab subroutine/macro [3, Chapter 12]. We see that the stationary point, defined by setting $\lambda = 0$, where a maximum $\hat{y} = 58.29$ is predicted, lies on this maximum path, because all the eigenvalues of $\mathbf{B}$ are negative, indicating response reduction everywhere *from* the stationary point. The behavior on this path is of the third type listed above.

In Table 1, we followed the path of maximum response. There are six stationary paths in this particular example ($2k$ in general) and we next provide calculations for two more of them. We consider λ values *below* the largest eigenvalue -1.77 . At this eigenvalue, $R = \infty$, so we have to back λ off downward from -1.77 and monitor a smooth path for decreasing R . At some value of λ that is hard to calculate explicitly, R reaches a minimum value, and the ridge path we are following has reached its starting point in the predictor variable space. As λ is decreased further, a new ridge path begins from the shared starting point and R increases again, approaching ∞ as λ declines

4 Ridge Analysis in Experimental Design

Table 1 Ridge trace for the maximum response on widening spheres of radius R about the origin $(0, 0, 0)$

λ	x_1	x_2	x_3	R	$\hat{y}$
∞	0	0	0	0	57.31
100	0.007	-0.010	0.009	0.015	57.36
50	0.015	-0.020	0.016	0.030	57.40
25	0.028	-0.039	0.030	0.057	57.48
12.5	0.053	-0.071	0.052	0.102	57.60
5	0.114	-0.141	0.089	0.203	57.83
1	0.289	-0.314	0.139	0.449	58.19
0	0.460	-0.464	0.151	0.671	58.29
-0.25	0.539	-0.530	0.151	0.771	58.27
-0.50	0.650	-0.621	0.146	0.911	58.18
-0.75	0.816	-0.753	0.134	1.119	57.91
-1.00	1.092	-0.968	0.105	1.463	57.11
-1.10	1.261	-1.097	0.083	1.673	56.42
-1.25	1.637	-1.383	0.029	2.143	54.30

Table 2 Two additional ridge traces on widening spheres of radius R about the origin $(0, 0, 0)$ associated with λ values between the highest and second highest eigenvalues^(a)

λ	x_1	x_2	x_3	R	$\hat{y}$
-2.3	-1.697	0.979	0.805	2.118	45.41
-2.4	-1.444	0.778	0.784	1.818	48.18
-2.5	-1.261	0.628	0.778	1.609	49.93
-2.6	-1.123	0.509	0.784	1.461	51.08
-2.7	-1.017	0.411	0.799	1.357	51.86
-2.8	-0.932	0.327	0.823	1.286	52.38
-2.9	-0.865	0.253	0.855	1.242	52.69
-3.0	-0.810	0.184	0.895	1.221	52.84
-3.04	-0.791	0.157	0.914	1.2194	52.86
-3.04*	-0.791	0.157	0.914	1.2194	52.86
-3.1*	-0.766	-0.118	0.946	1.223	52.83
-3.4*	-0.683	-0.084	1.180	1.366	51.61
-3.6*	-0.666	-0.252	1.452	1.617	48.98
-3.8*	-0.689	-0.501	1.926	2.106	42.22

^(a)Read up from the gap for the second trace. Read down from the gap for the third trace, whose λ values are indicated by asterisks

and approaches the next-lowest eigenvalue (-4.37). Care is needed in monitoring λ to ensure that the two paths are not confused by an inadvertent leap across the minimum R value (see [3, Chapter 12] for further details and a diagram of the behavior of R). Table 2 shows some selected calculations that take R from its 1.2194 lowest value to just above 2. We see that the “bottoming out point” for R occurs at about $\lambda = -3.04$, after which R begins to ascend again. So $\lambda = -3.04$, approximately, provides the

point $\mathbf{x} = (-0.791, 0.157, 0.914)'$ from which adjacent R versus λ curves rise. These two ridge paths begin at this $\mathbf{x}$ location, and *not at the origin*. Both paths have declining responses.

Extensions of Basic Ridge Analysis

Basic ridge analysis, as described above, has some important practical extensions. It is possible to track

a ridge path on *any* single or on *any* set of designated linear restrictions (i.e., restrictive planes in the experimental space) that may be encountered by the initial best path. One such type of linear restriction occurs, for example, when the experiment is conducted in mixture ingredients adding to 1 (or to any other level specified between 0 and 1) (see [3, Chapter 19] for these extensions).

Historical Notes

Ridge analysis was first introduced in the context of general response surface methodology by Hoerl [4–6]. It was further investigated by Draper [2], who proved several results that Hoerl had suggested without proof. A wide-ranging discussion is given by Hoerl [7]. Applications in **mixture experiments** were discussed by Hoerl [8] and by Draper and Pukelsheim [9–11]. For a complete account, see [3].

References

- [1] Draper, N.R. & Smith, H. (1998). *Applied Regression Analysis*, John Wiley & Sons, New York.
- [2] Draper, N.R. (1963). Ridge analysis of response surfaces, *Technometrics* **5**, 469–479.
- [3] Box, G.E.P. & Draper, N.R. (2007). *Response Surfaces, Mixtures and Ridge Analyses*, 2nd Edition, John Wiley & Sons, New York.
- [4] Hoerl, A.E. (1959). Optimum solution of many variables equations, *Chemical Engineering Progress* **55**(11), 69–78.
- [5] Hoerl, A.E. (1962). Application of ridge analysis to regression problems, *Chemical Engineering Progress* **58**, 54–59.
- [6] Hoerl, A.E. (1964). Ridge analysis, *Chemical Engineering Progress Symposium Series* **60**, 67–77.
- [7] Hoerl, R.W. (1987). Ridge analysis 25 years later, *The American Statistician* **39**, 186–192.
- [8] Hoerl, R.W. (1987). The application of ridge techniques to mixture data: ridge analysis, *Technometrics* **29**, 161–172.
- [9] Draper, N.R. & Pukelsheim, F. (2000). Ridge analysis of mixture response surfaces, *Statistics and Probability Letters* **48**, 131–140.
- [10] Draper, N.R. & Pukelsheim, F. (2002). Generalized ridge analysis under linear restrictions, with particular applications to mixture experiments problems, *Technometrics* **44**, 250–259; See also 2003 **45**, 185.
- [11] Draper, N.R. & Pukelsheim, F. (2003). Canonical reduction of second-order fitted models subject to linear restrictions, *Statistics and Probability Letters* **63**, 401–410.

Related Articles

Blocking; Central Composite Designs.

NORMAN R. DRAPER

Risk Processes

Introduction

The term *risk* is being used in many fields and in different connotations. Risk has been used for many years in decision theory as the expected loss associated with a specific action, under a particular state of nature. Google search yields thousands of entries for “risk processes”. In this article, we focus attention on a class of stochastic processes that may lead to certain undesirable, if not catastrophic events, for example, the claims process in a certain insurance portfolio. The claims arrive from different insured persons at random times. The amount of damage is different from one claim to another. The insurance company starts with certain capital, which steadily grows with the incoming premiums. On the other hand, whenever a payment is done for a claim, its monetary reserve is decreased. The risk to the insurance company is that the monetary reserve will drop below zero. This event is called a *ruin*. It is important to be able to compute the “ruin probability”, and the joint distribution of the “ruin” time and the amount of deficit. The reader is referred to the book of Asmussen entitled *Ruin Probability* [1], for a comprehensive presentation of the theory of ruin probabilities under different stochastic models, and a long list of related references. Another important book on this subject of risk processes is that of Rolski *et al.* entitled *Stochastic Processes for Insurance and Finance* [2]. In this article, we present, in addition to ruin probability, the joint distribution of the ruin time and the amount of deficit for the case of a claiming process that follows a **compound Poisson process** (CPP). In the literature, there are many papers on ruin probability and the distribution of ruin time which are based on other types of stochastic processes. In particular, see Kluppelberg *et al.* [3] for general **Lévy** processes.

Another type of risk process discussed here is that of a *cumulative damage process*, which leads to failure of systems. The damage occurs at random times, and the amount of damage is random. When the cumulative damage reaches a given threshold, the system collapses. We derive the life distribution of systems that are subjected to such failure processes, and the associated hazard functions. The cumulative **damage processes** considered are of the *compound renewal* type.

In the next section, we present the compound **renewal processes** (CRPs) and associated distributions. Then, we study the distributions of stopping times with linear boundaries and CPPs. Some results about the ruin probability and the distributions of its time and its deficit, when the process of claims is a CPP are presented after this. Finally, we discuss the distribution of failure times due to cumulative damage.

Compound Renewal Processes

Let $0 < \tau_1 < \tau_2 < \dots$ be a random sequence of “arrival” times of certain events. Let $T_n = \tau_n - \tau_{n-1}$, $n \geq 1$, where $\tau_0 \equiv 0$. If the random variables $\{T_n, n \geq 1\}$ are independent and identically distributed (i.i.d.), we say that $\{\tau_n, n \geq 0\}$ is a *renewal process*. Let $N(t)$, $t \geq 0$, denote the number of arrivals in the interval $(0, t]$, that is,

$$N(t) = \max(n \geq 0 : \tau_n \leq t) \quad (1)$$

The stochastic process $\{N(t), t \geq 0\}$ is a counting process associated with the renewal process.

Let $F(t)$ be the **cdf** of T_1 . One can immediately show (see [4, pp. 99]) that

$$P\{N(t) = n\} = F^{(n)}(t) - F^{(n+1)}(t), \quad n \geq 0 \quad (2)$$

where $F^{(n)}(t)$ denotes the n -fold convolution of F , that is, $F^{(0)}(t) = 1$ for all $t \geq 0$, $F^{(1)}(t) = F(t)$, and for $n \geq 2$,

$$F^{(n)}(t) = \int_0^t F^{(n-1)}(t-x) dF(x) \quad (3)$$

Let $\{X_n, n \geq 1\}$ be a sequence of i.i.d. random variables, independent of $\{\tau_n, n \geq 0\}$. The process

$$Y(t) = \sum_{n=0}^{N(t)} X_n, \quad t \geq 0 \quad (4)$$

where $X_0 \equiv 0$, is called a *compound renewal process*. In this article we are interested in CRPs with positive jumps. Let $G(x)$ denote the cdf of X_1 . We assume that $G(0) = 0$ and that G is absolutely continuous, with a density g . The cdf of $Y(t)$, at time t , is $H(x; t) = P\{Y(t) \leq x\}$. Notice that $H(0; t) = \bar{F}(t) = 1 - F(t)$ and for $0 < x < \infty$, $H(x; t) = \bar{F}(t) + \int_0^x h(y; t) dy$,

2 Risk Processes

where $h(y; t)$ is the density of $H(\cdot; t)$ on $(0, \infty)$, given by

$$h(y; t) = \sum_{n=1}^{\infty} P\{N(t) = n\} g^{(n)}(y) \quad (5)$$

and $g^{(n)}(y)$ is the n -fold convolution of g . An important subclass of CRP is the CPP, where $\{N(t), t \geq 0\}$ is an ordinary *Poisson process*. In this case,

$$P\{N(t) = n\} = e^{-\lambda t} \frac{(\lambda t)^n}{n!}, \quad n = 0, 1, \dots \quad (6)$$

is the probability function of the Poisson distribution, with mean λt . The parameter $\lambda = E\{N(1)\}$ is called the *intensity of the arrival process*. In this case, the interarrival times $\{T_n, n \geq 1\}$ are exponentially distributed with $E\{T_1\} = 1/\lambda$. We denote by $p(n; \mu)$ and $P(n; \mu)$ the **pdf** and cdf of the Poisson distribution with mean μ . Thus, in the CPP case, the density of $Y(t)$, on $(0, \infty)$ is

$$h(y; t) = \sum_{n=1}^{\infty} p(n; \lambda t) g^{(n)}(y) \quad (7)$$

The corresponding cdf is

$$H(x; t) = \sum_{n=0}^{\infty} p(n; \lambda t) G^{(n)}(x) \quad (8)$$

where $G^{(n)}$ is the n -fold convolution of G . **Moments** of $Y(t)$ can be obtained from the Laplace–Stieltjes transform (LST)

$$\begin{aligned} E\{e^{-\theta Y(t)}\} &= e^{-\lambda t} + \int_0^{\infty} e^{-\theta x} h(x; t) dx \\ &= \exp\{-\lambda t(1 - \phi(\theta))\} \end{aligned} \quad (9)$$

where $\phi(\theta)$ is the Laplace transform (LT) of $g(x)$. Equation (9) is applied to define the Wald martingale associated with a CPP.

Distributions of Stopping Times for CPP

Let $B_L(t, \beta_1)$ and $B_U(t, \beta_2)$, for $0 < \beta_1, \beta_2 < \infty$ be the linear boundaries

$$B_L(t; \beta_1) = -\beta_1 + t \quad (10)$$

and

$$B_U(t; \beta_2) = \beta_2 + t \quad (11)$$

For a CPP $Y(t)$, we define the stopping times

$$T_L(\beta_1) = \inf\{t > 0 : Y(t) = -\beta_1 + t\} \quad (12)$$

and

$$T_U(\beta_2) = \inf\{t > 0 : Y(t) \geq \beta_2 + t\} \quad (13)$$

Notice that $P\{T_L(\beta_1) \geq \beta_1\} = 1$ and $P\{T_L(\beta_1) = \beta_1\} = e^{-\lambda\beta_1}$. As proved by Borovkov [5, p. 66], the density of $T_L(\beta_1)$ on (β_1, ∞) is

$$\psi_L(t; \beta_1) = \frac{\beta_1}{t} h(t - \beta_1; t), \quad t > \beta_1 \quad (14)$$

Let $\rho = \lambda E\{X_1\}$. If $\rho < 1$, then $P\{T_L(\beta_1) < \infty\} = 1$. Since $Y(T_L(\beta_1)) = -\beta_1 + T_L(\beta_1)$, we obtain from the Wald martingale the identity

$$\begin{aligned} E\{\exp\{-\theta Y(T_L(\beta_1)) + \lambda(1 - \phi(\theta))T_L(\beta_1)\}\} \\ = E\{e^{\theta\beta_1} \exp\{-T_L(\beta_1)(\theta - \lambda(1 - \phi(\theta)))\}\} = 1 \end{aligned} \quad (15)$$

for all θ in the domain of convergence of $\phi(\theta)$. Define $\omega(\theta) = \theta - \lambda(1 - \phi(\theta))$. Then, from equation (15), we get

$$E\{e^{-\omega T_L(\beta_1)}\} = e^{-\beta_1 \theta(\omega)} \quad (16)$$

where $\theta(\omega)$ is the inverse function of $\omega(\theta)$. Moments of $T_L(\beta_1)$ can be obtained from equation (16). Indeed,

$$\begin{aligned} E\{T_L(\beta_1)\} &= -\left\{ \frac{d}{d\omega} e^{-\beta_1 \theta(\omega)} \Big|_{\omega=0} \right\} \\ &= \frac{\beta_1}{1 + \lambda\phi'(0)} = \frac{\beta_1}{1 - \rho} \end{aligned} \quad (17)$$

Similarly,

$$E\{T_L^2(\beta_1)\} = \frac{\beta_1^2}{(1 - \rho)^2} + \frac{\beta_1 \lambda E\{X_1^2\}}{(1 - \rho^3)} \quad (18)$$

Moreover, since

$$\begin{aligned} E\{T_L(\beta_1)\} &= \beta_1 e^{-\lambda\beta_1} + \int_0^{\infty} t \psi_L(t; \beta_1) dt \\ &= \beta_1 e^{-\lambda\beta_1} + \beta_1 \int_{\beta_1}^{\infty} h(t - \beta_1; t) dt \\ &= \beta_1 (e^{-\lambda\beta_1} + \int_0^{\infty} h(t; t + \beta_1) dt) \end{aligned} \quad (19)$$

we obtain from equations (17) and (19) the identity

$$\int_0^\infty h(t; t + \beta_1) dt = \frac{1}{1 - \rho} - e^{-\lambda\beta_1} \quad (20)$$

which holds for all distributions G such that $\rho = \lambda E(X_1) = \lambda \int_0^\infty x dG(x) < 1$. In particular for $\beta_1 = 0$, we obtain when $\rho < 1$ the identity

$$\int_0^\infty h(t; t) dt = \frac{1}{1 - \rho} - 1 = \frac{\rho}{1 - \rho} \quad (21)$$

Stadje and Zacks [6] derived the following formula for the (defective) density of $T_U(\beta_2)$, namely

$$\begin{aligned} \psi_U(t; \beta_2) = & \lambda e^{-\lambda t} \bar{G}(t + \beta_2) + \lambda \int_0^{\beta_2+t} h(y; t) \bar{G}(\beta_2 \\ & + t - y) dy \\ & - \lambda \int_0^t h(y + \beta_2; y) \cdot e^{-\lambda(t-y)} \bar{G}(t - y) dy \\ & - \lambda \int_0^t \frac{1}{u} h(t - u + \beta_2; t - u) \\ & \times \left(\int_0^u z \bar{G}(z) h(u - z; u) dz \right) du \end{aligned} \quad (22)$$

This density assumes in some special cases more explicit formulae. For example, if $G(x) = 1 - e^{-\mu x}$, $x \geq 0$, then

$$h(y; t) = \mu \sum_{n=1}^\infty p(n; \lambda t) p(n - 1; \mu y) \quad (23)$$

Moreover, in this exponential case, if $\beta_2 = 0$,

$$\begin{aligned} \psi_U(t; 0) &= \lambda \sum_{n=0}^\infty \frac{1}{n+1} p(n; \lambda t) p(n; \mu t) \\ &= \lambda e^{-(\lambda+\mu)t} \sum_{n=0}^\infty \frac{(\lambda\mu t^2)^n}{(n+1)!n!} \end{aligned} \quad (24)$$

As expected,

$$\begin{aligned} \int_0^\infty \psi_U(t; 0) dt &= \frac{\lambda}{\lambda + \mu} \sum_{n=0}^\infty \frac{(2n)!}{(n+1)!n!} \frac{(\lambda\mu)^n}{(\lambda + \mu)^{2n}} \\ &= \rho \end{aligned} \quad (25)$$

To verify equation (25), use MAPLE, and recall that $\rho = \lambda/\mu$.

Ruin Probability and Associated Distribution under CPP

Consider an insurance company that starts with an initial capital $[\$]\beta$ for a given portfolio. Suppose that the premium paid by the customers is $[\$]p$ per time unit. Thus, after t units of time, the capital reserve, if there were no claims, is $\beta + pt$. Suppose now that claims for damage arrive according to a CPP $Y(t) = \sum_{n=0}^{N(t)} X_n$. The capital reserve at time t is then

$$R(t) = \beta + tp - Y(t), \quad t \geq 0 \quad (26)$$

Ruin occurs at the first time in which $R(t) \leq 0$. This is the stopping time

$$T(\beta, p) = \inf\{t > 0 : Y(t) \geq \beta + tp\} \quad (27)$$

Making the time transformation $t' = tp$, and changing λ to $\lambda' = \lambda/p$, the stopping time $T(\beta, p)$ is reduced to the stopping time $T'_U(\beta) = \inf\{t' > 0 : Y'(t') \geq \beta + t'\}$, where $Y'(t')$ is a CPP with intensity parameter $\lambda' = \lambda/p$. Accordingly, we assume, without loss of generality, that $p = 1$, and apply the results obtained in the section titled “Distributions of Stopping Times for CPP”. The *ruin probability* is defined as

$$\begin{aligned} q(\beta) &= P\{T_U(\beta) < \infty\} \\ &= \int_0^\infty \psi_U(t; \beta) dt \end{aligned} \quad (28)$$

If $\rho \geq 1$ then $q(\beta) = 1$. We are interested in formulae for $q(\beta)$, when $\rho < 1$. Notice that if $p \neq 1$, the condition is $\rho < p$. Let $q^*(s) = \int_0^\infty e^{-s\beta} q(\beta) d\beta$ denote the LT of $q(\beta)$. The formula of $q^*(s)$ is (see [1, p. 65], [4, p. 24])

$$q^*(s) = \frac{1}{s} \cdot \frac{\rho s - \lambda(1 - \phi(s))}{s - \lambda(1 - \phi(s))} \quad (29)$$

where $\phi(s)$ is the LT of the density g . By inversion of $q^*(s)$ one can obtain $q(\beta)$. Accordingly, if $g(x) = \mu e^{-\mu x}$, then $\phi(s) = \mu/(\mu + s)$. Substituting this in equation (29) and inverting yields for the exponential case

$$q(\beta) = \rho \exp\{-\mu(1 - \rho)\beta\} \quad (30)$$

Notice that $q(0) = \rho$, as obtained in equation (25). If G is the Erlang $(2, \mu)$ distribution, $g(x) =$

$\mu^2 x e^{-\mu x}$, $\phi(s) = \mu^2/(\mu + s)^2$, and

$$q(\beta) = \rho e^{-\mu(1-\rho/4)\beta} \left[\cosh\left(\frac{\beta}{2}(\lambda(\lambda + 4\mu))^{1/2}\right) + \frac{(\lambda + \mu) \sinh\left(\frac{\beta}{2}(\lambda(\lambda + 4\mu))^{1/2}\right)}{(\lambda(\lambda + 4\mu))^{1/2}} \right] \quad (31)$$

The exact formula of $q(\beta)$ when G is Erlang $(3, \mu)$ is very complicated. It involves the roots of a cubic polynomial. We therefore develop an alternative formula for $q(\beta)$.

Theorem For $0 < \rho < 1$,

$$q(\beta) = (1 - \rho) \int_0^\infty h(u + \beta; u) du \quad (32)$$

Proof See Perry *et al.* [7].

We derive an explicit formula of equation (32) for the case where G is an Erlang $(2, \mu)$. Notice that $g^{(n)}(x) = [\mu^{2n}/(2n-1)!]x^{2n-1}e^{-\mu x}$, $n \geq 1$. Accordingly

$$h(u + \beta; u) = \mu \sum_{n=1}^\infty p(n; \lambda u) p(2n-1; \mu(u + \beta)) \quad (33)$$

Hence,

$$\begin{aligned} \int_0^\infty h(u + \beta; u) du &= e^{-\mu\beta} \sum_{n=1}^\infty \frac{\lambda^n \mu^{2n}}{n!(2n-1)!} \int_0^\infty y^n (y + \beta)^{2n-1} \\ &\quad \times e^{-(\lambda+\mu)y} dy \end{aligned} \quad (34)$$

From equations (32) and (34), we obtain the formula

$$\begin{aligned} q(\beta) &= (1 - \rho) e^{-\mu\beta} \sum_{n=1}^\infty \left(\frac{\lambda \mu^2}{(\lambda + \mu)^3} \right)^n \binom{3n-1}{n} \\ &\quad \times \left[1 + \frac{1}{(3n-1)!} \sum_{j=0}^{2n-2} \binom{2n-1}{j} \right. \\ &\quad \left. \times (n+j)! (\beta(\lambda + \mu))^{2n-1-j} \right] \end{aligned} \quad (35)$$

Formulae (31) and (35) yield the same results. For example, for $\beta = 5$ and $\rho = 0.6$ we have $q(5) = 0.0383$. If $\rho = 0.8$ we have $q(5) = 0.20958$. Formulae similar to equation (35) can be derived for the cases of G equal to Erlang $(3, \mu)$, Erlang $(4, \mu)$, and so on, while formulae like equation (31) are unavailable.

We remark that the convergence in equation (35) is slow, and one needs to calculate hundreds of terms. The density of the ruin time, $T_U(\beta)$, when $p = 1$, is given by equation (22). We can write this density as

$$\begin{aligned} \psi_U(t; \beta) &= \lambda e^{-\lambda t} \bar{G}(t + \beta) \\ &\quad + \lambda \int_0^{t+\beta} p_U(x; t, \beta) \bar{G}(t + \beta - x) dx \end{aligned} \quad (36)$$

where $p_U(x; t, \beta)$ is the defective density

$$p_U(x; t, \beta) = \frac{d}{dx} P\{Y(t) \leq x, T_U(\beta) > t\} \quad (37)$$

Stadje and Zacks [6] proved that

$$p_U(x; t, 0) = \frac{t-x}{t} h(x; t), \quad t > x \quad (38)$$

and

$$\begin{aligned} p_U(x; t, \beta) &= h(x; t) - e^{-\lambda(t-x+\beta)} \cdot h(x; x - \beta) \\ &\quad - (t - x + \beta) \int_{t-x+\beta}^t \frac{1}{u} h(u - t \\ &\quad + x - \beta; u) h(t - u + \beta; t - u) du \end{aligned} \quad (39)$$

Substituting equation (39) in equation (36), we obtain equation (22). The *deficit* at ruin time, when $p = 1$, is

$$D_U(\beta) = Y(T_U(\beta)) - \beta - T_U(\beta) \quad (40)$$

This random variable is generally dependent on $T_U(\beta)$. The joint density of $(T_U(\beta), D_U(\beta))$ is

$$\begin{aligned} \tilde{\psi}_U(t, d; \beta) &= \lambda e^{-\lambda t} g(t + \beta + d) \\ &\quad + \lambda \int_0^{t+\beta} p_U(x; t, \beta) g(t + \beta \\ &\quad + d - x) dx \end{aligned} \quad (41)$$

One can immediately verify that, in the exponential case, when $g(x) = \mu e^{-\mu x}$,

$$\tilde{\psi}_U(t, d; \beta) = \psi_U(t; \beta) \cdot \mu e^{-\mu d}, \quad \text{all } 0 < t, d < \infty \quad (42)$$

that is, $T_U(\beta)$ and $D_U(\beta)$ are independent. This is the only case where $T_U(\beta)$ and $D_U(\beta)$ are independent. On this topic, see also Gerber and Shiu [8].

We conclude this section with a few comments. It is often the case that the parameter ρ may increase at some unknown *change point*. In such a case, the ruin probability might increase dramatically. Insurance companies should detect such a change point as soon as possible in order to avoid ruin by changing the premium p . Other methods common these days are based on a bonus–malus system. In such a system the individual premium increases or decreases according to the history of claims. For a Markovian bonus–malus system, see Zacks and Levikson [9].

Failure Times due to Cumulative Damage

Consider a damage process, which, at random times, brings a random amount of damage to a system. We model such a process by a CRP, $Y(t) = \sum_{n=0}^{N(t)} X_n$, where $X_0 \equiv 0$, and $\{X_n, n \geq 1\}$ are i.i.d. positive random variables, representing the damage to the system. The cdf of $Y(t)$ is

$$H(y; t) = \sum_{n=0}^{\infty} (F^{(n)}(t) - F^{(n+1)}(t)) G^{(n)}(y) \quad (43)$$

where F is the cdf of the interarrival time and G is the cdf of X_1 . Let β , $0 < \beta < \infty$, be a known specified threshold of the system. The system fails as soon as $Y(t)$ exceeds β . Thus, let

$$T(\beta) = \inf\{t > 0 : Y(t) \geq \beta\} \quad (44)$$

Since $\{Y(t), t \geq 0\}$ is an upward jump process, it is a nondecreasing function of t . Accordingly

$$P\{T(\beta) > t\} = H(\beta - 0; t) \quad (45)$$

Obviously, $P\{T(\beta) < \infty\} = 1$, and thus $\lim_{t \rightarrow \infty} H(\beta - 0; t) = 0$. The density function of $T(\beta)$ is

$$p_T(t; \beta) = f(t) \bar{G}(\beta) + \sum_{n=2}^{\infty} f^{(n)}(t) (G^{(n-1)}(\beta) - G^{(n)}(\beta)) \quad (46)$$

Let $S_n = \sum_{i=0}^n X_i$, and define

$$N^*(t) = \max\{n \geq 0 : S_n \leq t\}, \quad t \geq 0 \quad (47)$$

$\{N^*(t), t \geq 0\}$ is a renewal process associated with $\{X_i, i \geq 0\}$. Thus, $P\{N^*(\beta) = n - 1\} = G^{(n-1)}(\beta) - G^{(n)}(\beta)$, $n \geq 1$. We thus obtain

$$p_T(t; \beta) = \sum_{n=1}^{\infty} f^{(n)}(t) P\{N^*(\beta) = n - 1\} \quad (48)$$

The following theorem was proved by Zacks [10].

Theorem *If the interarrival time T_1 has a moment of order m , μ_m ($m \geq 1$), then*

$$E\{T^m(\beta)\} = \sum_{n=1}^{\infty} M_{n,m} P\{N^*(\beta) = n - 1\} \quad (49)$$

where

$$M_{n,m} = E \left\{ \left(\sum_{i=1}^n T_i \right)^m \right\} \quad (50)$$

Hence,

$$V\{T(\beta)\} = \sigma^2(1 + E\{N^*(\beta)\}) + \mu_1^2 V\{N^*(\beta)\} \quad (51)$$

where μ_1 is the expected value of T_1 and $\sigma^2 = V\{T_1\}$.

We consider now the special case of CPP with exponential damage, that is, $G(x) = 1 - e^{-\mu x}$, $x \geq 0$. Let $\zeta = \mu\beta$. The reliability function is $R(t; \lambda, \zeta) = P\{T(\beta) > t\}$. One can verify that

$$R(t; \lambda, \zeta) = \sum_{n=0}^{\infty} p(n; \zeta) P(n; \lambda t) \quad (52)$$

where $P(n; \lambda t) = \sum_{j=0}^n p(j; \lambda t)$. Accordingly, the density of the failure time $T(\beta)$ is

$$f_T(t; \lambda, \zeta) = \lambda \sum_{n=0}^{\infty} p(n; \zeta) p(n; \lambda t) \quad (53)$$

Furthermore, in this case

$$E\{T(\beta)\} = \frac{1 + \zeta}{\lambda} \quad (54)$$

and

$$V\{T(\beta)\} = \frac{1 + 2\zeta}{\lambda^2} \quad (55)$$

The hazard rate function in this case is

$$h_z(t; \lambda, \zeta) = \lambda \frac{\sum_{n=0}^{\infty} p(n; \zeta) p(n; \lambda t)}{\sum_{n=0}^{\infty} p(n; \zeta) P(n; \lambda t)} \quad (56)$$

One can show that $t \mapsto h_z(t; \lambda, \zeta)$ is strictly increasing, with $h_z(0; \lambda, \zeta) = \lambda e^{-\zeta}$ and $\lim_{t \rightarrow \infty} h_z(t; \lambda, \zeta) = \lambda$.

Zacks [11] studied another class of failure time distributions caused by nonhomogeneous compound Poisson damage processes. For a nonhomogeneous Poisson process with **intensity function** $\lambda(t)$, if $m(t) = \int_0^t \lambda(s) ds$, the cumulative damage at time t has a cdf

$$H(y; t) = \sum_{n=0}^{\infty} p(n; m(t)) G^{(n)}(y) \quad (57)$$

The special case with $m(t) = (\lambda t)^\nu$, $0 < \nu < \infty$, is called a *compound Weibull process (CWP)*. A CWP with exponential damage distribution G , yields the reliability function

$$R(t; \lambda, \nu, \zeta) = \sum_{n=0}^{\infty} p(n; \zeta) P(n; (\lambda t)^\nu) \quad (58)$$

Equation (52) is a special case for $\nu = 1$. As before, $\zeta = \mu\beta$. The failure time associated with equation (58) has the density function

$$f(t; \lambda, \nu, \zeta) = \lambda \nu (\lambda t)^{\nu-1} \sum_{j=0}^{\infty} p(j; \zeta) p(j; \lambda t)^\nu \quad (59)$$

Notice that

$$\lim_{\zeta \rightarrow 0} f(t; \lambda, \nu, \zeta) = \lambda \nu (\lambda t)^{\nu-1} \cdot e^{-(\lambda t)^\nu} \quad (60)$$

which is the common Weibull distribution. It shows that the Weibull life distribution is the case where the first shock leads to failure (zero threshold).

In the literature there are other stochastic models of damage and wearout. Abdel-Hameed [12] used *gamma processes*. This was done also by Durfresne *et al.* [13]. Faddy [14] used *phase-type* distributions. See also the article of Singpurwalla [15].

References

- [1] Asmussen, S. (2000). *Ruin Probabilities*, World Scientific, Singapore.
- [2] Rolski, T., Schmidli, H., Schmidt, V. & Tiegels, J. (2000). *Stochastic Processes for Insurance and Finance*, Wiley Series in Probability and Statistics, John Wiley & Sons, New York.
- [3] Kluppelberg, C., Kyprianou, A.E. & Maller, R.A. (2004). Ruin probabilities and overshoots for general Levy insurance risk processes, *Annals of Applied Probability* **14**, 1766–1801.
- [4] Kao, E.P.C. (1997). *An Introduction to Stochastic Processes*, Duxbury Press, Belmont.
- [5] Borovkov, A.A. (1976). *Stochastic Processes in Queuing Theory*, Springer, New York.
- [6] Stadje, W. & Zacks, S. (2003). Upper first-exit times of compound Poisson processes revisited, *Probability in the Engineering and Informational Sciences* **17**, 459–465.
- [7] Perry, D., Stadje, W. & Zacks, S. (2002). Hitting and ruin probabilities for compound Poisson processes and the cycle maximum of the M/G/1 queue, *Stochastic Models* **18**, 553–564.
- [8] Gerber, H.U. & Shiu, E.S.W. (1997). The joint distribution of the time to ruin, the surplus immediately before ruin, and the deficit at ruin, *Insurance Mathematics and Economics* **21**, 129–137.
- [9] Zacks, S. & Levikson, B. (2004). Claiming strategies and premium levels for bonus malus systems, *Scandinavian Journal of Actuary* **14**, 14–27.
- [10] Zacks, S. (2006). Failure distributions associated with general compound renewal damage processes, in *Probability, Statistics and Modeling in Public Health*, M., Nikulin, D. Commenge & C. Huber, eds, Springer, pp. 466–475.
- [11] Zacks, S. (2004). Distributions of failure times associated with non-homogeneous compound Poisson damage processes, *A Festschrift for Herman Rubin, IMS Lecture Notes* **45**, 396–407.
- [12] Abdel-Hameed, M. (1975). A gamma wear process, *IEEE Transactions on Reliability* **24**, 152–153.
- [13] Durfresne, F., Gerber, H.U. & Shiu, E.S.W. (1991). Risk theory with the gamma process, *ASTIN Bulletin* **21**, 177–192.
- [14] Faddy, M.J. (1995). Phase-type distributions for failure times, *Mathematical and Computer Modelling* **22**, 63–70.
- [15] Singpurwalla, N.D. (1995). Survival in dynamic environments, *Statistical Science* **10**, 86–103.

Further Reading

Zacks, S. (2005). Some recent results on the distributions of stopping times of compound Poisson processes with linear boundaries, *Journal of Statistical Planning and Inference* **130**, 95–109.

Related Articles

Extremes in; Stochastic Deterioration.

Damage Processes; Degradation Processes; Lévy Processes; Poisson Processes; Risk Analysis,

SHELEMYAHU ZACKS

Brownian Motion

Introduction

Brownian motion is the continuous stochastic process *par excellence*. All continuous martingales are time-changed Brownian motions. All continuous **Markov processes** are constructed from Brownian motions. Many classical problems in partial differential equations are solved by probabilistic methods using Brownian motion. Many stochastic processes can be approximated well by Brownian motions or processes related to them. Brownian motion is central to the theory and application of stochastic processes and statistical analysis.

It is named after the botanist Robert Brown, who observed in 1826 that microscopic particles suspended in a fluid move with jumps and starts as if alive. Physicists use the term *Brownian motion* for this physical phenomenon itself, for which there have been many mathematical models built by Einstein, Perrin, Langevin, Smoluchovsky, Uhlenbeck, Wang, and Ornstein, to name a few. In the mathematical literature, the term is used for the model first introduced by Einstein in 1905, which is also the mathematical model introduced by Bachelier [1] for the evolution of stock prices. Nelson [2] has an enlightening account of this history. We shall use the term as it is customary in mathematics.

A real-valued stochastic process is called a *Brownian motion* if it is continuous, temporally homogeneous, and has independent increments. Continuity of paths is a regularity condition. Temporal homogeneity signals that the underlying mechanism producing the randomness remains unchanged over time. *Independence of increments* refers to the mutual independence of the displacements over disjoint time intervals.

Every such process is necessarily Gaussian. We use the term *Wiener process* for the standard case where the initial state is zero, the mean vanishes, and the variance increases with time at unit rate; we denote it by W . Then, every Brownian motion X has the form

$$X_t = X_0 + \mu t + \sigma W_t, \quad t \geq 0 \quad (1)$$

where μ and σ are deterministic constants and X_0 is independent of the Wiener process W . Here, μ is

called the *drift rate* and σ the *volatility* (and σ^2 is called the *diffusion coefficient* in earlier literature).

Symmetry, Scaling, and Time Reversal

The Wiener process has certain invariance properties that enhance its value. It is symmetric: W and its negative have the same law. It has scale invariance: setting

$$\widehat{W}_t = \frac{1}{\sqrt{c}} W_{ct}, \quad t \geq 0 \quad (2)$$

yields another Wiener process $\widehat{W}$. It has time reversibility:

$$\widetilde{W}_t = t W_{1/t}, \quad t \geq 0 \quad (3)$$

is again a Wiener process.

Scale invariance shows that a Wiener process is strictly stable with index 2; *self-similar* is another term. This property puts it at the center of symmetric stable Lévy processes: every such process is obtained from a Wiener process by subordination to an increasing stable process.

Martingale Connection

Wiener process is the basic continuous martingale. It is itself a martingale, and putting

$$Y_t = W_t^2 - t, \quad t \geq 0 \quad (4)$$

yields another martingale Y . Indeed, if the process W is a continuous martingale and Y is also a martingale, then W is necessarily a Wiener process; this is Lévy's characterization of the Wiener process as a martingale. Another, very useful characterization is *via* the exponential martingale:

$$X_t = \exp\left(r W_t - \frac{1}{2} r^2 t\right), \quad t \geq 0 \quad (5)$$

is a martingale for each real-number r if and only if W is a Wiener process.

Sample Paths and Stochastic Calculus

Some of the fascination with Brownian motion comes from the intricacies of its sample paths.

2 Brownian Motion

Although continuous, almost every path is nowhere differentiable and has infinite total variation over every open interval, but has finite quadratic variation over every interval (the value being equal to the length of the interval); see Freedman [3] for the precise details. Thus, the paths are highly irregular and highly oscillatory, but the oscillations are “small” in magnitude. Hence, the differential dW does not have a meaning in the ordinary sense, but Itô was able to give it a meaning in the sense of convergence in probability. The resulting stochastic calculus is an extension of the ordinary calculus and has proved to be of immense value in applications ranging from engineering to finance. We refer to Karatzas and Shreve [4] and Revuz and Yor [5] on these matters.

Hitting Times

For fixed $\alpha > 0$, let T_α denote the first time that the Wiener process W hits the point α . A probabilistic argument based on the strong Markov property shows that T_α has the same distribution as α^2/Z^2 , where Z is a standard Gaussian variable (normal with mean 0 and variance 1).

It follows that T_α is almost surely finite; the Wiener process will hit α with probability one, however large α maybe. The expected time it takes to hit α is equal to infinity (since α^2/Z^2 has infinite expectation), however small $\alpha > 0$ maybe. It is easy, using the well-known density for Z , to obtain the **probability density function** for T_α : it is

$$\frac{\alpha e^{-\alpha^2/2t}}{\sqrt{2\pi t^3}}, \quad t > 0 \quad (6)$$

Engineers and statisticians call this the *inverse Gaussian distribution*, probabilists have no name for it. It is an infinitely divisible distribution; it is strictly stable with index $1/2$. In fact, as a stochastic process indexed by α , the process (T_α) is an increasing pure-jump Lévy process which is, further, strictly stable with index $1/2$.

Maximum

Let M_t be the maximum achieved by the Wiener process W during the time interval from 0 to t . This defines an increasing continuous process M . Note that, for each time t and point $\alpha \geq 0$, the event that

$T_\alpha < t$ is the same as the event $M_t > \alpha$. Since T_α has the same distribution as α^2/Z^2 , this implies that

$$\begin{aligned} \mathbb{P}\{M_t > \alpha\} &= \mathbb{P}\left\{\frac{\alpha^2}{Z^2} < t\right\} \\ &= \mathbb{P}\left\{\sqrt{t} |Z| > \alpha\right\} = \mathbb{P}\{|W_t| > \alpha\} \end{aligned} \quad (7)$$

In other words, for each t fixed, M_t has the same distribution as $|W_t|$. In particular, then,

$$\mathbb{E}M_t = \sqrt{\frac{2t}{\pi}}, \quad \mathbb{E}M_t^2 = t \quad (8)$$

Since W touches every α with probability one, the maximum process M increases without bound. By symmetry, the minimum process must be decreasing without bound. This implies, in particular, that W crosses and recrosses the point 0 without end, that is, the zeros of W form an unbounded random subset of the positive real line.

Zeros of W

Consider the random set R of times t for which $W_t = 0$. The set R has many of the features of the Cantor set. Almost surely, R is closed, unbounded, and everywhere dense in itself, and also has Lebesgue measure zero, empty interior, and the power of the continuum.

The random set R is regenerative: for every stopping time T belonging to R , its future after T is an independent replica of R . It follows from a theorem of Maisonneuve [6] that there exists an increasing Lévy process $S = (S_\alpha)$ such that the set R is the range of S , that is,

$$R = \{t \geq 0 : S_\alpha = t \text{ for some } \alpha \geq 0\} \quad (9)$$

The process S turns out to have the same probability law as the hitting process (T_α) ; in other words, S is a strictly increasing pure-jump **Lévy process** and is stable with index $1/2$.

Local Times

For each time t , let L_t be the infimum of the set of numbers α for which S_α exceeds t . In short, L is the functional inverse of S . A closer examination of

this definition shows that L is a continuous increasing process that increases when and only when W is at the point 0. Thus, L may be considered to be a random clock whose mechanism is so rigged that the clock advances when and only when W is at 0. For this reason, L is called the *local time process at 0*.

The process L has the same probability law as the maximum process M . This is because L bears the same relation to (S_α) as M does to (T_α) , and the processes (S_α) and (T_α) have the same law. Indeed, the two-dimensional processes $(M, M - W)$ and $(L, |W|)$ have the same probability law. This implies that $L - |W|$ has the same probability law as $M - (M - W) = W$, that is, $L - |W|$ is again a Wiener process.

Markovian Connection

Brownian motion is a special Markov process. We consider the case of the d -dimensional Brownian motion $X = X_0 + W$, where

$$W = (W^{(1)}, \dots, W^{(d)}) \quad (10)$$

is the d -dimensional Wiener process, that is, the components $W^{(i)}$ are independent one-dimensional Wiener processes. Then, the increment over an interval of length t has the d -dimensional Gaussian distribution with mean 0 and variance–**covariance** matrix equal to t times the identity matrix. Thus, X is a Markov process with transition probability density

$$p(t, x, y) = \frac{1}{\sqrt{(2\pi t)^d}} \exp\left(-\frac{1}{2t}|y - x|^2\right) \quad (11)$$

here $|x|$ denotes the length of the d -dimensional vector x .

Transition Semigroup

Define the operators P_t by

$$P_t f(x) = \int dy \, p(t, x, y) f(y) \quad (12)$$

for bounded Borel functions f on the d -dimensional Euclidean space. They have the semigroup property:

$$P_t P_s = P_{s+t} \quad (13)$$

which is a consequence of the Chapman–Kolmogorov equation

$$\int dy \, p(s, x, y) p(t, y, z) = p(s+t, x, z) \quad (14)$$

Infinitesimal Generator

Let A denote the infinitesimal generator of the semigroup (P_t) . Let f have second-order partial derivatives that are continuous and bounded. Note that

$$P_t f(x) - f(x) = \int dy \, \frac{e^{-|y|^2/2}}{\sqrt{(2\pi)^d}} \times [f(x + \sqrt{t} y) - f(x)] \quad (15)$$

and expand f in a Taylor series in $\sqrt{t}$. This shows that

$$A f(x) = \lim_{t \rightarrow 0} \frac{1}{t} [P_t f(x) - f(x)] = \frac{1}{2} \Delta f(x) \quad (16)$$

where Δ is the Laplace operator:

$$\Delta = \sum_{i=1}^d \frac{\partial^2}{\partial x_i^2} \quad (17)$$

Then, the differential equation

$$\frac{d}{dt} P_t = A P_t \quad (18)$$

implies that $u(t, x) = P_t f(x)$ satisfies the heat equation

$$\frac{\partial}{\partial t} u = \frac{1}{2} \Delta u \quad (19)$$

Connections to Classical Analysis

The appearance of the Laplace operator and the heat equation signal that there are deep connections between Brownian motion and classical analysis. Many partial differential equations can be related to Brownian motion or its generalizations, the diffusions. Often, then, probabilistic methods can be used to obtain solutions, and/or simulation studies coupled with statistical analysis will yield good approximations.

In particular, as Doob and Kakutani have observed in the 1940s, there is a one-to-one correspondence between the theories of potentials and Brownian motion; see Port and Stone [7] for an account. This opened up the possibility of generalizing potential theory by replacing Brownian motion with more general Markov processes; see Hunt [8, 9] for the seminal work, and Blumenthal and Gettoor [10] for a careful account. Finally, Doob [11] is the masterwork on the general theory that connects potentials, martingales, and Markov processes.

References

- [1] Bachelier, L. (1900). Théorie de la speculation, *Annales scientifique de l'École Normale Supérieure* **17**, 21–86 (Re-issued in Editions Jacques Gabay, Paris, 1995).
- [2] Nelson, E. (1967). *Dynamical Theories of Brownian Motion*, Princeton University Press, Princeton.
- [3] Freedman, D. (1971). *Brownian Motion and Diffusion*, Holden-Day, San Francisco (Re-issued by Springer-Verlag, New York, 1983).
- [4] Karatzas, I. & Shreve, S.E. (1988). *Brownian Motion and Stochastic Calculus*, Springer-Verlag, New York.
- [5] Revuz, D. & Yor, M. (1991). *Continuous Martingales and Brownian Motion*, Springer-Verlag, Berlin.
- [6] Maisonneuve, B. (1971). Ensembles régénératifs, temps locaux et subordinateurs, *Séminaire de Probabilités V, Lecture Notes in Mathematics*, Springer-Verlag, Berlin, Vol. 191, pp. 147–169.
- [7] Port, S. & Stone, C. (1978). *Brownian Motion and Classical Potential Theory*, Academic Press, New York.
- [8] Hunt, G.A. (1957). Markov processes and potentials, *Illinois Journal of Mathematics* **1**, 44–93, 316–369.
- [9] Hunt, G.A. (1958). Markov processes and potentials, *Illinois Journal of Mathematics* **2**, 151–213.
- [10] Blumenthal, R.M. & Gettoor, R.K. (1968). *Markov Processes and Potential Theory*, Academic Press, New York.
- [11] Doob, J.L. (1984). *Classical Potential Theory and its Probabilistic Counterpart*, Springer-Verlag, New York.

Further Reading

Knight, F.B. (1981). *Essentials of Brownian Motion and Diffusion, Mathematical Surveys*, American Mathematical Society, Providence, Vol. 18.

ERHAN ÇINLAR

Poisson Processes

Introduction

Many processes in everyday life that “count” events up to a particular point in time can be accurately described by the so-called Poisson process, named after the French scientist Siméon Poisson (1781–1840; appointed as full professor at the Ecole Polytechnique, Paris, in 1806 as a successor of Fourier). An (ordinary) Poisson process is a special Markov process (see **Markov Processes**), in continuous time, in which the only possible jumps are to the next higher state. A Poisson process may also be viewed as a counting process that has particular, desirable, properties. In this article we shall first give a few equivalent definitions of the Poisson process. Subsequently, we describe the relation between the Poisson process and the (negative) exponential distribution; we then show relations between the Poisson process and the uniform distribution, and between the Poisson process and the binomial distribution. Next we discuss “Poisson conservation” results that are extremely useful in, for example, analyzing queuing networks in which customers arrive according to a Poisson process. The next section is devoted to the uniformization principle, a useful principle in studying the transient behavior of **Markov processes**. Finally, a generalization of the Poisson process, *viz.*, the compound Poisson process is considered.

Definitions of the Poisson Process

A counting process $\{C(t), t \geq 0\}$ is a stochastic process that keeps count of the number of events that have occurred up to time t . Obviously, $C(t)$ is nonnegative and integer valued for all $t \geq 0$. Furthermore, $C(t)$ is nondecreasing in t . $C(t) - C(s)$ equals the number of events in the time interval $(s, t]$, $s < t$.

$C(t)$ could, for example, denote the number of arrivals of customers at a railway station in $(0, t]$, or the number of accidents on a particular highway in that time interval, or the number of births of animals in a particular zoo in $(0, t]$, or the number of calls to a telephone call center during that period. A Poisson process is a counting process that has the desirable additional properties that the numbers of events

in disjoint intervals are independent (*independent increments*) and that the number of events in any given interval depends only on the length of that interval, and not on its particular position in time (*stationary increments*). In the case of the arrivals at the railway station, the stationarity assumption is clearly not fulfilled; there will be many more arrivals between 5 p.m. and 6 p.m. than between, say, 5 a.m. and 6 a.m. Still, one might wish to study the arrival process at the railway station during the rush hour. Restricting oneself to subsequent working days between 5 p.m. and 6 p.m. does allow one to use the stationary increments assumption.

Similarly, the independent increments assumption may be violated in the zoo example, but it will be a reasonably accurate representation of reality in many cases. From the viewpoint of mathematical tractability, these two properties are extremely important. We refer to Feller [1] for an extensive and lucid discussion of stochastic processes with stationary and independent increments. In the sequel we make an additional assumption, which reduces counting processes with stationary and independent increments to Poisson processes.

Definition 1 A Poisson process $\{N(t), t \geq 0\}$ is a counting process with the following additional properties:

1. $N(0) = 0$.
2. The process has stationary and independent increments.
3. $\mathbb{P}(N(h) = 1) = \lambda h + o(h)$ and $\mathbb{P}(N(h) \geq 2) = o(h)$, $h \downarrow 0$, for some $\lambda > 0$.

Above, the $o(h)$ symbol indicates that the ratio $\frac{\mathbb{P}(N(h) \geq 2)}{h}$ tends to zero for $h \downarrow 0$.

An equivalent definition is as follows:

Definition 2 A Poisson process $\{N(t), t \geq 0\}$ is a counting process with the following additional properties:

1. $N(0) = 0$.
2. The process has stationary and independent increments.
3. $\mathbb{P}(N(t) = n) = e^{-\lambda t} \frac{(\lambda t)^n}{n!}$, $n = 0, 1, 2, \dots$

Of course, the last property states that the number of events in any interval of length t is Poisson

2 Poisson Processes

distributed with mean λt . λ is called the *rate* of the Poisson process. It readily follows that the probability generating function of $N(t)$ is given by $\mathbb{E}[z^{N(t)}] = \sum_{n=0}^{\infty} z^n \mathbb{P}(N(t) = n) = e^{-\lambda(1-z)t}$. Differentiation yields $\mathbb{E}[N(t)] = \lambda t$, $\mathbb{E}[N(t)(N(t) - 1)] = (\lambda t)^2$, and hence $\text{Var}(N(t)) = \lambda t$.

The last property of Definition 1 may look awkward at first sight, but is insightful. It states that having two or more events in a small time interval is extremely unlikely, while the probability of a single event is approximately proportional to the length of that small interval.

Remark 1 A Poisson process arises naturally in large populations. Consider such a population of size n , and observe the number of events of a certain type, occurring during a unit time interval. For example, if the probability of an individual calling a call center between, say, 9:00 and 9:01 is p , then the total number of calls during that minute is binomially distributed with parameters n and p . If n becomes large and p gets small such that $np \rightarrow \lambda > 0$, the result is a Poisson process with rate λ of calls per minute.

We close this section by giving yet another, equivalent, definition of the Poisson process.

Definition 3 A Poisson process $\{N(t), t \geq 0\}$ is a counting process with the following additional properties:

1. $N(0) = 0$.
2. The only changes in the process are unit jumps upward. The intervals between jumps are independent, exponentially distributed random variables with mean $1/\lambda$, $\lambda > 0$.

Notice that part (3) of Definition 2 indeed implies that $\mathbb{P}(N(t) = 0) = e^{-\lambda t}$, that is, that the interval until the first event is exponentially distributed with mean $1/\lambda$, denoted $\exp(\lambda)$. The exponential distribution has the memoryless property (uniquely among all continuous distributions): If $T \sim \exp(\lambda)$, then $\mathbb{P}(T > s + t | T > s) = \mathbb{P}(T > t)$. It should also be noted that the memoryless property of the exponential distribution and the independence of successive jump intervals indeed imply that increments are stationary and independent.

Each of the above equivalent definitions has features that make it useful for deriving particular properties, as we shall see in the sequel.

Relation between the Poisson Process and the Exponential Distribution

There is an intimate relation between the Poisson process and the exponential distribution, as already revealed by Definition 3. In this section we go somewhat deeper into this relation.

Let T_1 denote the time of the first event of a Poisson process and let T_n denote the time between the $(n-1)$ th and the n th event, $n = 2, 3, \dots$. Let $S_n = \sum_{i=1}^n T_i$ denote the instant of the n th event. An important observation is that

$$N(t) \geq n \iff S_n \leq t \quad (1)$$

That is, the number of events during $(0, t]$ is at least n if the time until the n th event is no larger than t . Hence

$$\mathbb{P}(N(t) \geq n) = \mathbb{P}(S_n \leq t), \quad n = 1, 2, \dots, t \geq 0 \quad (2)$$

Definition 2 implies that $\mathbb{P}(N(t) \geq n) = \sum_{j=n}^{\infty} e^{-\lambda t} \frac{(\lambda t)^j}{j!}$, and hence

$$\mathbb{P}(S_n \leq t) = 1 - \sum_{j=0}^{n-1} e^{-\lambda t} \frac{(\lambda t)^j}{j!}, \quad n = 1, 2, \dots, (t \geq 0) \quad (3)$$

This is the well-known result that the sum of n independent, $\exp(\lambda)$ distributed, random variables is $\text{Erlang}(n; \lambda)$ distributed. In particular, $T_1 = S_1$ is $\exp(\lambda)$. The density of this Erlang distribution equals $\lambda e^{-\lambda t} \frac{(\lambda t)^{n-1}}{(n-1)!}$, $n = 1, 2, \dots, (t > 0)$, which also follows from the reasoning below:

$$\begin{aligned} \mathbb{P}(t < S_n \leq t + h) &= \mathbb{P}(N(t) = n - 1, \text{ one event in } (t, t + h]) \\ &\quad + o(h) \\ &= \mathbb{P}(N(t) = n - 1) \mathbb{P}(\text{one event in } (t, t + h]) \\ &\quad + o(h) \\ &= e^{-\lambda t} \frac{(\lambda t)^{n-1}}{(n-1)!} \lambda h + o(h) \end{aligned} \quad (4)$$

Now divide by h and let $h \downarrow 0$ to get the density.

A well-known property of the exponential distribution is that it has a constant failure rate equal to λ : if T is $\exp(\lambda)$ distributed, then

$$\mathbb{P}(T \leq t + h | T \geq t) = \lambda h + o(h), \quad h \downarrow 0 \quad (5)$$

implying that $\lim_{h \downarrow 0} \mathbb{P}(T \leq t + h | T \geq t)/h = \lambda$. We have also seen this property in Definition 1 of the Poisson process.

To illustrate the concept of failure rate, consider a probability distribution function $F(\cdot)$ describing the lifetime of a device and denote its density by $f(\cdot)$. Then the failure rate of $F(\cdot)$ is defined as $r(t) = f(t)/[1 - F(t)]$; $r(t) dt$ gives the probability that the device fails during $(t, t + dt]$, given that it is “alive” at time t .

Relations between the Poisson Process and the Uniform Distribution

In this section we discuss a property of the Poisson process that is often very useful in applications. If exactly one event of a Poisson process has occurred in $(0, t]$, then the time of that occurrence is uniformly distributed on $(0, t)$. The informal explanation is that, because of the stationary and independent increments, each subinterval of equal length in $(0, t)$ has the same probability to contain that event. The formal derivation is

$$\begin{aligned} \mathbb{P}(T_1 \leq s | N(t) = 1) &= \frac{\mathbb{P}(\text{one event in } (0, s], \text{ no event in } (s, t])}{\mathbb{P}(N(t) = 1)} \\ &= \frac{\mathbb{P}(N(s) = 1) \mathbb{P}(N(t-s) = 0)}{\mathbb{P}(N(t) = 1)} = \frac{s}{t}, \\ & \quad 0 \leq s \leq t \end{aligned} \quad (6)$$

More generally, the following can be proved for a Poisson process: If $N(t) = n$, then the event times $S_1, \dots, S_n$ are distributed like the order statistics of n independent random variables that are uniformly distributed on $(0, t)$.

The property that a Poisson arrival “is just as likely to occur in any interval” has proved to be extremely useful in, for example, queuing theory. Firstly, it forms the basis of the well-known PASTA (*Poisson Arrivals See Time Averages*) property. In queuing terms this property states that an outside observer, arriving to a queue according to a Poisson process,

sees the system as if it were in steady state. Consider, for example, an $M/G/s$ queue. This is a queuing model in which arrivals occur according to a Poisson process (M), requiring a generally distributed (G) service time at one of s servers. The PASTA property implies for the $M/G/s$ queue that the number of customers seen by an arriving customer has the same distribution as the steady-state number of customers. The PASTA property was first rigorously proved by Wolff [2].

Another well-known application of the property that Poisson arrivals are uniformly distributed over an interval is provided by the $M/G/\infty$ queue [3]. This is an $M/G/s$ queue with an ample ($s = \infty$) number of servers. The $M/G/\infty$ system may be used, for example, to model certain production or transportation systems; the customers might then be pallets moving on a conveyor belt, or cars on a road. Given that n arrivals have occurred in the $M/G/\infty$ system up to t , the above property specifies the distribution of the arrival epochs. One can then easily determine the probability that such a customer, who has arrived in $(0, t]$, is still present at t . Thus one can obtain the distribution of the number of customers $L(t)$ in the $M/G/\infty$ system at t , starting from an empty system at 0 (*cf.* Takács [3, Chapter 3]). It turns out that $L(t)$ has a Poisson distribution with time-dependent rate $\lambda \int_0^t \mathbb{P}(B > y) dy$, B denoting service time of an individual customer. When $t \rightarrow \infty$, one gets the so-called stationary distribution of $L(t)$, which is Poissonian with rate $\lambda \mathbb{E}B$. Furthermore, the number of customers served in $(0, t]$ is also Poisson distributed, with rate $\lambda \int_0^t \mathbb{P}(B \leq y) dy$, and it is statistically independent of the process $\{L(t), t \geq 0\}$.

Relations between the Poisson Process and the Binomial Distribution

Two important consequences of the stationary and independent increments properties are the following:

1. For any $t > u \geq 0$ and integers $n \geq k$,

$$\begin{aligned} \mathbb{P}(N(u) = k | N(t) = n) &= \binom{n}{k} \left(\frac{u}{t}\right)^k \\ &\quad \times \left(1 - \frac{u}{t}\right)^{n-k}, \quad k = 0, 1, \dots, n \end{aligned} \quad (7)$$

That is, given that n events occurred in $(0, t]$, the probability that k of them occurred in

- (0, u] is given by the binomial distribution with parameters n and “success” probability $p = u/t$.
2. For two independent Poisson processes $N_1(t)$ and $N_2(t)$, having rates λ_1 and λ_2 , respectively,

$$\begin{aligned} \mathbb{P}(N_1(t) = k | N_1(t) + N_2(t) = n) &= \binom{n}{k} \\ &\times \left(\frac{\lambda_1}{\lambda_1 + \lambda_2} \right)^k \left(\frac{\lambda_2}{\lambda_1 + \lambda_2} \right)^{n-k}, \\ k &= 0, 1, \dots, n \end{aligned} \quad (8)$$

This result implies, in particular, that if two such processes “compete” on which one will be the first to occur, the probability that $N_i(t)$ is “quicker” is given by $\frac{\lambda_i}{\lambda_1 + \lambda_2}$, $i = 1, 2$.

See again Remark 1 for yet another relation between the Poisson process and the binomial distribution.

Conservation Properties

The Poisson process satisfies some conservation properties that greatly enhance its applicability. In particular, random splitting of a Poisson process results in independent Poisson processes; similarly, merging independent Poisson processes again produces a Poisson process. We now formalize these statements, without proof.

Proposition 1 *Suppose that events occur according to a Poisson process $\{N(t), t \geq 0\}$ of rate λ . Further, suppose that each event is classified as a type i event with probability $p_i \in (0, 1)$, $i = 1, \dots, K$. Let $N_i(t)$ denote the number of type- i events in $(0, t]$. Then $\{N_1(t), t \geq 0\}, \dots, \{N_K(t), t \geq 0\}$, are independent Poisson processes with corresponding rates $\lambda p_1, \dots, \lambda p_K$.*

Indeed, it is easy to see that an arbitrary interval of length h will contain a type- i event with probability $\lambda p_i h + o(h)$, $h \downarrow 0$. The independence of the various Poisson processes is less obvious. But remember that the knowledge that, say, $N_1(t) = j$ just implies that j out of the t/h intervals of length h contain a type-1 event; it does not really imply anything about the occurrence of events of other types.

Now consider the dual situation, in which there are K independent Poisson processes $\{N_1(t),$

$t \geq 0\}, \dots, \{N_K(t), t \geq 0\}$ with corresponding rates $\lambda_1, \dots, \lambda_K$, and in which we count together all events of all these processes.

Proposition 2 *The sum process $\{N(t), t \geq 0\}$ of K independent Poisson processes $\{N_i(t), t \geq 0\}$, $i = 1, \dots, K$, with $N(t) = \sum_{i=1}^K N_i(t)$, $t \geq 0$, is a Poisson process with rate equal to the sum $\sum_{i=1}^K \lambda_i$ of the rates of the individual processes.*

This proposition can be easily proved using any of the three Definitions 1, 2, or 3. For example, it is straightforward (and intuitive) to verify the following property of Definition 1: $\mathbb{P}(N(h) = 1) = \lambda h + o(h)$, from the similar property of the individual Poisson processes. The proposition also follows using the multiplication property of probability generating function (PGF) of independent random variables and using the fact that the PGF of $N_i(t)$ equals $e^{-\lambda_i(1-z)t}$. Alternatively, Definition 3 also works well here; use the fact that the minimum of K independent $\exp(\lambda_i)$ distributed random variables is $\exp\left(\sum_{i=1}^K \lambda_i\right)$ distributed, combined with the memoryless property of the exponential distribution.

Networks of Queues

In 1956 Paul Burke [4] proved the output theorem, a result that plays a crucial role in the study of networks of **queues**. This theorem states the following: Consider an $M/M/s$ queue, viz., a queuing system with s servers, first-come-first-served service operation, a $\text{Poisson}(\lambda)$ arrival process of customers, and independent $\exp(\mu)$ distributed service times. Assume that $\lambda < s\mu$, implying that the process of number of customers in the system reaches an equilibrium distribution. Then the output process of the $M/M/s$ queue, counting the number of departures in $(0, t]$, is a Poisson process with rate λ . It can be shown that, when $s < \infty$, the above “conservation of Poisson flows” property only holds for the multiserver queue when service times are exponential. However, when $s = \infty$, it holds for general service times.

In 1957 Reich [5] provided an extremely elegant proof of the output theorem. He observed that the process constituted by the number of customers in the $M/M/s$ system is a reversible stochastic process. Viewed backward in time, this process is statistically indistinguishable from the original process. The output process of the reversed-in-time $M/M/s$ queue

coincides with the Poisson input process of the original $M/M/s$ queue, and hence is also Poissonian. Using once more the reversibility property leads to the conclusion that the output process of the original $M/M/s$ queue is a Poisson process.

The previously discussed properties of conservation of Poisson flows under merging, splitting, and passing through an $M/M/s$ system allow one to analyze the queue-length process in a network of queues $Q_1, \dots, Q_K$, all having external Poisson arrival processes, exponential service times, possibly multiple servers, and with Markovian routing of customers from queue to queue. Indeed, the output of an $M/M/s_i$ queue Q_i is Poisson. If a fraction p_{ij} is routed to Q_j , then the resulting flow is again Poisson; and if it merges with another, independent, Poisson process, the resulting process is also Poisson. This is only part of the explanation why such networks of exponential multiserver queues have a joint queue-length distribution that exhibits a product form, the i th term of the product corresponding to an $M/M/s_i$ queue Q_i in isolation. For an excellent exposition of the theory of product-form queueing networks we refer to Kelly [6].

Uniformization

In many application areas it is important to determine the *transient* behavior of a Markov process. Consider a Markov process $\{X(t), t \geq 0\}$, with transition probabilities P_{ij} and visit times to all states being exponentially distributed with the *same* mean $1/\lambda$. Hence the number of transitions in $(0, t]$ is Poisson distributed with mean λt . Denoting the corresponding Poisson process by $\{N(t), t \geq 0\}$ we have the following:

$$\mathbb{P}(X(t) = j | X(0) = i) = \sum_{n=0}^{\infty} P_{ij}^{(n)} e^{-\lambda t} \frac{(\lambda t)^n}{n!} \quad (9)$$

where $P_{ij}^{(n)}$ denote the n -step transition probabilities between the states (with $P_{(i,i)}^{(0)} = 1$); put differently, these are the n -step transition probabilities of the discrete-time Markov chain underlying the Markov process. The above formula is derived by conditioning on the number of transitions in $(0, t]$.

Equation (9) is computationally advantageous, as the infinite sum can be truncated while the $P_{ij}^{(n)}$ can be evaluated efficiently. An interesting idea, apparently

due to Jensen [7], allows extension of this principle to the case of *unequal* mean visit times. Let us assume that the mean visit time to state i is $1/\lambda_i$. Let λ be such that $\lambda \geq \lambda_i$ for all i . Consider a Poisson process with rate λ . Now consider a Markov process that spends $\exp(\lambda)$ in any state i and then jumps with probability $\hat{P}_{ij} = \frac{\lambda_i}{\lambda} P_{ij}$ to j and with

probability $\hat{P}_{ii} = 1 - \frac{\lambda_i}{\lambda}$ back to i ; one might call this a fictitious transition. Remember that a sum of a geometrically distributed number of independent exponentially distributed random variables is again exponentially distributed; hence the sum of consecutive visit times to state i , before another state is visited, is $\exp(\lambda_i)$. A little thought now shows that this new Markov process is really the same as the original Markov process. The transient behavior of the original Markov process can hence be inferred from that of the new Markov process, by using equation (9), with $P_{ij}^{(n)}$ being replaced by the n -step transition probabilities of the new discrete-time Markov chain that has transition probabilities $\hat{P}_{ij}$.

As an application of this uniformization technique, consider a system that alternates between up and down periods, up (down) periods being exponentially distributed with parameter μ_u (μ_d). We leave it to the reader to verify that by taking $\lambda = \mu_u + \mu_d$ as uniformization parameter, one gets new transition probabilities $\hat{P}_{uu} = \hat{P}_{du} = \frac{\mu_d}{\mu_u + \mu_d}$, and hence

$$\hat{P}_{uu}^{(n)} = \hat{P}_{du}^{(n)} = \frac{\mu_d}{\mu_u + \mu_d}, \quad n = 1, 2, \dots, \text{ leading to}$$

$$\begin{aligned} P(X(t) = u | X(0) = u) \\ = \frac{\mu_d}{\mu_u + \mu_d} + \frac{\mu_u}{\mu_u + \mu_d} e^{-(\mu_u + \mu_d)t} \end{aligned} \quad (10)$$

and by symmetry,

$$\begin{aligned} P(X(t) = d | X(0) = d) \\ = \frac{\mu_u}{\mu_u + \mu_d} + \frac{\mu_d}{\mu_u + \mu_d} e^{-(\mu_u + \mu_d)t} \end{aligned} \quad (11)$$

There is a large variety of other applications of the uniformization property in stochastic analysis and simulation, that may be found in various textbooks; the present example is taken from Ross [8, Section 5.8].

The Compound Poisson Process

A limitation of the Poisson process is that the jumps are always of unit size. A stochastic process $\{X(t), t \geq 0\}$ is called a *compound Poisson process* if it can be represented by

$$X(t) = \sum_{i=0}^{N(t)} Y_i, \quad t \geq 0 \quad (12)$$

where $\{N(t), t \geq 0\}$ is a Poisson process and $Y_1, Y_2, \dots$ are independent, identically distributed random variables that are also independent of $\{N(t), t \geq 0\}$. $X(t)$ could, for example, represent the accumulated workload input into a queuing system in $(0, t]$: Customers arrive according to a Poisson process $\{N(t), t \geq 0\}$, and the i th customer requires a service time of length Y_i . Alternatively, the Poisson arrival process might represent the number of insurance claims in $(0, t]$, while the Y_i represent independent claim sizes. $X(t)$ is then the total amount of monetary claims up to time t .

As another example one can envision a device subject to a series of independent random shocks. If the damage caused by the i th shock is Y_i and the number of shocks in $(0, t]$ is $N(t)$, then the total accumulated damage up to time t is given by $X(t)$.

It is easily seen that, if the rate of the Poisson process equals λ and the Y_i have a common Laplace–Stieltjes transform $\beta(s) = \mathbb{E}[e^{-sY_i}]$, then

$$\mathbb{E}[e^{-sX(t)}] = e^{-\lambda(1-\beta(s))t} \quad (13)$$

Differentiation then readily yields that $\mathbb{E}[X(t)] = \lambda t \mathbb{E}Y_1$ and $\text{Var}[X(t)] = \lambda t \mathbb{E}[Y_1^2]$.

Compound Poisson processes are an important subclass of Lévy processes (see **Lévy Processes**). We refer to Bertoin [9] for a detailed account of the theory of Lévy processes.

Epilogue

The Poisson process is a stochastic counting process that arises naturally in daily life situations in which

there is a large population of individuals who, more or less independently of each other, have a small probability of contributing to the count in the next small time interval. The (compound) Poisson process has beautiful mathematical properties, which make it a very powerful tool for stochastic modeling and analysis. We refer the interested reader to the monograph of Kingman [10]. Excellent textbook treatments can be found in the books of Ross [8] and Tijms [11].

References

- [1] Feller, W. (1966). *An Introduction to Probability Theory and its Applications*, John Wiley & Sons, New York, Vol. II.
- [2] Wolff, R.W. (1982). Poisson arrivals see time averages, *Operations Research* **30**, 223–231.
- [3] Takács, L. (1962). *Introduction to the Theory of Queues*, Oxford University Press, Oxford.
- [4] Burke, P.J. (1956). The output of a queuing system, *Operations Research* **4**, 699–704.
- [5] Reich, E. (1957). Waiting times when queues are in tandem, *Annals of Mathematical Statistics* **28**, 768–773.
- [6] Kelly, F.P. (1979). *Reversibility and Stochastic Networks*, John Wiley & Sons, New York.
- [7] Jensen, A. (1953). Markoff chains as an aid in the study of Markoff processes, *Skandinavisk Aktuarietidskrift* **36**, 87–91.
- [8] Ross, S.M. (1983). *Stochastic Processes*, John Wiley & Sons, New York.
- [9] Bertoin, J. (1996). *Lévy Processes*, Cambridge University Press.
- [10] Kingman, J.F.C. (1993). *Poisson Processes*, Oxford Studies in Probability, 3, The Clarendon Press, Oxford University Press, New York.
- [11] Tijms, H.C. (1994). *Stochastic Models. An Algorithmic Approach*, John Wiley & Sons, Chichester.

Related Articles

Intensity Function; Intensity Functions for Non-homogeneous Poisson Processes.

ONNO J. BOXMA AND URI YECHIALI

Lévy Processes

Definitions and Assumptions

On a given probability space $(\Omega, \mathcal{F}, P)$, a continuous-time *stochastic process* is a collection of random variables $\{X_t | t \geq 0\}$. A Lévy process $\{X_t | t \geq 0\}$ is a continuous-time stochastic process having the following three properties:

1. $X_{t+s} - X_s$ is distributed like X_t for any nonnegative s and t .
2. For every integer $n \geq 2$ and $0 = t_0 < t_1 < \dots < t_n$, the random variables $\{X_{t_i} - X_{t_{i-1}} | 1 \leq i \leq n\}$ are independent.
3. $X_t \rightarrow 0$ in probability as $t \downarrow 0$.

Since for $0 < t < s$, $X_{t+s} - X_s$ and $X_s - X_{s-t}$ are distributed like X_t , it is clear that $X_{t+s} \rightarrow X_s$ and $X_{t-s} \rightarrow X_s$ in probability as $t \downarrow 0$ for any $s \geq 0$. Therefore, it is customary to refer to these three properties as 1: stationary increments, 2: independent increments, and 3: continuous in probability, respectively. Thus a Lévy process is a continuous-time stochastic process which has stationary, independent increments and is continuous in probability.

A more general definition of a Lévy process is as follows. Let $\{\mathcal{F}_t | t \geq 0\}$ be a nondecreasing family of σ algebras all of which are contained in $\mathcal{F}$. Such a family is called a *filtration*. In addition assume that the following two conditions are satisfied: (a) $\lim_{s \downarrow t} \mathcal{F}_s = \mathcal{F}_t$ and (b) $\mathcal{F}_0$ contains all the P null sets $\mathcal{F}$ (in particular, so does $\mathcal{F}$). Such a right continuous and augmented filtration is said to satisfy the *usual conditions* or is called a *standard filtration*. A Lévy process with respect to $\{\mathcal{F}_t | t \geq 0\}$ is a process satisfying the first and third condition above and where condition 2 is replaced by:

2. $X_{t+s} - X_t$ is independent of $\mathcal{F}_t$ for every $s, t \geq 0$.

It is easy to check that this condition implies that $\{X_t | t \geq 0\}$ has independent increments. Also, it can be shown that a Lévy process has a right continuous left limit (càdlàg) version, that is, there exists a càdlàg process $\{\tilde{X}_t | t \geq 0\}$ such that $P[\tilde{X}_t = X_t] = 1$ for each $t \geq 0$. It can be shown that the filtration generated by a right continuous process having

stationary independent increments is right continuous. A nice exposition of Lévy processes emphasizing such properties can be found in Section I.4 of [1]. See also [2], which has become a standard reference regarding Lévy processes.

Infinite Divisibility and Basic Structure

A random variable X has an infinitely divisible distribution if and only if for every $n \geq 2$ there are independent and identically distributed (i.i.d.) random variables, $X_1^{(n)}, \dots, X_n^{(n)}$, such that $\sum_{i=1}^n X_i^{(n)}$ is distributed like X . It is well known (e.g., Section 7.6 of [3] and Section 9.5 of [4]) that if X has an infinitely divisible distribution, then $Ee^{i\alpha X} = e^{\varphi(\alpha)}$, where, with A^c being the complement of a set A ,

$$\begin{aligned} \varphi(\alpha) = i\alpha c - \frac{\sigma^2 \alpha^2}{2} + \int_{[-1,1]^c} (e^{i\alpha x} - 1) \nu(dx) \\ + \int_{[-1,1]} (e^{i\alpha x} - 1 - i\alpha x) \nu(dx) \end{aligned} \quad (1)$$

and ν is a measure satisfying $\nu([-1,1]^c) < \infty$, $\int_{[-1,1]} x^2 \nu(dx) < \infty$, and $\nu(\{0\}) = 0$. The measure ν is called the *Lévy measure* and equation (1) is often referred to as the *Lévy–Khintchine formula*. It can be shown that for any measure satisfying these assumptions and for any c and σ^2 , there is an infinitely divisible distribution with a characteristic function having an exponent φ of this form. In fact, letting

$$I_n = \begin{cases} (n^{-1}, (n-1)^{-1}) & n = 2, 3, \dots \\ [n^{-1}, (n-1)^{-1}) & n = 2, 3, \dots \\ (1, \infty) & n = 1 \\ (-\infty, -1) & n = -1 \end{cases} \quad (2)$$

we set $\lambda_n = \nu(I_n)$, and when $\lambda_n > 0$, let $F_n(x) = \nu(I_n \cap (-\infty, x]) / \lambda_n$ (when $\lambda_n = 0$, F_n can be arbitrary). Let N_n be a Poisson random variable with mean λ_n and $X_k^{(n)}$ have distribution F_n . In addition, let Y be a normal random variable with mean c and variance σ^2 . Assume that all of these random variables are independent. Finally, let $S_0^{(n)} = 0$ and for $m \geq 1$, $S_m^{(n)} = \sum_{k=1}^m X_k^{(n)}$. Then,

$$Y + S_{N_1}^{(-1)} + S_{N_1}^1 + \sum_{\substack{n=-\infty \\ n \neq -1, 0, 1}}^{\infty} \left(S_{N_n}^{(n)} - \lambda_n E X_1^{(n)} \right) \quad (3)$$

2 Lévy Processes

has a distribution with the desired exponent and it can be shown that the infinite sum converges almost surely. Owing to this decomposition, it is easy to verify that when $\nu(-\infty, 0) = 0$, then E^{zX} is finite for every complex z with $\Re z \leq 0$ and is equal to $e^{\varphi(-iz)}$, where $\varphi(-iz)$ is the Lévy–Khintchine formula in which $i\alpha$ is replaced by z , that is,

$$\begin{aligned} \varphi(-iz) = & zc + \frac{\sigma^2 z^2}{2} + \int_{(1, \infty)} (e^{zx} - 1) \nu(dx) \\ & + \int_{(0, 1]} (e^{zx} - 1 - zx) \nu(dx) \end{aligned} \quad (4)$$

and in particular, for $\alpha \geq 0$ real,

$$\begin{aligned} \log Ee^{-\alpha X} = & -\alpha c + \frac{\sigma^2 \alpha^2}{2} + \int_{(1, \infty)} (e^{-\alpha x} - 1) \nu(dx) \\ & + \int_{(0, 1]} (e^{-\alpha x} - 1 + \alpha x) \nu(dx) \end{aligned} \quad (5)$$

Replacing Y by a **Brownian motion** $\{Y_t | t \geq 0\}$ and N_n by a **Poisson process** $\{N_n(t) | t \geq 0\}$ with rate λ_n , the resulting process is a Lévy process $\{X_t | t \geq 0\}$ satisfying $Ee^{i\alpha X_1} = e^{\varphi(\alpha)}$.

For a Lévy process $\{X_t | t \geq 0\}$, $X_t = \sum_{k=1}^n (X_{kt/n} - X_{(k-1)t/n})$, and thus by the stationary increments property, it follows that X_t has an infinite divisible distribution. In particular, if $Ee^{i\alpha X_t} = e^{\varphi(\alpha)t}$, then $Ee^{i\alpha X_t} = e^{\varphi(\alpha)t}$. φ is called the *characteristic exponent* (or just *exponent*) of the Lévy process. In particular, for any Lévy process, there is an exponent which satisfies the Lévy–Khintchine formula, and for any such exponent, there is a Lévy process having this exponent.

When $\int_{[-1, 1]} x \nu(dx) < \infty$, the non-Brownian part of the Lévy process almost surely has bounded variation sample paths; otherwise, it does not. In the case that it does, it is of the form $J_1(t) - J_2(t) + ct$, where J_1 and J_2 are independent nondecreasing pure jump Lévy processes. A nondecreasing Lévy process is called a *subordinator*. From the preceding, any subordinator is a finite or infinite sum of independent compound Poisson processes. It either has a finite number of jumps on any finite interval or an infinite number of jumps on any finite interval, but not both. An example of a subordinator with an infinite number of jumps on finite intervals is the gamma process, with $\nu(dx) = ae^{-bx}/x dx$. For this process, X_t has a gamma distribution,

that is, with density $\frac{b^{at}}{\Gamma(at)} e^{-bx} x^{at-1}$ for $x > 0$, and zero otherwise.

When $\nu(-\infty, 0) = 0$, the Lévy process has no negative jumps and for every complex z with $\Re z \leq 0$, $Ee^{zX_t} = e^{\varphi(-iz)t}$ and is finite. In particular, when $z = -\alpha$, where $\alpha \geq 0$ (real valued), we denote the Laplace–Stieltjes exponent by $\varphi_L(\alpha) = \varphi(i\alpha)$. It is easy to check that φ_L is a convex function which is differentiable on $(0, \infty)$ of all orders. In particular, if the Lévy process is not deterministic (necessarily a linear function), then φ_D is a strictly convex function. It also holds that φ_D is nonincreasing if and only if the Lévy process is a subordinator and that when the process is neither a subordinator nor deterministic then necessarily $\varphi_D(\alpha) \rightarrow \infty$ as $\alpha \rightarrow \infty$ and thus for every $\beta > 0$ there exists an $\alpha > 0$ for which $\varphi_D(\alpha) = \beta$. Let us denote this value by the function $\psi_D(\beta)$, that is, $\varphi_D(\psi_D(\beta)) = \beta$, so that ψ_D is an inverse function of φ_D when the domain of φ_D is $\{\alpha | \varphi_D(\alpha) > 0\}$ and the range is $(0, \infty)$. In this case, when $\varphi'_D(0) < 0$, there are two roots to the equation $\varphi_D(\alpha) = 0$ and when $\varphi'_D(0) \geq 0$ there is only one.

In general, it can be shown that if $\int_{[-1, 1]^c} |x| \nu(dx) < \infty$ then

$$EX_t/t = c + \int_{[-1, 1]^c} x \nu(dx) = -i\varphi'(0) \quad (6)$$

if $\int_{[-1, 1]^c} x^2 \nu(dx) < \infty$ then

$$\text{Var}(X_t)/t = \sigma^2 + \int_{(-\infty, \infty)} x^2 \nu(dx) = -\varphi''(0) \quad (7)$$

and if $\int_{[-1, 1]^c} |x|^n \nu(dx) < \infty$ for $n > 2$ then

$$(-i)^n \varphi^{(n)}(0) = \int_{(-\infty, \infty)} |x|^n \nu(dx) \quad (8)$$

There is a close relationship between Lévy processes and Poisson random measures. In particular, if $A_1, \dots, A_n$ are mutually disjoint Borel subsets of $[0, \infty) \times (-\infty, \infty)$, then if $N(A_i)$ counts the number of instances where $(t, \Delta X_t) \in A_i$, then $N(A_1), \dots, N(A_n)$ are independent Poisson random variables with $EN(A_i) = \nu(A_i)$. If the closure of A_i is disjoint from $[0, \infty) \times \{0\}$ and it is contained in $[0, t] \times (-\infty, \infty)$ for some $t < \infty$, then $\nu(A_i) < \infty$.

Martingales Associated with Lévy Processes

In this section, we consider a right continuous Lévy process $\{X_t | t \geq 0\}$ with respect to some filtration $\{\mathcal{F}_t | t \geq 0\}$ satisfying the usual conditions. The Wald martingale associated with $\{X_t | t \geq 0\}$ is the process $e^{i\alpha X_t - \varphi(\alpha)t}$. It is easy to check that this process is indeed a martingale. Similarly, for the case with no negative jumps $e^{-\alpha X_t - \varphi_L(\alpha)t}$ is a martingale for every $\alpha \geq 0$, and thus we have, for every $\beta > 0$, $e^{-\psi(\beta)X_t - \beta t}$ is also a martingale.

This last martingale can be used to show that for $T_x = \inf\{t | X_t < -x\}$ we have $Ee^{-\beta T_x} = e^{-\psi_D(\beta)x}$. $\{T_x | x \geq 0\}$ is in fact a subordinator with exponent $-\psi_D(\beta)$. Letting $\psi_D(0) = \psi_D(0+) = \lim_{\beta \downarrow 0} \psi_D(\beta)$, we have, if $\varphi'_D(0) \geq 0$ then $\psi_D(0) = 0$ and thus $P[T_x < \infty] = 1$, and if $\varphi'_D(0) < 0$ then $\psi_D(0) > 0$ and then $P[T_x < \infty] = e^{-\psi_D(0)x}$. In particular, this implies that $-\inf\{X_t | t \geq 0\}$ is exponentially distributed with rate $\varphi_D(0)$ as $T_x < \infty$ if and only if $-\inf\{X_t | t \geq 0\} > x$. If $S_x = \inf\{t | X_t = -x\}$, then it is an easy exercise to show that $P[S_x = T_x] = 1$, so that the same results hold with respect to S_x .

In [5], a useful martingale was established that has been used extensively in the literature. The results are as follows.

Theorem 1 *Let $\{X_t | t \geq 0\}$ be a Lévy process and let $\{Y_t | t \geq 0\}$ be right continuous with left limits. Denote $Y_{t-} = \lim_{s \uparrow t} Y_s$, $\Delta Y_t = Y_t - Y_{t-}$*

$$Y_t^c = Y_t - \sum_{0 < s \leq t} \Delta Y_s \quad (9)$$

Assume that $Y^c(\cdot)$ has bounded expected variation on compact intervals, that is, assume that it is a difference of two nonnegative, nondecreasing, continuous processes each having a finite expected value at each time epoch. Also assume that

$$E \sum_{0 < s \leq t} 1_{\{|\Delta Y_s| > 1\}} < \infty \quad (10)$$

and that

$$E \sum_{0 < s \leq t} |\Delta Y_s| 1_{\{|\Delta Y_s| \leq 1\}} < \infty \quad (11)$$

Finally, set $Z_t = X_t + Y_t$ (in particular, $Z_0 = Y_0$). Then the following is a zero mean martingale:

$$\begin{aligned} \varphi(\alpha) \int_0^t e^{i\alpha Z_s} ds + e^{i\alpha Z_0} - e^{i\alpha Z_t} + i\alpha \int_0^t e^{i\alpha Z_s} dY_s^c \\ + \sum_{0 < s \leq t} e^{i\alpha Z_s} (1 - e^{-i\alpha \Delta Y_s}) \end{aligned} \quad (12)$$

If in addition $\nu(-\infty, 0) = 0$ (no negative jumps) and $\{Z_t | t \geq 0\}$ is a nonnegative process, then the following is a zero mean martingale:

$$\begin{aligned} \varphi_L(\alpha) \int_0^t e^{-\alpha Z_s} ds + e^{-\alpha Z_0} - e^{-\alpha Z_t} - \alpha \int_0^t e^{-\alpha Z_s} dY_s^c \\ + \sum_{0 < s \leq t} e^{-\alpha Z_s} (1 - e^{\alpha \Delta Y_s}) \end{aligned} \quad (13)$$

It is noted that equations (10) and (11) are weaker than the condition

$$E \sum_{0 < s \leq t} 1_{\{|\Delta Y_s| > 0\}} < \infty \quad (14)$$

which is assumed in equation [5], but the results remain valid with this weaker condition. This martingale was successfully applied in a variety of models in queuing, storage models, and networks, risk, and finance.

Reflected Lévy Processes

A special case where the martingale of [5] applies is where $Y_t = z + L_t$, $z \geq 0$, and L_t is defined as follows.

1. $\{L_t | t \geq 0\}$ is nonnegative, nondecreasing, and right continuous.
2. $\{Z_t | t \geq 0\}$ is nonnegative.
3. One (and then all) of the following equivalent conditions holds (see [6]):
 - (a) For each $t \geq 0$, if $L_{t-} < L_s$ for $s > t$ then $Z_t = 0$.
 - (b) $\int_0^\infty Z_s dL_s = 0$.
 - (c) $\{L_t | t \geq 0\}$ is the minimal process satisfying 1 and 2.
 - (d) For each $t \geq 0$, $L_t = - \inf_{0 \leq s \leq t} \min(z + X_s, 0)$.

4 Lévy Processes

For this case $\{Z_t | t \geq 0\}$ is called a *reflected Lévy process*. When the Lévy process has no negative jumps, the process $\{Y_t | t \geq 0\}$ is continuous with $Y_0 = z$ and it can be verified that $EY_t < \infty$ for each $t \geq 0$. These properties together with the martingale in [5] can be used to show that

$$\varphi_L(\alpha) \int_0^t e^{-\alpha Z_s} ds + e^{-\alpha z} - e^{-\alpha Z_t} - \alpha L_t \quad (15)$$

is a zero mean martingale, which implies in a very simple way that if $EX_1 < 0$ then the distribution of Z_t converges in law to a distribution which has the following generalized Polachek–Khinchine formula (see [5] for details).

$$Ee^{-\alpha Z_\infty} = \frac{\alpha \varphi'_L(0)}{\varphi_L(\alpha)} \quad (16)$$

This result is well known and has also been proved by various other methods (e.g., see [7–9]). Recently, in [10], a class of martingales associated with a reflected Lévy process has been established, of which the martingale of [5] is a special case when applied to the reflected case.

The Wiener–Hopf Factorization

It is well known (e.g., [7, 10]) that the following decomposition holds. A similar decomposition holds for random walks. For every $q > 0$, there are two unique functions φ_q^+ and φ_q^- having the following properties:

1. φ_q^+ and φ_q^- are the characteristic functions of infinitely divisible laws.
2. φ_q^+ is analytic on $\{z | \Im(z) \geq 0\}$ and φ_q^- is analytic on $\{z | \Im(z) \leq 0\}$.
3. For every real α

$$\frac{q}{q + \varphi(\alpha)} = \varphi_q^+(\alpha) \varphi_q^-(\alpha) \quad (17)$$

Equation (17) is known as the celebrated *Wiener–Hopf factorization*. An interesting fact is that this result can be proved using the martingale from [5]. This is explicitly done in [10]. The Wiener–Hopf factorization has been used to prove various results about Lévy processes, the generalized Polachek–Khinchine formula, and results regarding first passage times in the same way that it has been applied to

random walks (e.g., [7]). Many of these results are now known to have independent martingale proofs.

Multidimensional Lévy Processes

An infinitely divisible distribution has a multivariate analog. In particular, it can be shown that a random vector has an infinitely divisible distribution if and only if it has a characteristic function of the form $e^{\varphi(\alpha)}$ where $\alpha \in \mathbb{R}^d$ and

$$\begin{aligned} \varphi(\alpha) = i\alpha^T c - \frac{1}{2}\alpha^T \Sigma \alpha + \int_{B^c} (e^{i\alpha^T x} - 1) \nu(dx) \\ + \int_B (e^{i\alpha^T x} - 1 - i\alpha^T x) \nu(dx) \end{aligned} \quad (18)$$

where B is a unit ball (open or closed), ν is a measure satisfying $\nu(B^c) < \infty$, $\int_B x^T x \nu(dx) < \infty$, and $\nu(\{0\}) = 0$, where T means transposition and α , x , and c are in $\mathbb{R}^d$. B can be replaced by other bounded sets like $[-1, 1]^d$, of which the interior contains the origin.

If $\{X_t | t \geq 0\}$ is a multidimensional Lévy process, then for every $\alpha \in \mathbb{R}^d$, $\{\alpha^T X_t | t \geq 0\}$ is a one-dimensional Lévy process with exponent $\varphi^*(\beta) = \varphi(\beta\alpha)$. In particular, inserting $\beta = 1$ immediately implies that the martingale result of [5] applies without change in the multivariate case, that is, one simply replaces α by α^T in equations (12) and (13), and assumes that each coordinate of Y^c has bounded expected variation on finite intervals and that

$$E \sum_{0 < s \leq t} 1_{\{\|\Delta Y_s\| > 1\}} < \infty \quad (19)$$

and that

$$E \sum_{0 < s \leq t} \|\Delta Y_s\| 1_{\{\|\Delta Y_s\| \leq 1\}} < \infty \quad (20)$$

where $\|x\| = \sqrt{x^T x}$. The last assumption is to assure that

$$E \left| \sum_{0 < s \leq t} e^{i\alpha^T Z_s} (1 - e^{-i\alpha^T \Delta Y_s}) \right| < \infty \quad (21)$$

and the same for the case where the Lévy process has no negative jumps and the coordinates of α are nonnegative.

Multidimensional Reflected Lévy Process

If $\{X_t | t \geq 0\}$ is a multidimensional Lévy process and $A = (a_{ij})$ is a nonnegative matrix with spectral radius less than 1 (i.e., $A^n \rightarrow 0$ as $n \rightarrow \infty$), then $Z_t = z + X_t + (I - A)L_t$ is a multidimensional reflected Lévy process if the following are satisfied.

1. Each coordinate of $\{L_t | t \geq 0\}$ is nonnegative, nondecreasing, and right continuous.
2. Each coordinate of $\{Z_t | t \geq 0\}$ is nonnegative.
3. One (and then all) of the following equivalent conditions holds (see [6]):
 - (a) For each i and each $t \geq 0$, if $L_{t-}^i < L_s^i$ for $s > t$ then $Z_t^i = 0$.
 - (b) For each i , $\int_0^\infty Z_s^i dL_s^i = 0$.
 - (c) $\{L_t | t \geq 0\}$ is the minimal process satisfying 1 and 2.
 - (d) $\{L_t | t \geq 0\}$ is the unique process such that for each $t \geq 0$,

$$L_t^i = - \inf_{0 \leq s \leq t} \min \left(z + X_s^i - \sum_j a_{ij} L_s^j, 0 \right) \quad (22)$$

See, for example, Lemma 14.2.6 in [11].

For the case where the Lévy process is a Brownian motion, this process is an instance of a Brownian network and has been shown, under various conditions, to be the weak limit of general Jackson-type queueing networks, that is, a network of **queues** where customers upon completion of service in a given station are routed to another station or out of the network according to some lottery which is independent of the past and can change from station to station (Markovian routing). The literature related to this is vast. The earliest general result seems to be [12] and a recent one is [13] which contains an impressive bibliography.

Another application is a fluid network with Lévy input. In this model, a content which is referred to as *fluid* arrives at the various stations of the network according to a nondecreasing multidimensional Lévy process $\{J_t | t \geq 0\}$. Each station has a maximal processing rate r_i and processes fluid at the maximal rate whenever its buffer content is not zero. A fixed proportion p_{ij} of each unit of processed fluid is instantly transferred from station i to station j . It can be shown that such a model is a special case of a multidimensional reflected Lévy process with

$A = P^T$ and $X_t = J_t - (I - A)rt$. This model, under various assumptions, has been considered in the literature. For a recent reference with a rather complete set of bibliography to related literature, see [14].

Markov Additive Processes and Associated Martingales

A generalization of a Lévy process is the *Markov additive process*. See [15]. This is more or less a Lévy process subject to a Markovian random environment. When the Markovian environment changes according to a right continuous, continuous-time, finite state space Markov chain $\{E_t | t \geq 0\}$, this process has a more explicit representation.

Let $\{1, \dots, K\}$ be the states and let $\{X^i | 1 \leq i \leq K\}$ be independent right continuous Lévy processes, that is, having stationary increments with $X^i(t) - X^i(s)$ being independent of $\mathcal{F}_s$ for every $0 \leq s < t$, with $X^i(0) = 0$, each having an exponent $\varphi_i(\alpha) = \log(Ee^{i\alpha X_1})$. Assume that $\{E_t | t \geq 0\}$ is independent of $X^1, \dots, X^K$. In addition, let $\{U_n^{ij} | n \geq 1, 1 \leq i, j \leq K\}$ be independent random variables, which are also independent of $X^1, \dots, X^K$, $\{E_t | t \geq 0\}$, and $\mathcal{F}_0$, such that for every i, j, n , $\{U_n^{ij} | n \geq 1\}$ are identically distributed. We assume that $U_{ii}^{ij} \equiv 0$. Letting $\{T_i | i \geq 0\}$ be the jump epochs of $\{E_t | t \geq 0\}$ (with $T_0 = 0$), it is assumed that for every i, j, n , U_n^{ij} is $\mathcal{F}_{T_n}$ measurable. Then the Markov additive process has the following representation.

$$\begin{aligned} X_t = X_0 + \sum_{n \geq 1} \sum_{1 \leq i, j \leq K, i \neq j} (X_{T_n}^i - X_{T_{n-1}}^i + U_n^{ij}) \\ \times 1_{\{E_{T_{n-1}} = i, E_{T_n} = j, T_n \leq t\}} + \sum_{n \geq 1} \sum_{1 \leq i \leq K} (X_t^i - X_{T_{n-1}}^i) \\ \times 1_{\{E_{T_{n-1}} = i, T_{n-1} \leq t < T_n\}} \end{aligned} \quad (23)$$

where $X_0 \in \mathcal{F}_0$.

Let Q be the infinitesimal transition matrix of $\{E_t | t \geq 0\}$, $G_{ij}(\alpha)$ be the characteristic function of U_n^{ij} and $\mathbf{1}_j$ be a row vector where the j th coordinate is one and the rest are zero. Finally let $F(\alpha) = (F_{ij}(\alpha))$,

$$F_{ij}(\alpha) = q_{ij}G_{ij}(\alpha) + \varphi_i(\alpha)\delta_{ij} \quad (24)$$

where $\delta_{ii} = 1$ and $\delta_{ij} = 0$, when $i \neq j$.

In [16], the following results are shown and various applications are discussed. Note that for a

matrix A , e^A is the matrix exponential associated with A , that is, $e^A = \sum_{n=0}^{\infty} \frac{1}{n!} A^n$.

Lemma 1 *For any initial distribution of (X_0, E_0) and every real α ,*

$$M^W(t, \alpha) = e^{i\alpha X_t} \mathbf{1}_{E_t} e^{-F(\alpha)t} \quad (25)$$

is a (row) vector-valued martingale. Consequently, if $h(\alpha)$ and $\lambda(\alpha)$ are a right eigenvector and eigenvalue of the matrix $F(\alpha)$, then with every initial distribution of (X_0, E_0)

$$e^{i\alpha X_t - \lambda(\alpha)t} h_{E_t}(\alpha) \quad (26)$$

is a martingale.

For the next theorem, we extend the definition of $F(\alpha)$ to the appropriate part of the complex plane (as discussed in the theorem) and we denote $\tilde{F}(\alpha) = F(-i\alpha)$ for complex α .

Theorem 2 *Let $Y = \{Y_t | t \geq 0\}$ be an adapted continuous process having a finite expected variation on compact intervals. Set $Z_t = X_t + Y_t$ and assume either of the following conditions:*

1. X^i have no negative jumps, U_{ij} are nonnegative, $Z = \{Z_t | t \geq 0\}$ is a nonnegative process and $\Re(\alpha) \leq 0$,
2. X^i have no positive jumps, U_{ij} are nonpositive, Z is nonpositive and $\Re(\alpha) \geq 0$,
3. $\Re(\alpha) = 0$,
4. X^i are Brownian motions or linear and U_{ij} are zero (α is unrestricted).

Then, for every initial distribution of (X_0, E_0) ,

$$\begin{aligned} M(\alpha, t) = & \int_0^t e^{\alpha Z_s} \mathbf{1}_{E_s} ds \tilde{F}(\alpha) + e^{\alpha Z_0} \mathbf{1}_{E_0} - e^{\alpha Z_t} \mathbf{1}_{E_t} \\ & + \alpha \int_0^t e^{\alpha Z_s} \mathbf{1}_{E_s} dY_s \end{aligned} \quad (27)$$

is a (row) vector-valued, zero mean martingale. Consequently, with $\lambda(\alpha)$ and $h(\alpha)$ as in Lemma 1

$$\begin{aligned} \lambda(\alpha) \int_0^t e^{\alpha Z_s} h_{E_s}(\alpha) ds + e^{\alpha Z_0} h_{E_0}(\alpha) - e^{\alpha Z_t} h_{E_t}(\alpha) \\ + \alpha \int_0^t e^{\alpha Z_s} h_{E_s}(\alpha) dY_s \end{aligned} \quad (28)$$

is a zero mean martingale.

References

- [1] Protter, P. (2004). *Stochastic Integration and Differential Equations*, 2nd Edition, Springer.
- [2] Bertoin, J. (1996). *Lévy Processes*, Cambridge University Press.
- [3] Chung, K.L. (1968). *A Course in Probability Theory*, 2nd Edition, Academic Press.
- [4] Breiman, L. (1992). *Probability, Classics in Applied Mathematics Vol. 7*, SIAM.
- [5] Kella, O. & Whitt, W. (1992). Useful martingales for stochastic storage processes with Levy input, *Journal of Applied Probability* **29**, 396–403.
- [6] Kella, O. (2006). Reflecting thoughts, *Statistics and Probability Letters*, **76**(16), 1808–1811.
- [7] Bingham, N.H. (1975). Fluctuation theory in continuous time, *Advances in Applied Probability* **7**, 705–766.
- [8] Harrison, J.M. (1977). The supremum distribution of a Lévy process with no negative jumps, *Advances in Applied Probability* **9**, 417–422.
- [9] Zolotarev, V.M. (1964). The first passage time to a level and the behavior at infinity for a class of processes with independent increments, *Theory of Probability and its Applications* **9**, 653–662.
- [10] Nguyen-Ngoc, L. & Yor, M. (2004). Some martingales associated to reflected Lévy processes, in *Lecture Notes in Mathematics: Séminaire de Probabilités XXXVIII*, M. Émery, M. Ledoux & M. Yor, eds, Springer, Vol. 1857, pp. 42–69.
- [11] Whitt, W. (2002). *Stochastic Process Limits*, Springer.
- [12] Reiman, M.I. (1984). Open queueing networks in heavy traffic, *Mathematics of Operations Research* **9**, 441–458.
- [13] Gamarnik, D. & Zeevi, A. (2006). Validity of heavy traffic steady-state approximations in generalized Jackson networks, *Annals of Applied Probability* **16**, 56–90.
- [14] Dębicki, K., Dieker, A.B. & Rolski, T. (2006). Quasi-Product Forms for Lévy-Driven Fluid Networks. Preprint: <http://arxiv.org/abs/math.PR/0512119>. To appear in *Mathematics of Operations Research*.
- [15] Çinlar, E. (1972). Markov additive processes I and II. *Probability Theory and Related Fields* **24**, 85–93, 95–121 (resp.). Springer.
- [16] Asmussen, S. & Kella, O. (2000). A multi-dimensional martingale for Markov additive processes and its applications, *Advances in Applied Probability* **32**, 376–393.

OFFER KELLA

Markov Processes

Basic Definitions

A *Markov process* (MP) is a stochastic process $(X_t)_{t \in T}$ (where usually $T = \mathbb{Z}_+$ or $T = \mathbb{R}_+$) with values in some state space S that has the *Markov property*: for any $n \in \mathbb{N}$, time indices $t_1 < \dots < t_n < t$ and states $x_1, \dots, x_n \in S$ the conditional distribution of X_t given that $X_{t_1} = x_1, \dots, X_{t_n} = x_n$ coincides with that under the condition $X_{t_n} = x_n$. In other words, the probability law of the evolution of X_t after any fixed time instant s depends on its past only through the value of X_s , so that the “present state” contains all the relevant information about the future. The Markov property is equivalent to the following conditional independence property. Let $\mathcal{F}_s$ and $\mathcal{F}$ be the σ -algebras generated by $(X_t)_{t \leq s}$ and $(X_t)_{t \geq s}$, describing the past up to time s and the future starting at s , respectively. Then for all s and events $A \in \mathcal{F}_s$ and $B \in \mathcal{F}$,

$$\mathbb{P}(A \cap B \mid X_s) = \mathbb{P}(A \mid X_s) \mathbb{P}(B \mid X_s) \quad (1)$$

Depending on whether the time index set T and the state space S are discrete or continuous, one can distinguish four types of MPs. We will for simplicity only consider Borel subsets $S \subset \mathbb{R}$ and the time ranges $T = \mathbb{R}_+$ or $T = \mathbb{Z}_+$. If either the state space or the time range or both are countable, one also uses the term *Markov chain*.

The feature that the future development depends on the past only through the current state lies behind the ubiquitous occurrence of MPs in applied probability. It is analogous to the well-known property of deterministic dynamical systems that the solution of their motion equations are uniquely determined by the data (i.e., the positions and velocities of the particles involved) at one time instant. Arguably most stochastic models for real-life systems are based on underlying MPs. Important classes of MPs are random walks, birth-and-death processes, many of the processes occurring in classical queuing, dam, storage and reliability theory, Lévy processes (e.g., compound Poisson processes and **Brownian motion**), and diffusions. In this article we can give only a condensed survey of the basic mathematical machinery for the analysis of MPs. In the section titled “Discrete-Time Markov Processes” we deal with a case of discrete

time, describing in particular the classical theory for countable state space and the main results for Harris chains. The section titled “Continuous-Time Markov Processes” is devoted to the basics of the continuous-time case. In the section titled “Example: The Performance of a Replacement Policy” we present the computation of performance measures of a replacement policy in a concrete model of a failure system; this is meant as an example to show the general techniques at work. Finally, in the section titled “Continuous-Time MPs with Countable State Space” we discuss the basic constructions and analysis of continuous-time MPs with countable state space.

Among the many excellent monographs and textbooks on different levels we mention [1–4] for the general theory of MPs and [5–8] for an in-depth analysis of Markov chains with a view toward applications. See also [9, 10] for the case of discrete state space and [11] for Harris chains. We do not discuss Lévy processes [12, 13], Brownian motion [4, 14] or diffusions [1, 2, 15]. The example in the section titled “Example: The Performance of a Replacement Policy” is taken from [16]. We start with the case of discrete time.

Discrete-Time Markov Processes

If $T = \mathbb{Z}_+$, the Markov property can be formulated as

$$\mathbb{P}(X_{n+1} \in B \mid X_0, \dots, X_n) = \mathbb{P}(X_{n+1} \in B \mid X_n) \quad \text{almost surely} \quad (2)$$

for all $n \in \mathbb{Z}_+$ and all $B \in \mathcal{B}(S)$ (the Borel σ -algebra in S). For every n one can choose a stochastic kernel $p_n(x, B)$ which is a version of $\mathbb{P}(X_{n+1} \in B \mid X_n = x)$. These kernels give the *transition probabilities* of the MP. Together with the *initial distribution*,

$$\pi(B) = \mathbb{P}(X_0 \in B) \quad (3)$$

they determine the finite-dimensional distributions and thus the law of the MP: one can show that for all n and $B_0, \dots, B_n \in \mathcal{B}(S)$,

$$\begin{aligned} \mathbb{P}(X_0 \in B_0, \dots, X_n \in B_n) &= \int_{B_{n-1}} \dots \int_{B_0} \\ &\times p_{n-1}(x_{n-1}, B_{n-1}) p_{n-2}(x_{n-2}, dx_{n-1}) \\ &\dots p_0(x_0, dx_1) \pi(dx_0) \end{aligned} \quad (4)$$

2 Markov Processes

Conversely, for any sequence $p_n(x, B)$, $n \in \mathbb{Z}_+$, of stochastic kernels and any distribution π on S one can construct an MP on the infinite product space $(S^{\mathbb{N}}, \mathcal{B}(S)^{\mathbb{N}})$ of countably many factors $(S, \mathcal{B}(S))$ with the X_n s being the coordinate mappings, having the $p_n(x, B)$ as transition probabilities and π as distribution of X_0 . Let $\mathbb{P}_\pi$ and $\mathbb{E}_\pi$ be the corresponding probability measure on $(S^{\mathbb{N}}, \mathcal{B}(S)^{\mathbb{N}})$ and expectation, respectively. This standard measure theoretic construction allows to consider simultaneously a whole collection of MPs having the same transition probabilities but variable initial distributions. If π is the point mass at some x , we write $\mathbb{P}_\pi = \mathbb{P}_x$ and $\mathbb{E}_\pi = \mathbb{E}_x$ and speak of the MP starting from the point x .

In the following we will restrict ourselves to the case when $p(x, B) = p_n(x, B)$ can be chosen independently of n so that the time index n can be omitted. The MP is then called *homogeneous*; $p(x, B)$ is the probability of going from x to B in one step. For homogeneous MPs equation (2) takes the form

$$\mathbb{P}_\pi(X_{n+1} \in B \mid X_0, \dots, X_n) = \mathbb{P}_{X_n}(X_1 \in B) \quad \mathbb{P}_\pi\text{-almost surely} \quad (5)$$

For homogeneous discrete-time MPs a generalization of their defining property (5) from deterministic times n to certain nonanticipatory random times, the so-called *strong Markov property*, holds. A *stopping time* is a random variable with values in $\mathbb{Z}_+ \cup \{\infty\}$ for which $\{\tau = n\} \in \sigma(X_0, \dots, X_n)$ (the σ -algebra generated by $X_0, \dots, X_n$) for all $n \in \mathbb{Z}_+$. Then $\mathcal{F}_\tau$ is defined to be the set of all events A for which $A \cap \{\tau = n\} \in \sigma(X_0, \dots, X_n)$ for all $n \in \mathbb{Z}_+$; $\mathcal{F}_\tau$ is the σ -algebra consisting of all events whose occurrence or nonoccurrence is known at time τ . If $\tau < \infty$, the MP is at state X_τ at time τ . The strong Markov property states that the evolution of the process after time τ is conditionally independent of $\mathcal{F}_\tau$, given X_τ , and follows the same law as the process started at the state X_τ . Formally,

$$\mathbb{P}_\pi((X_\tau, X_{\tau+1}, \dots) \in C \mid \mathcal{F}_\tau) = \mathbb{P}_{X_\tau}((X_0, X_1, \dots) \in C) \quad \mathbb{P}_\pi\text{-almost surely on } \{\tau < \infty\} \text{ for every } C \in \mathcal{B}(S)^{\mathbb{N}} \quad (6)$$

For example, the strong Markov property is basic for studying the recurrence behavior of MPs, because

it implies that following any visit to some state the process behaves as if starting anew from this state.

An MP is called *asymptotically stationary* if there is a stationary process $(Y_n)_{n \in \mathbb{Z}_+}$ such that the joint distribution of the tail sequence $(X_{k+n})_{n \in \mathbb{Z}_+}$ converges to that of $(Y_n)_{n \in \mathbb{Z}_+}$ in the sense that

$$\lim_{k \rightarrow \infty} \mathbb{P}_\pi((X_{k+n})_{n \in \mathbb{Z}_+} \in C) = \mathbb{P}((Y_n)_{n \in \mathbb{Z}_+} \in C) \quad \text{for all } C \in \mathcal{B}(S)^\infty \quad (7)$$

under any initial distribution π , as $k \rightarrow \infty$. To prove asymptotic stationarity, it turns out to be sufficient to show a one-dimensional convergence result. A probability measure ρ on the state space is called *stationary initial distribution* (s.i.d.) if

$$\rho(B) = \int_S p(x, B) d\rho(x) \quad \text{for all } B \in \mathcal{B}(S) \quad (8)$$

If an MP has an s.i.d. ρ , then it is a stationary process under $\mathbb{P}_\rho$, and if it can be shown that for some probability measure ρ , $\mathbb{P}_x(X_n \in B)$ converges to $\rho(B)$ for all x and B , then the MP is asymptotically stationary (the limit being $(X_n)_{n \in \mathbb{Z}_+}$ under $\mathbb{P}_\rho$). There is at most one s.i.d. if the MP is not decomposable, i.e., if the state space cannot be partitioned into two disjoint sets A_1, A_2 such that $p(x, A_i) = 1$ for all $x \in A_i$, $i = 1, 2$. If an s.i.d. ρ exists, one can show that equation (7) holds with limit $\mathbb{P}_\rho((X_n)_{n \in \mathbb{Z}_+} \in \cdot)$ if

1. for every $n = 1, 2, \dots$, indecomposability holds under all transition probabilities $p^{(n)}(x, B) = P_x(X_n \in B)$, and
2. the distribution $p(x, \cdot)$ is absolutely continuous with respect to ρ for every x .

In general the problems of the existence of an s.i.d. and of asymptotic stationarity are difficult. A complete solution can be given for countable state spaces, whereas for uncountable state spaces there is the theory of Harris chains. We will now discuss these two special cases.

Countable State Space

In this subsection we take S to be a subset of the integers. The transition probabilities can then be represented as a matrix $P = (p_{ij})_{i,j \in S}$, where $p_{ij} = p(i, \{j\})$, whereas the initial distribution becomes a vector $\pi = (\mathbb{P}(X_0 = i))_{i \in S}$. The *n-step transition*

probabilities $p_{ij}^{(n)} = \mathbb{P}(X_n = j \mid X_0 = i)$ satisfy the Chapman–Kolmogorov equations

$$p_{ij}^{(n+1)} = \sum_{k \in S} p_{ik}^{(n)} p_{kj} \quad (9)$$

This means that the matrix $p_{ij}^{(n)}$ is just the n th power P^n of the transition matrix P . Thus, the stationarity and asymptotics of MPs can be tackled by studying the limiting behavior of the powers of stochastic matrices. If the state space is finite, this is a classical topic in matrix theory. In the general countable case one has to study the sequences of visits of the MP to its states because by the strong Markov property at each time a state i is visited, the MP restarts anew from i independently of the past.

A state i is called *recurrent* if $\mathbb{P}_i(X_n = i \text{ for some } n \geq 1) = 1$ and *transient* otherwise. Let N_i denote the number of visits in i . It can be shown that i is recurrent if and only if $\mathbb{P}_i(N_i < \infty) = 0$, which in turn is equivalent to

$$\mathbb{E}_i(N_i) = \sum_{n=1}^{\infty} p_{ii}^{(n)} = \infty \quad (10)$$

If i is recurrent, the $(\mathbb{P}_i$ -almost surely nonterminating) sequence of times of visits in i forms a renewal process, i.e., the times between visits are i.i.d. random variables. Let $\tau_i = \inf\{n \geq 1 \mid X_n = i\}$ be the time of the first visit of i (with $\inf \emptyset = \infty$). i is called *positive recurrent* if $\mathbb{E}_i(\tau_i) < \infty$, and *null recurrent* if $\mathbb{E}_i(\tau_i) = \infty$. The largest integer d for which $\mathbb{P}_i(\tau_i \in \{d, 2d, 3d, \dots\}) = 1$ is called the *period* of i , and if the period is 1, i is called *aperiodic*. Two states i and j are said to *communicate* if there are integers $n, m \geq 1$ such that $p_{ij}^{(n)} > 0$ and $p_{ji}^{(m)} > 0$. If two states communicate, they have the same period and are simultaneously either positive recurrent, or null recurrent, or transient. On the set R of recurrent states communication is an equivalence relation, so that R consists of disjoint communication classes. The states in one class have the same period d , and if $d > 1$ the class can be split into disjoint “periodic” subsets $C_1, \dots, C_d$ such that for any state $i \in C_l$ the next transition goes to some state in C_{l+1} (where $C_{d+1} = C_1$). Finally, an MP is called *irreducible* if all states communicate with each other.

On the basis of these classifications of the states a complete description of the asymptotic behavior of

the transition probabilities can be given. First, let us call a measure $\rho = (\rho_j)_{j \in S}$ *stationary* if $0 \leq \rho_j < \infty$ and $\rho_j = \sum_{i \in S} \rho_i p_{ij}$ for all $j \in S$. For any fixed recurrent state i one can define a stationary measure by setting ρ_j equal to the expected number of visits of j between two consecutive visits of i . If X_n is irreducible and recurrent, this ρ is the only stationary measure up to a multiplicative constant and $\rho_j > 0$ for all j . In particular, if we additionally assume positive recurrence, there is exactly one stationary distribution π , and it is given by $\pi_j = 1/\mathbb{E}_j(\tau_j)$. Now let us look at the limiting behavior of $p_{ij}^{(n)}$.

1. Let j be null recurrent or transient. Then $\lim_{n \rightarrow \infty} p_{ij}^{(n)} = 0$ for every state i .
2. Let j be positive recurrent with period d , and let i communicate with j . Then if $p_{ij}^{(nd+l)} > 0$ for some $n \geq 0$ and $l \in \{1, \dots, d\}$, we have $p_{ij}^{(N)} = 0$ for $N \not\equiv l \pmod{d}$ and $\lim_{n \rightarrow \infty} p_{ij}^{(nd+l)} = d/\mathbb{E}_j(\tau_j)$.
3. If j is positive recurrent and $p_{ij}^{(n)} > 0$ for some n , then either i and j communicate (so that we are in case (2)), or i is transient. Let i be transient. Then if j is aperiodic, $\lim_{n \rightarrow \infty} p_{ij}^{(n)} = r(i)/\mathbb{E}_j(\tau_j)$, where $r(i) = \mathbb{P}_i(X_n = j \text{ for some } n)$. If j has period $d > 1$ the limiting behavior depends on which (if any) of the periodic subsets $C_1, \dots, C_d$ of the communication class C of j is entered first after the start from i , and after how many steps this happens. Let $\sigma_j = \inf\{n \geq 1 \mid X_n \in C\}$. Then for all $N \geq 1$ and $l \in \{1, \dots, d\}$,

$$\begin{aligned} \lim_{n \rightarrow \infty} \mathbb{P}_i(X_{N+l+nd} = j \mid \sigma_j = N, X_{\sigma_j} \in C_{d-l+1}) \\ = d/\mathbb{E}_j(\tau_j) \end{aligned} \quad (11)$$

$$\begin{aligned} \mathbb{P}_i(X_{N+l+nd} = j \mid \sigma_j = N, X_{\sigma_j} \in C_{d-l+1}) \\ = 0 \text{ if } n \not\equiv N+l \pmod{d} \end{aligned} \quad (12)$$

The limiting behavior of $p_{ij}^{(n)}$ is then obtained by deconditioning with respect to the pair (σ_j, X_{σ_j}) .

The main result in this description is the ergodic theorem for discrete-time MPs. If X_n is irreducible, aperiodic and positive recurrent, then

$$\lim_{n \rightarrow \infty} p_{ij}^{(n)} = 1/\mathbb{E}_j(\tau_j) = \pi_j \text{ for all } i, j \in S \quad (13)$$

Note that the limit is the stationary distribution of X_n . Actually, a stronger result is valid: if all states

4 Markov Processes

communicate, are aperiodic and recurrent, we have

$$\lim_{n \rightarrow \infty} \sum_{j \in I} |\mathbb{P}_i(X_n = j) - \mathbb{P}_k(X_n = j)| = 0$$

for all $i, k \in S$ (14)

Equation (13) follows from equation (14) because for any fixed states $i, j_0 \in S$

$$\begin{aligned} 0 &= \lim_{n \rightarrow \infty} \sum_{k \in I} \pi_k \sum_{j \in I} |\mathbb{P}_i(X_n = j) - \mathbb{P}_k(X_n = j)| \\ &\geq \lim_{n \rightarrow \infty} \sum_{k \in I} \pi_k |\mathbb{P}_i(X_n = j_0) - \mathbb{P}_k(X_n = j_0)| \\ &\geq \lim_{n \rightarrow \infty} \left| \sum_{k \in I} \left(\pi_k \mathbb{P}_i(X_n = j_0) - \pi_k \mathbb{P}_k(X_n = j_0) \right) \right| \\ &= \lim_{n \rightarrow \infty} |\mathbb{P}_i(X_n = j_0) - \mathbb{P}_\pi(X_n = j_0)| \\ &= \lim_{n \rightarrow \infty} |p_{ij_0}^{(n)} - \pi_{j_0}| \end{aligned} \quad (15)$$

Continuous State Space

For a discrete-time MP X_n with state space $\mathbb{R}$ there is in most cases no simple regeneration structure, since single states are usually not revisited with positive probability. Thus one has to consider visits to subsets $R \subset \mathbb{R}$: we call R *recurrent* if $\tau_R = \inf\{n \in \mathbb{N} \mid X_n \in R\}$ satisfies $\mathbb{P}_x(\tau_R < \infty) = 1$ for all $x \in R$. A recurrent set R is said to be a *regeneration set* if for some $n \in \mathbb{N}$ the n -step transition probability measures $p^{(n)}(x, \cdot) = \mathbb{P}_x(X_n \in \cdot)$, $x \in R$, share some component, i.e., if there is an $n_0 \in \mathbb{N}$, a probability measure μ on the state space and an $\varepsilon > 0$ such that

$$p^{(n_0)}(x, \cdot) \geq \varepsilon \mu(\cdot) \quad \text{for all } x \in R \quad (16)$$

If X_n possesses a regeneration set, it is called a *Harris chain*. For given transition probabilities with a regeneration set R we can construct a corresponding Harris chain with an embedded regenerative structure as follows. First take the standard version of the MP from an arbitrary initial state x until time τ_R when R is visited for the first time. Then define $X_{\tau_R+n_0}$ by **randomization**: take a random variable which has either distribution μ with probability ε or distribution $(1 - \varepsilon)^{-1}[p^{(n_0)}(X_{\tau_R}, \cdot) - \varepsilon \mu(\cdot)]$ with probability $1 - \varepsilon$. Next, construct the missing part

$(X_{\tau_R+n})_{0 < n < n_0}$, according to the conditional distribution of $(X_n)_{0 < n < n_0}$, given that the pair of “boundary values” (X_0, X_{n_0}) is equal to the pair of values obtained before for X_{τ_R} and $X_{\tau_R+n_0}$. Finally, continue the construction from the new “initial” value $X_{\tau_R+n_0}$, and so on. The process constructed this way is an MP with transition probability measures $p(x, \cdot)$ which, due to the randomization, has distribution μ with probability ε n_0 steps after any visit to R . Let us start this MP with μ as initial distribution for X_0 . The times between two consecutive recurrences of the distribution μ obtained by the randomizations are called *cycles*.

Using this cycle decomposition, the following results can be proved for Harris chains:

1. There is exactly one stationary measure up to a positive factor; this measure ρ can be defined by setting $\rho(A)$ equal to the expected number of visits to A during a cycle started with initial distribution μ .
2. X_n is called *aperiodic* if the distribution of the length of a cycle started with μ is aperiodic. It is called *positive recurrent (null recurrent)* if the stationary measure has finite (infinite) total mass. For an aperiodic, positive recurrent Harris chain $p^{(n)}(x, \cdot)$ converges to the stationary probability measure in total variation.

For an aperiodic, null recurrent Harris chain, $\lim_{n \rightarrow \infty} p^{(n)}(x, B) = 0$ for all x and all Borel sets B with finite stationary measure.

Continuous-Time Markov Processes

In the case $T = \mathbb{R}_+$, $S \subset \mathbb{R}$ the Markov property reads as

$$\mathbb{P}(X_{t+s} \in B \mid (X_u)_{u \in [0, t]}) = \mathbb{P}(X_{t+s} \in B \mid X_t)$$

$\mathbb{P}$ -almost surely for all $t, s \geq 0$, $B \in \mathcal{B}(S)$ (17)

For all $t > s \geq 0$ one can choose a stochastic kernel $p_{s,t}(x, B)$ which is a version of $\mathbb{P}(X_t \in B \mid X_s = x)$. These kernels give the *transition probabilities* of the MP. Together with the *initial distribution*,

$$\pi(B) = \mathbb{P}(X_0 \in B) \quad (18)$$

they determine the finite-dimensional distributions and thus the law of the MP: one can show that for

all n , $0 \leq t_0 < t_1 < \dots < t_n$ and $B_0, \dots, B_n \in \mathcal{B}(S)$,

$$\begin{aligned} \mathbb{P}(X_{t_0} \in B_0, \dots, X_{t_n} \in B_n) &= \int_{B_{n-1}} \dots \int_{B_0} \\ &\times p_{t_{n-1}, t_n}(x_{n-1}, dx_n) p_{t_{n-2}, t_{n-1}}(x_{n-2}, dx_{n-1}) \\ &\dots p_{t_0, t_1}(x_0, dx_1) \pi(dx_0) \end{aligned} \quad (19)$$

The Chapman–Kolmogorov equations hold in the form

$$p_{t_0, t_1}(x, B) = \int_S p_{s, t_1}(y, B) p_{t_0, s}(x, dy) \quad (20)$$

almost surely with respect to the distribution of X_{t_0} for all $t_1 > s > t_0 \geq 0$.

Usually, the transition probabilities and the initial distribution of an MP are specified, and one wants to study the properties of an MP in these terms. Suppose a probability measure π on S and a collection of stochastic kernels $p_{s, t}(x, B)$, $t > s \geq 0$, are given. The kernels are assumed to satisfy equation (20) for all $x \in S$, $B \in \mathcal{B}(S)$ and all $t_1 > s > t_0 \geq 0$. Then one can show that the finite-dimensional distributions defined by equation (19) are consistent, so that they uniquely determine a probability measure on the product space $(S^{[0, \infty)}, \mathcal{B}(S)^{[0, \infty)})$ such that its coordinate mappings X_t form an MP having the functions $p_{s, t}(x, B)$ as transition probabilities and π as distribution of X_0 . Thus the only requirement to create an MP from transition kernels is that they satisfy the Chapman–Kolmogorov equations identically.

From now on we will again only consider homogeneous MPs: these are MPs whose transition probabilities satisfy

$$p_{s, t}(x, B) = p_{0, t-s}(x, B) \quad \text{for all } t > s \geq 0 \quad (21)$$

We write $p_t(x, B) = p_{0, t}(x, B)$. The Chapman–Kolmogorov equations now take the form

$$p_{t+s}(x, B) = \int_S p_s(y, B) p_t(x, dy) \quad (22)$$

for all $t, s > 0$ and all x and B . In the following all conditional probabilities and expectations are assumed to be obtained from the $p_t(x, B)$.

Let $\mathbb{P}_\pi$ and $\mathbb{E}_\pi$ be the corresponding probability measure and expectation on $(S^{[0, \infty)}, \mathcal{B}(S)^{[0, \infty)})$ constructed from the $p_t(x, B)$ and the initial distribution π . If π is the point mass at some x , we write

$\mathbb{P}_\pi = \mathbb{P}_x$ and $\mathbb{E}_\pi = \mathbb{E}_x$ and speak of the MP starting from the point x . We now consider only this coordinate representation process (so that the measurability of sets and functions will not depend on the selection of a starting point) simultaneously for variable initial distributions.

By equation (22), the transition probabilities $p_t(x, B)$ are completely determined for all $t > 0$ by their values for t in an arbitrarily small time interval $(0, \varepsilon)$, $\varepsilon > 0$. Thus it seems that the distribution of an MP can be specified by their infinitesimal behavior at zero. From now on we only consider MPs for which $p_t(x, (x - \varepsilon, x + \varepsilon)) \rightarrow 1$ for all $\varepsilon > 0$ as $t \rightarrow 0$, i.e., $X_t \rightarrow X_0$ in probability. The *generator* of an MP is defined to be the operator $\mathcal{G}$, given by

$$(\mathcal{G}f)(x) = \lim_{t \rightarrow 0} t^{-1} [\mathbb{E}_x(f(X_t)) - f(x)] \quad (23)$$

whose domain $D_{\mathcal{G}}$ is the set of all bounded measurable functions f on S for which the limit in equation (23) exists for all $x \in S$. For example, if X_t is a compound **Poisson process** with rate λ and jump size distribution function F , then $D_{\mathcal{G}}$ is the set of all bounded measurable functions on S and

$$(\mathcal{G}f)(x) = \lambda \int_{\mathbb{R}} [f(x+u) - f(x)] dF(u) \quad (24)$$

whereas if X_t is standard Brownian motion,

$$(\mathcal{G}f)(x) = \frac{1}{2} f''(x) \quad (25)$$

for all bounded $f : \mathbb{R} \rightarrow \mathbb{R}$ having a bounded continuous second derivative (in this case it is not obvious what $D_{\mathcal{G}}$ is).

In general, the generator $(\mathcal{G}, D_{\mathcal{G}})$ does not determine the transition probabilities uniquely. However, it leads to Kolmogorov's so-called *backward equations* and *forward equations*, which are valid under smoothness conditions on $p_t(x, B)$. Setting $f_{t, B}(x) = p_t(x, B)$, the backward equations are

$$(\mathcal{G}f_{t, B})(x) = \frac{\partial f_{t, B}}{\partial t}(x) \quad (26)$$

The forward equations are

$$\int (\mathcal{G}f)(u) p_t(x, du) = \int f(u) \frac{\partial}{\partial t} p_t(x, du), \quad f \in D_{\mathcal{G}} \quad (27)$$

Uniqueness can be achieved by restricting the domain of $\mathcal{G}$ to the set $\overline{D}_{\mathcal{G}}$ of all bounded measurable functions on S for which (a) $\lim_{t \rightarrow 0} \mathbb{E}_x(f(X_t)) = f(x)$ for all $x \in S$ and (b) the ratio on the right side of equation (23) converges to $(\mathcal{G}f)(x)$ and is uniformly bounded in x for sufficiently small t . Then if two MPs have the same generator $\mathcal{G}$ with the same restricted domain $\overline{D}_{\mathcal{G}}$, they have the same law.

The strong Markov property does not hold in general for MPs in continuous time. It is valid if the stopping time takes only countably many values. Arbitrary stopping times can be approximated by those with countably many values under regularity conditions on the sample paths, and the strong Markov property can be carried over to more general cases. For example, it holds if X_t is a Feller process and has right-continuous paths. An MP with state space $\mathbb{R}$ is called a *Feller process* if for every continuous function f on the state space satisfying $\lim_{|x| \rightarrow \infty} f(x) = 0$ the function $x \mapsto \mathbb{E}_x(f(X_t))$ has the same properties and $\lim_{t \rightarrow 0} \mathbb{E}_x(f(X_t)) = f(x)$ for all x . For example, Lévy processes, diffusions, and Markov jump processes are Feller processes. Every Feller process has a version that has right-continuous paths with left-hand limits. A very useful property of Feller processes is that for any $f \in D_{\mathcal{G}}$ the process

$$f(X_t) - \int_0^t (\mathcal{G}f)(X_s) ds$$

is a martingale. *Dynkin's formula* states that

$$\mathbb{E}_x(f(X_\tau)) = f(x) + \mathbb{E}_x \left(\int_0^\tau (\mathcal{G}f)(X_s) ds \right) \quad (28)$$

for every stopping time τ for which $\mathbb{E}_x(\tau) < \infty$. In the next section we present a typical example of how these methods can be employed in stochastic modeling.

Example: The Performance of a Replacement Policy

Consider a technical system subject to fatal failure at some random time σ . The state of the system X_t , $t \geq 0$, is described by a real number indicating its extent of attrition. The process $(X_t)_{t \geq 0}$ starts at $X_0 = 0$

and increases linearly at rate $c(J_t)$, where $(J_t)_{t \geq 0}$ is a “modulating” irreducible Markov chain having state space $\{1, \dots, n\}$ and the rates $c(j)$ satisfy $c(1) > c(2) > \dots > c(n) \geq 0$. The failure rate $r(x)$ is assumed to be a function of the state variable x . The system has to run indefinitely; upon failure it has to be replaced by a new identical one. However, it can also be replaced preventively when its attrition reaches a certain threshold $a > 0$ which has to be specified by the controller. Thus, setting $T_a = \inf\{t > 0 \mid X_t = a\}$, the first replacement takes place at time $\min[T_a, \sigma]$. Suppose that after replacement the modulating chain is also restarted at some fixed state i_0 . We want to determine the performance of the replacement strategy T_a as a function of a .

The underlying MP of this system is two dimensional: $Z_t = (X_t, J_t)$. The generator of Z_t is of course defined analogously to the section titled “Continuous-Time Markov Processes” by $(\mathcal{G}f)(x, i) = \lim_{t \rightarrow 0} t^{-1} [\mathbb{E}_{x,i}(f(Z_t)) - f(x, i)]$ for functions $f: \mathbb{R}_+ \times \{1, \dots, n\} \rightarrow \mathbb{R}$. It can be shown that

$$\begin{aligned} (\mathcal{G}f)(x, i) &= c(i)f'(x, i) + \sum_{j \neq i} q_{ij}f(x, j) \\ &\quad - (q_i + r(x))f(x, i) \\ &\quad + r(x) \int_{[0,x)} f(y, i) \mu_x(dy), \\ i &= 1, \dots, n \end{aligned} \quad (29)$$

or in matrix form

$$\begin{aligned} (\mathcal{G}f)(x) &= Cf'(x) + (Q - r(x)E)f(x) \\ &\quad + r(x)D_f(x) \end{aligned} \quad (30)$$

where $f'(x, i)$ denotes the derivative of f with respect to x and

1. $f(x) = (f(x, 1), \dots, f(x, n))^t$;
2. $Q = (q_{ij})_{i,j \in \{1, \dots, n\}}$ is the generator of the Markov chain J_t and $q_i = -q_{ii}$;
3. C and $D_f(x)$ are diagonal matrices with diagonal entries $c(i)$ and $\int_{[0,x)} f(y, i) \mu_x(dy)$, respectively;
4. E is the $n \times n$ identity matrix.

Now assume a replacement of a still functioning system costs C_1 and a replacement upon disaster costs C_2 (where $C_1 < C_2$). Then the long-run average cost

of running the system when using the replacement policy T_a is given by

$$C(a) = \frac{C_1 \mathbb{P}_{(0,i_0)}(\min[T_a, \sigma] = T_a) + C_2(1 - \mathbb{P}_{(0,i_0)}(\min[T_a, \sigma] = T_a))}{\mathbb{E}_{(0,i_0)}(\min[T_a, \sigma])} \quad (31)$$

Hence, we have to compute

$$\mathbb{P}_{(x,i)}(\min[T_a, \sigma] = T_a) \quad \text{and} \quad \mathbb{E}_{(x,i)}(\min[T_a, \sigma]).$$

Once these quantities are known as *functions of a* one can try to minimize $C(a)$.

This problem can be tackled as follows. Suppose the process is killed at time σ by entering into a “coffin state” ∂ . By Dynkin’s formula, we have

$$f(x, i) + \mathbb{E}_{(x,i)}\left(\int_0^T (\mathcal{G}f)(Z_t) dt\right) = \mathbb{E}_{(x,i)}(f(Z_T)) \quad (32)$$

for f bounded and in the domain of $\mathcal{G}$ and T any integrable stopping time. Now apply equation (32) to $T = \min[T_a, \sigma]$ in two cases:

1. $f = f_1$ such that $(\mathcal{G}f_1)(x, i) = 0$, $x \in (0, a)$, and $f_1(\partial) = 0$, $f_1(a, i) = 1$.
2. $f = f_2$ such that $(\mathcal{G}f_2)(x, i) = -1$, $x \in (0, a)$, and $f_2(\partial) = 0$, $f_2(a, i) = 0$.

A moment’s reflection shows that if f_1, f_2 have these properties, then

$$f_1(x, i) = \mathbb{P}_{(x,i)}(\min[T_a, \sigma] = T_a) \quad (33)$$

$$f_2(x, i) = \mathbb{E}_{(x,i)}(\min[T_a, \sigma]) \quad (34)$$

Hence, we have to solve

$$Cf'(x) + (Q - r(x)E)f(x) = (0, \dots, 0)^t \quad (35)$$

subject to

$$f(a, \cdot) = (1, \dots, 1)^t \quad (36)$$

and

$$Cf'(x) + (Q - r(x)E)f(x) = (-1, \dots, -1)^t \quad (37)$$

subject to

$$f(a, \cdot) = (0, \dots, 0)^t \quad (38)$$

These are linear systems of differential equations that can be dealt with using standard methods. For

special failure rate functions $r(x)$ one can even obtain closed-form solutions and thus determine the long-run average cost $C(a)$ explicitly.

Continuous-Time MPs with Countable State Space

Let the state space of X_t be a set of integers. Then the transition probabilities are given by the set of functions $p_{ij}(t) = \mathbb{P}_i(X_t = j)$, $t > 0$, where $i, j \in S$, which have the following properties:

1. $p_{ij}(t) \geq 0$, $\sum_{j \in S} p_{ij}(t) = 1$;
2. $p_{ij}(t + s) = \sum_{k \in S} p_{ik}(t)p_{kj}(s)$;
3. $\lim_{t \rightarrow 0} p_{ij}(t) = \delta_{ij}$.

It follows from 1 to 3 that the limits $\lim_{t \rightarrow 0} t^{-1} p_{ij}(t) = \lambda_{ij}$, $i \neq j$, and $\lim_{t \rightarrow 0} t^{-1} [1 - p_{ii}(t)] = \lambda_i$ exist, but the latter ones may be infinite. If λ_i is finite for some i , the functions $p_{ij}(t)$, $j \in S$, are continuously differentiable on $(0, \infty)$ and we have the linear approximations

$$p_{ij}(t) = \lambda_{ij}t + o(t), \quad j \neq i \quad (39)$$

$$p_{ij}(t) = 1 - \lambda_i t + o(t) \quad (40)$$

as $t \rightarrow 0$. Thus, λ_{ij} can be interpreted as the transition rate from i to j and λ_i as the total transition rate away from i . From now on we assume that $\sum_{j \neq i} \lambda_{ij} = \lambda_i < \infty$ for all states i . Then the backward equations are valid in the form

$$p'_{ij}(t) = -\lambda_i p_{ij}(t) + \sum_{k \neq i} \lambda_{ik} p_{kj}(t) \quad (41)$$

The forward equations are

$$p'_{ij}(t) = -\lambda_j p_{ij}(t) + \sum_{k \neq j} \lambda_{kj} p_{ik}(t) \quad (42)$$

Equation (42) does not hold in general; its derivation additionally requires that

$$\lim_{t \rightarrow 0} \sup_{i \neq j} |t^{-1} p_{ij}(t) - \lambda_{ij}| = 0 \quad (43)$$

a uniformity property that is satisfied in most concrete models. Let Λ be the matrix with entries λ_{ij} , $i \neq j$, and $-\lambda_i$ on the diagonal, and let $P(t)$ and $P'(t)$ be the matrices with entries $p_{ij}(t)$ and $p'_{ij}(t)$, respectively.

8 Markov Processes

Then equation (41) and equation (42) can be written as the linear matrix differential equations

$$P'(t) = P(t)\Lambda \quad (44)$$

$$P'(t) = \Lambda P(t) \quad (45)$$

with the initial condition $P(0) = (\delta_{ij})$. A formal solution to equation (44) and equation (45) is

$$P(t) = e^{\Lambda t} \quad (46)$$

If the state space is finite, equation (46) is the unique solution of equation (41) (and of equation (42)), and when Λ can be diagonalized, $e^{\Lambda t}$ can be expressed in terms of eigenvectors and eigenvalues. The case of an infinite state space is more complicated.

In many stochastic systems transition probabilities for the successively visited states and total transition rates λ_i are given by the intuitive formulation of the model. The standard construction of MPs with countable state space starts from an arbitrary discrete-time MP $(Y_n)_{n \in \mathbb{Z}_+}$, giving the sequence of visited states, and arbitrary non-negative sojourn time parameters λ_i for the holding times (which turn out to be exponential). Let $Q = (q_{ij})_{i,j \in S}$ be a transition matrix satisfying $q_{ii} = 0$ for all i . We introduce an extra point $\partial \notin S$ to account for possible “explosions” of the MP in finite time. Let $S_\partial = S \cup \{\partial\}$ and extend Q to a transition matrix Q_∂ on S_∂ by making ∂ an absorbing state, i.e., setting $q_{\partial,\partial} = 1$. Define the underlying sample space by

$$\Omega = (0, \infty]^{\mathbb{N}} \times (S_\partial)^{\mathbb{N}} \quad (47)$$

and let $T_0, T_1, \dots : \Omega \rightarrow (0, \infty]$ and $Y_0, Y_1, \dots : \Omega \rightarrow S_\partial$ be the coordinate projections. Then one can, for any probability vector $\pi = (\pi_i)_{i \in S}$, construct probability measures $\mathbb{P}_\pi$ on Ω , endowed with its product σ -algebra, such that, under $\mathbb{P}_\pi$:

1. $(Y_n)_{n \in \mathbb{Z}_+}$ is a discrete-time MP on S_∂ governed by Q_∂ and satisfying $\mathbb{P}_\pi(Y_0 = i) = \pi_i$.
2. Conditional on $(Y_n)_{n \in \mathbb{Z}_+}$, the random variables $T_0, T_1, \dots$, are independent and T_n is exponential with rate λ_{Y_n} .

Then

$$X_t = \begin{cases} Y_n, & \text{if } T_0 + \dots + T_{n-1} \leq t < T_0 + \dots + T_n, \\ n \in \mathbb{Z}_+ \\ \partial, & \text{if } t \geq \sup_{n \in \mathbb{Z}_+} (T_0 + \dots + T_n) \end{cases} \quad (48)$$

defines an MP which runs through the states $Y_0, Y_1, \dots$ and has for every visit to state i an exponential sojourn time (with mean λ_i^{-1}) which is independent of its history before and after the current visit. The rates λ_{ij} in equation (39) are given by $\lambda_{ij} = \lambda_i q_{ij}$, $i \neq j$.

This process can be *explosive*, i.e. $\mathbb{P}_i(M < \infty) > 0$ for some i , where $M = \sup_{n \in \mathbb{Z}_+} (T_0 + \dots + T_n)$. According to definition (48), it will stay in ∂ forever after time M , but there is an infinity of other possibilities to define the process after an explosion. The MP constructed by equation (48) is called *minimal* as it minimizes $\mathbb{P}_\pi(X_t = i)$ for all $i \in S$, $t > 0$ and π . Instead of introducing an absorbing extra state, one could define that at every explosion time the process moves to state i with some probability $\alpha_i \geq 0$, where $\sum_{i \in S} \alpha_i = 1$ (in this case one does not need ∂). Formally, starting with M and X_t from the above construction,

$$p_{ij}^{(0)}(t) = \mathbb{P}_i(X_t = j, M > t) \quad (49)$$

is the probability to go from i to j before the first explosion. Now define recursively the probability to go from i to j in time t with exactly n explosions in between by setting

$$p_{ij}^{(n)}(t) = \int_0^t \sum_{k \in S} p_{ik}^{(n-1)}(t-s) \alpha_k \mathbb{P}_k(M \in ds), \quad n \geq 1 \quad (50)$$

and let

$$p_{ij}(t) = \sum_{n=0}^{\infty} p_{ij}^{(n)}(t) \quad (51)$$

It can be shown that these $p_{ij}(t)$ define a family of transition probability functions which satisfy the backward but not the forward equations.

The MP constructed by equation (48) is nonexplosive (that is, $\mathbb{P}_i(M < \infty) = 0$ for every $i \in S$) if and only if the vector equation $\Lambda a = a$, $a = (a_i)_{i \in S}$ admits only the trivial solution $a = 0$. Sufficient conditions to rule out explosions are: (a) $\sup_{i \in S} \lambda_i < \infty$ or (b) Y_n is recurrent.

The asymptotic behavior of X_t is slightly simpler than in the discrete-time case because periodicity phenomena are smoothed out by the exponential sojourn times. In the following we consider a minimal MP X_t and summarize the main results. X_t is called

irreducible if $p_{ij}(t) > 0$ for some $t > 0$ for all $i, j \in S$ (this is equivalent to the irreducibility of Y_n). A state i is called *recurrent (transient)* if it is recurrent (transient) for Y_n , and the MP itself is said to be recurrent (transient) if all states have this property.

1. If X_t is irreducible and recurrent, it possesses a stationary measure $\rho = (\rho_i)_{i \in S}$, i.e., $\rho \neq 0$, $0 \leq \rho_i < \infty$ and $\rho_i = \sum_{j \in S} \rho_j p_{ji}(t)$ for all $i \in S$ and $t > 0$. Such a ρ is unique up to positive factors and is given by $\tilde{\rho}_i = \mu_i / \lambda_i$, where $(\tilde{\rho}_i)_{i \in S}$ is stationary for Y_n .
2. Alternatively, ρ_i can be characterized (up to a constant factor) as the expected time spent in i between two consecutive visits of j , where j is an arbitrary fixed state, or as solution of the equation $\rho \Lambda = 0$.
3. If the stationary measure has finite total mass, the MP is called *ergodic*. If the MP is irreducible and nonexplosive, it is ergodic if and only if there is a stationary probability distribution; for an irreducible MP to be ergodic it is sufficient that there exists a probability vector $\pi = (\pi_i)_{i \in S}$ solving $\pi \Lambda = 0$ and satisfying $\sum_i \pi_i \lambda_i < \infty$. In this case $\lim_{t \rightarrow \infty} p_{ij}(t) = \pi_j$ for all states i and j . If X_t is irreducible and recurrent but not ergodic, the limits are all equal to zero.

Spelled out, the vector equation $\pi \Lambda = 0$ is tantamount to $\pi_j \lambda_j = \sum_{i \neq j} \rho_i \lambda_{ij}$ for all j ; this means that in steady state the outflow rate of any state is equal to its inflow rate.

References

- [1] Dynkin, E.B. (1965). *Markov Processes*, Springer, Vols 1–2, English translation.
- [2] Breiman, L. (1968). *Probability*, Addison-Wesley.
- [3] Ethier, S.N. & Kurtz, T.G. (1986). *Markov Processes: Characterization and Convergence*, John Wiley & Sons.
- [4] Kallenberg, O. (2001). *Foundations of Modern Probability*, 2nd Edition, Springer.
- [5] Feller, W. (1968, 1971). *An Introduction to Probability Theory and its Applications*, Vol. 1 (3rd Edition), Vol. 2 (2nd Edition), John Wiley & Sons.
- [6] Karlin, S. & Taylor, H.M. (1975). *A First Course in Stochastic Processes*, 2nd Edition, Academic Press.
- [7] Karlin, S. & Taylor, H.M. (1981). *A Second Course in Stochastic Processes*, Academic Press.
- [8] Asmussen, S. (2003). *Applied Probability and Queues*, 2nd Edition, Springer.
- [9] Chung, K.L. (1960). *Markov Chains with Stationary Transition Probabilities*, Springer.
- [10] Freedman, D. (1971). *Markov Chains*, Holden-Day.
- [11] Meyn, S.P. & Tweedie, R.L. (1993). *Markov Chains and Stochastic Stability*, Springer.
- [12] Bertoin, J. (1996). *Lévy Processes*, Cambridge University Press.
- [13] Sato, K. (1999). *Lévy Processes and Infinitely Divisible Distributions*, Cambridge University Press.
- [14] Revuz, P. & Yor, M. (1999). *Continuous Martingales and Brownian Motion*, 3rd Edition, Springer.
- [15] Rogers, L.C.G. & Williams, D. (2000). *Diffusions, Markov Processes, and Martingales*, Cambridge University Press, Vols 1–2.
- [16] Boxma, O., Perry, D., Stadjé, W. & Zacks, S. (2006). A Markovian growth-collapse model, *Advances in Applied Probability* **76**, 221–243.

Further Reading

Karatzas, I. & Shreve, S.E. (1991). *Brownian Motion and Stochastic Calculus*, 2nd Edition, Springer.

WOLFGANG STADJE

Renewal Theory

Basic Definitions

A *renewal process* (RP) is a random sequence $0 \leq S_0 < S_1 < S_2 < \dots$, interpreted as the occurrence times of some phenomenon, for which the interarrival times $X_n = S_n - S_{n-1}$, $n \geq 1$, and $X_0 = S_0$ are independent and $X_1, X_2, \dots$ (but not necessarily X_0) have the same cumulative distribution function (cdf) F . It provides a mathematical model for a system that is completely renewed at times $S_0, S_1, \dots$. The system could for example consist of a single item with lifetime distribution F , which is replaced by a new item of the same kind at every failure time. The item functioning at time zero may not be new so that its residual lifetime S_0 could have a different cdf, say G . If $S_0 = 0$ almost surely (i.e., G corresponds to the point mass at zero), the RP is called *pure*; otherwise it is called *delayed*.

An RP can be equivalently described by its corresponding counting process

$$\begin{aligned} N_t &= \#\{n \geq 0 \mid S_n \leq t\} \\ &= \inf\{n \geq 0 \mid S_n > t\}, \quad t \geq 0 \end{aligned} \quad (1)$$

which gives the number of renewals in $[0, t]$. Note that $N_{S_n} = n + 1$ and that N_t and S_n are related by

$$\{N_t \leq n\} = \{S_n > t\} \quad (2)$$

so that for fixed t the distribution of N_t can be expressed in terms of the n -fold convolutions F^{*n} of F with itself and of G . Renewal theory has as its core the study of the process $(N_t)_{t \geq 0}$. As the basic theory of recurrent events, it is a universal mathematical tool of stochastic modeling. Extensive treatments with a wealth of applications can be found in [1–4]. We will focus on the standard framework and not consider extensions and ramifications like **Markov RPs** or nonlinear or multidimensional renewal theory. It can be shown that $\mathbb{E}[\exp\{aN_t\}] < \infty$ for some $\alpha > 0$; this implies that all **moments** of N_t are finite. The main results of renewal theory concern the convolution series

$$U(t) = \sum_{n=0}^{\infty} F^{*n} \quad (3)$$

the so-called renewal function. Here F^{*n} denotes the n -fold convolution of F with itself and $F^{*0} = 1_{[0, \infty)}$. For a pure RP, $U(t)$ is equal to $\mathbb{E}[N_t]$, the expected number of renewals in $[0, t]$.

Explicit formulas for $U(t)$ are available in a few cases; for example if F is uniform [1, 5] Erlang [6], truncated exponential [7], or Weibull [8]; for the phase-type case see [3]. The counting process of a pure RP with $\exp(\lambda)$ -distributed interarrival times is of course nothing but a homogeneous **Poisson process**; in this case $U(t) = 1 + \lambda t$.

The Renewal Equation and the Renewal Theorem

In this section we consider a pure RP. The function $U(t)$ satisfies the equation

$$U(t) = 1 + \int_0^t U(t-x) dF(x), \quad t \geq 0 \quad (4)$$

To see equation (4), condition N_t on the first renewal time S_1 and note that $\mathbb{E}[N_t \mid S_1 = x] = 1$ for $t < x$ and $\mathbb{E}[N_t \mid S_1 = x] = 1 + U(t-x)$ for $t \geq x$.

Equation (4) is a special case of the *renewal equation*

$$Z(t) = z(t) + \int_0^t Z(t-x) dF(x) \quad (5)$$

In equation (5) the function $z(t)$ and the cdf F are given and $Z(t)$ is an unknown function on $[0, \infty)$. If z is bounded on every finite interval, it can be shown that

$$Z(t) = \int_0^t z(t-x) dU(x) \quad (6)$$

is a solution of equation (5), and indeed the only one that is bounded on finite intervals.

The *renewal theorem* describes the asymptotic behavior of $U(t)$ and of the solutions of corresponding renewal equations. Set $\mu = \mathbb{E}[X_1] \in (0, \infty]$. Intuitively, for $\mu < \infty$ one expects $U(t)$ to grow like t/μ because, on an average, there will be a renewal every μ time units, and indeed it is not difficult to show that

$$\lim_{t \rightarrow \infty} t^{-1} U(t) = \mu^{-1} \quad (7)$$

The renewal theorem extends this idea in a very strong sense. One has to distinguish two cases: (a) If F is concentrated on the multiples $\{a, 2a, 3a, \dots\}$

2 Renewal Theory

of some number $a > 0$, it is said to be *lattice*, and the largest a with this property is called its *period*; (b) if no such a exists, F is called *nonlattice*.

Blackwell's renewal theorem states that if F is nonlattice, then for every $s > 0$,

$$U(t + s) - U(t) \longrightarrow \frac{s}{\mu}, \text{ as } t \longrightarrow \infty \quad (8)$$

while if F is lattice with period a , then

$$U(na) - U((n-1)a) \longrightarrow \frac{a}{\mu}, \text{ as } n \longrightarrow \infty \quad (9)$$

If $\mu = 0$, the limits have to be interpreted as zero, so that we also have a result for this case.

The *key renewal theorem* concerns the solutions of equation (5) for a certain class of functions $z : [0, \infty) \rightarrow \mathbb{R}$. It states that for nonlattice F and any directly Riemann integrable (DRI) z ,

$$Z(t) = \int_0^t z(t-x) dU(x) \longrightarrow \frac{1}{\mu} \int_0^\infty z(x) dx, \text{ as } t \longrightarrow \infty \quad (10)$$

Here a nonnegative function z is called *DRI* if its upper and lower Riemann sums over the entire interval $[0, \infty)$ converge properly. Specifically, let

$$\begin{aligned} M_\alpha^{(n)}(z) &= \sup_{n\alpha \leq x < (n+1)\alpha} z(x), \quad m_\alpha^{(n)}(z) \\ &= \inf_{n\alpha \leq x < (n+1)\alpha} z(x) \end{aligned} \quad (11)$$

for $\alpha > 0$ and $n = 0, 1, 2, \dots$. Then $z \geq 0$ is DRI if

$$\lim_{\alpha \rightarrow 0} \alpha \sum_{n=0}^{\infty} (M_\alpha^{(n)}(z) - m_\alpha^{(n)}(z)) = 0 \quad (12)$$

A real-valued function z on $[0, \infty)$ is DRI if its negative and its positive parts are DRI. For z to be DRI, it is necessary that it is bounded and continuous Lebesgue almost everywhere. Sufficient conditions are (a) z is nonincreasing and Lebesgue integrable or (b) z is continuous Lebesgue almost everywhere and $0 \leq z \leq \tilde{z}$ for some function $\tilde{z}$ satisfying (a).

A direct analytical proof of the renewal theorem can e.g., be found in [1, 3]. A Fourier analytical derivation is given in [9, 10]. An innovative probabilistic approach that has also led to new proofs of the renewal theorem and extensions is the coupling method [11, 12].

Refinements and Applications

We now give a survey of classical results linked with the renewal theorem and consider a pure nonlattice RP.

1. The *age process* and the *residual lifetime process* corresponding to an RP are defined by

$$A_t = t - S_{N_t-1}, \quad R_t = S_{N_t} - t, \quad t \geq 0 \quad (13)$$

They give the time elapsed since the last renewal and the waiting time until the next renewal, respectively. A_t and B_t are strong **Markov processes**. The cdf's $\mathbb{P}(A_t \leq x)$ and $\mathbb{P}(B_t \leq x)$ of A_t and B_t , considered as functions of t for fixed x , both satisfy the renewal equation (5) (with $z(t) = 1_{[0,x]}(t)(1 - F(t))$ and $z(t) = F(t+x) - F(t)$, respectively). Applying the key renewal theorem yields that if $\mu < \infty$ they both have the same limiting stationary cdf,

$$\lim_{t \rightarrow \infty} \mathbb{P}(A_t \leq x) = \lim_{t \rightarrow \infty} \mathbb{P}(B_t \leq x) = F_e(x) \quad (14)$$

which is given by

$$F_e(x) = \frac{1}{\mu} \int_0^x (1 - F(u)) du \quad (15)$$

If the second moment $\mathbb{E}(X_1^2)$ is finite, the expected values $\mathbb{E}(A_t)$ and $\mathbb{E}(B_t)$ converge to $\int_0^\infty x dF_e(x) = \mathbb{E}(X_1^2)/2\mu$, and since $N_t = t + B_t$, it follows that in this case

$$\lim_{t \rightarrow \infty} \left(U(t) - \frac{t}{\mu} - \frac{\mathbb{E}(X_1^2)}{2\mu} \right) = 0 \quad (16)$$

2. The *inspected lifetime* L_t , $t \geq 0$, is defined by $L_t = X_{N_t} = A_t + R_t$. It gives the lifetime of the item active at time t . It can be shown that its limiting cdf is $F_0(x) = \int_0^x (u/\mu) dF(u)$. The fact that the cdf of the inter-renewal times containing a fixed time t is different from F is called the *inspection paradox*. The reason is that sampling the inter-renewal times containing a fixed t favors the larger ones. In the limit, the expected inspected lifetime is $\mathbb{E}[X_1^2]/\mu$, which is greater than μ . In fact, F_0 is stochastically larger than F . One can even prove [13] that the backward sequence $X_{N_t}, X_{N_t-1}, X_{N_t-2}, \dots$ of renewal epochs up to $t+0$ is stochastically decreasing, i.e.

$$\begin{aligned} \mathbb{P}(X_{N_t-n} \leq x) &\geq \mathbb{P}(X_{N_t-m} \leq x) \text{ for all } x \geq 0 \\ &\text{and all } n \geq m \geq 0 \end{aligned} \quad (17)$$

where we set $X_k = 0$ for $k < 0$.

3. If F is absolutely continuous with density f , there is a *renewal density* $u(t)$, i.e. $U(t) = \int_0^t u(s) ds$. It is given by $u = \sum_{n=1}^{\infty} f^{*n}$ and satisfies the renewal equation (5) with $z = f$. In view of Blackwell's renewal theorem one expects the corresponding result for the density:

$$\lim_{t \rightarrow \infty} u(t) = \frac{1}{\mu} \quad (18)$$

Conditions on f for equation (18) to hold were given in [1, 14–18]; the simplest are (a) $\lim_{t \rightarrow \infty} f(t) = 0$ and $\int_0^{\infty} f(x)^{1+\delta} dx < \infty$ for some $\delta > 0$ or (b) f is DRI. If f is just bounded, we have [1]

$$\lim_{t \rightarrow \infty} (u(t) - f(t)) = \frac{1}{\mu} \quad (19)$$

which means that the renewal density asymptotically compensates the possible oscillations of f . Equation (19) also holds if f is unbounded but $\int_0^{\infty} f(x)^2 dx < \infty$ and $\int_0^{\infty} f(x)x \log(1+x) dx < \infty$, and if the second moment of F is finite, we even have

$$u(t) - f(t) - \mu^{-1} = O(t^{-1/2}) \quad (20)$$

see [19, 20]. Further estimates for the rate of convergence and exact upper bounds can be found in [19, 21, 22].

4. The convergence results above can be considerably sharpened if a renewal density exists. Actually only a somewhat weaker condition is needed. The cdf F is called *spread out* if $F^{*n} \geq \varepsilon H$ for some $n \in \mathbb{N}$, $\varepsilon > 0$ and some absolutely continuous cdf H . If F has this property, then *Stone's decomposition* holds: the renewal function U can be written as $U(t) = U_1(t) + U_2(t)$, where U_1 and U_2 are non-decreasing, right-continuous functions on $[0, \infty)$, U_1 has a bounded continuous density $u_1 = dU_1/dt$ satisfying $\lim_{t \rightarrow \infty} u_1(t) = 1/\mu$, and U_2 is bounded. It can then be concluded that the key renewal theorem holds for every Lebesgue integrable bounded function z for which $\lim_{t \rightarrow \infty} z(t) = 0$ (so that z need not be DRI). Moreover, A_t and R_t converge to their common limit F_e in total variation (not just in distribution as in equation (14)). One gets exponential rates of convergence for the total variation distances if F is spread out and additionally $\mathbb{E}[e^{\alpha X_1}] < \infty$ for some $\alpha > 0$. In this situation, if z is bounded and of order $O(e^{-\beta t})$ as $t \rightarrow \infty$, the solution of the corresponding renewal equation converges to $\mu^{-1} \int_0^{\infty} z(t) dt$ exponentially fast.

Delayed Renewal Processes

The results of the sections titled “The Renewal Equation and the Renewal Theorem” and “Refinements and Applications” essentially carry over to delayed RPs. The expected number of renewals in $[0, t]$ is given by

$$\mathbb{E}[N_t] = (G * U)(t) \quad (21)$$

If F and G are lattice with period a , then

$$\mathbb{E}[N_{na}] - \mathbb{E}[N_{(n-1)a}] \longrightarrow \frac{a}{\mu}, \text{ as } n \longrightarrow \infty \quad (22)$$

If F is nonlattice,

$$\mathbb{E}[N_{t+s}] - \mathbb{E}[N_t] \longrightarrow \frac{s}{\mu} \text{ for every } s > 0 \quad (23)$$

and for every DRI function z ,

$$\int_0^t z(t-x) d(G * U)(x) \longrightarrow \frac{1}{\mu} \int_0^{\infty} z(x) dx, \quad \text{as } t \longrightarrow \infty \quad (24)$$

By choosing the “right” cdf for the initial renewal instant S_0 one can make the delayed RP become stationary in the sense that the distribution of the shifted process $(N_{t+s} - N_s)_{t \geq 0}$ is independent of $s \geq 0$, i.e., it is equal to that of $(N_t)_{t \geq 0}$. For any $t \geq 0$, after S_{N_t} (the first S_n in $[t, \infty)$) the RP continues with independent and identically distributed (i.i.d.) inter-renewal times. Thus, if the residual lifetime $R_t = S_{N_t} - t$ has the same distribution for every t , the delayed RP is stationary. If $\mu < \infty$, this can be achieved by assigning to $R_0 = S_0$ the limiting distribution of B_t , having cdf F_e (cf. equation (14)). For this reason F_e is called the *equilibrium distribution* of the RP. In this case $\mathbb{E}[N_t] = t/\mu$. In fact there is a nice construction of a stationary RP with interarrival cdf F . Let L and U be independent random variables where L has cdf F_0 and U is uniform on $(0, 1)$. Let L and U be independent of the i.i.d. sequence $X_1, X_2, \dots$. Construct the RP based on X_n and starting with the initial age $A_0 = UL$ and the initial residual waiting time $R_0 = (1 - U)L$. Then this RP and the associated Markov process (A_t, R_t) are stationary.

Normal and Other Approximations

On the basis of equation (2) between S_n and N_t , one can deduct a normal approximation for N_t

(after standardization) from the standard central limit theorem for S_n . If $\sigma^2 = \text{Var}(X_1) < \infty$, one can conclude that

$$\lim_{t \rightarrow \infty} t^{-1} \text{Var}(N_t) = \frac{\sigma^2}{\mu^3} \quad (25)$$

and, setting $N_t^* = (N_t - \mu^{-1}t)/\mu^{-3/2}\sigma t^{1/2}$,

$$\lim_{t \rightarrow \infty} \mathbb{P}(N_t^* \leq x) = \int_{-\infty}^x (2\pi)^{-1/2} e^{-x^2/2} dx, \quad x \in \mathbb{R} \quad (26)$$

Equation (26) can be extended to an invariance principle. For given $t > 0$ define the process $Y = (Y^{(t)}(s))_{s \in [0,1]}$ by $Y^{(t)}(s) = N_{st}^*$. Y takes values in the function space $D[0,1]$ and it can be shown that it converges weakly to **Brownian motion** as $t \rightarrow \infty$. For weak and strong functional limit theorems for RPs, we refer to [23–26]. A related invariance principle for the process $X^{(n)} = (S_{[nt]})_{t \in [0,1]}$, conditioned on $S_n = s$ for fixed $s > 0$, has been proved in [13]. Here we consider the pure case $S_0 = 0$. The conditional distribution of $X^{(n)}$, given $S_n = s$, on $D[0,1]$ converges to the point mass at the linear function x_s , defined by $x_s(t) = st$, if $\lim_{n \rightarrow \infty} n \mathbb{E}[X_1^2] = 0$. This latter condition is for example satisfied if F is strongly unimodal (i.e., it has a log concave density). Thus, under the assumption that the n th renewal takes place at some fixed time s , the n renewal times in $[0, s]$ tend to be uniformly dispersed in this interval as $n \rightarrow \infty$. Technically, let us measure this dispersion by

$$U_n = \sup_{0 \leq t_1 \leq t_2 \leq s} |n^{-1}[N_{t_2} - N_{t_1}] - s^{-1}(t_2 - t_1)| \quad (27)$$

Then $\mathbb{P}(U_n > \varepsilon \mid S_n = s) = 0$ for all $\varepsilon > 0$.

For a certain range of values of t and n , the probability that $N_t = n + 1$ can be approximated in terms of the Laplace–Stieltjes transform $\hat{F}(\alpha) = \int_0^\infty e^{-\alpha t} dF(t)$ of F , using the analytic method of steepest descent. Assume that for $\ell(\alpha) = \ln \hat{F}(\alpha)$ the equation $\ell'(\alpha) = -t/n$ has exactly one real root ζ . Then for the pure RP with renewal epoch cdf F the probability $P(N_t = n + 1)$ can be approximated by

$$\frac{e^{\zeta t} (1 - \hat{F}(\zeta)) (\hat{F}(\zeta))^n}{(2\pi n \ell''(\zeta))^{1/2} \zeta}$$

for those values of t and n for which $n \ell''(\zeta)$ is large. This technique yields good results in various special cases.

References

- [1] Feller, W. (1971). *An Introduction to Probability Theory and its Applications*, 2nd Edition, John Wiley & Sons, New York, Vol. 2.
- [2] Alsmeyer, G. (1991). *Erneuerungstheorie*, Teubner, Stuttgart.
- [3] Asmussen, S. (2001). *Applied Probability and Queues*, Springer, New York.
- [4] Ross, S. (1996). *Stochastic Processes*, John Wiley & Sons.
- [5] Jensen, U. (1984). Some remarks on the renewal function of the uniform distribution, *Advances in Applied Probability* **16**, 214–215.
- [6] Gnedenko, B.V., Belyayev, Yu.K. & Solov'yev, A.D. (1969). *Mathematical Methods of Reliability Theory*, Academic Press, New York.
- [7] Stadje, W. (1987). A note on the renewal process for truncated exponential variables, *Statistics and Probability Letters* **6**, 61–66.
- [8] Smith, W.L. & Leadbetter, M.R. (1963). On the renewal function for the Weibull distribution, *Technometrics* **5**, 393–396.
- [9] Breiman, L. (1968). *Probability*, Addison-Wesley.
- [10] Woodroffe, M. (1982). *Nonlinear Renewal Theory in Sequential Analysis*, SIAM.
- [11] Lindvall, T. (1992). *Lectures on the Coupling Method*, John Wiley & Sons.
- [12] Thorisson, H. (2000). *Coupling, Stationarity and Regeneration*, Springer.
- [13] Schulte-Geers, E. & Stadje, W. (1988). Some results on the joint distribution of the renewal epochs prior to a given time instant, *Stochastic Processes and their Applications* **30**, 85–104.
- [14] Feller, W. (1941). On the integral equation of renewal theory, *The Annals of Mathematical Statistics* **12**, 243–267.
- [15] Täcklind, S. (1945). Fourieranalytische Behandlung vom Erneuerungsproblem, *Skandinavisk Aktuarietidskrift* **28**, 68–105.
- [16] Smith, W.L. (1954a). Asymptotic renewal theorems, *Proceedings of the Royal Society of Edinburgh* **A64**, 9–48.
- [17] Smith, W.L. (1954b). Extensions of a renewal theorem, *Proceedings of the Cambridge Philosophical Society* **51**, 629–638.
- [18] Smith, W.L. (1962). On necessary and sufficient conditions for the convergence of the renewal density, *Transactions of the American Mathematical Society* **104**, 79–100.
- [19] Stadje, W. (1989). The asymptotic behavior of the renewal density, *Methods of Operations Research* **58**, 307–317.
- [20] Stadje, W. (2006). Integral representations and convergence of the renewal density, *Journal of Mathematical Analysis and Applications* **323**, 974–984.

-
- [21] Grübel, R. (1982). Über das asymptotische Verhalten von Erneuerungsdichten, *Mathematische Nachrichten* **108**, 39–48.
- [22] Grübel, R. (1989). Stochastic models as functionals: some remarks on the renewal case, *Journal of Applied Probability* **26**, 296–303.
- [23] Lamperti, J. (1962). An invariance principle in renewal theory, *The Annals of Mathematical Statistics* **33**, 685–696.
- [24] Alex, M. & Steinebach, J. (1995). Invariance principles for renewal processes and some applications, *Theory of Probability and Mathematical Statistics* **50**, 23–54.
- [25] Steinebach, J. (1988). Invariance principles for renewal processes when only moments of low order exist, *Journal of Multivariate Analysis* **26**, 169–183.
- [26] Csörgö, M. & Horváth, L. (1993). *Weighted Approximations in Probability and Statistics*, John Wiley & Sons, New York.

Related Articles

Markov Renewal Processes in Reliability Modeling.

WOLFGANG STADJE

Software Reliability Modeling and Analysis

Introduction

It is a well-known fact that modern society depends more and more heavily on software systems, and this trend will certainly continue in the future. In order to improve the reliability of a software product being developed, it is a common practice to have the software tested extensively prior to its release to customers. The testing of the software is a verification process for software reliability improvement, which is featured by the phenomenon of failure observation and the activity of fault removal. Software reliability models are commonly used to monitor the testing process and to measure and predict present and future software reliability.

Since the early 1970s, over 100 software reliability models have been proposed and studied. In this chapter we give a brief overview of software reliability models and their applications. More detailed reviews of software reliability modeling can be found in Xie [1], Lyu [2], Pham [3], and Xie *et al.* [4]. The readers are also encouraged to read some recent research articles for more in-depth ideas on various issues related to software reliability modeling and applications.

Overview of Some Traditional Software Reliability Models

Software reliability models can be classified into different categories according to different classification systems. Xie [1] provided an extensive coverage of most software reliability models that were widely discussed till that time. Recently, research has focused more on the application and minor extension of those models. In this section we adopt the classification system based on model assumptions.

Markov Models

The Jelinski–Moranda model [5] is the first published Markov model, which has profound influences on software reliability modeling thereafter. The main assumptions of this model are

1. The number of initial faults is an unknown but fixed constant.
2. A detected fault is removed immediately and no new faults are introduced.
3. Times between failures are independent, exponentially distributed random quantities.
4. Each remaining software fault contributes the same amount to the software failure intensity.

Denote by N_0 the number of initial faults in the software before the testing starts, then the initial failure intensity is $N_0\phi$, where ϕ is a constant of proportionality denoting the failure intensity contributed by each fault. Denote by $T_i, i = 1, 2, \dots, N_0$, the time between $(i - 1)$ th and i th failures, then T_i s are independent, exponentially distributed random variables with parameter

$$\lambda_i = \phi[N_0 - (i - 1)], i = 1, 2, \dots, N_0 \quad (1)$$

Many modified versions of the Jelinski–Moranda model have been studied in the literature. Schick and Wolverton [6] proposed a model assuming that times between failures are not exponential but follow Rayleigh distribution. Shanthikumar [7] generalized the Jelinski–Moranda model by using general time-dependent transition probability function. Xie [8] developed a general decreasing failure intensity model allowing for the possibility of different fault sizes. Whittaker *et al.* [9] considered a model which allows use of prior testing data to cover the real-world scenario in which the release build is constructed only after a succession of repairs to buggy prerelease builds. Boland and Singh [10] recently considered a birth-process approach to a related software reliability model.

Most software reliability models are based on the assumption of statistical independence among successive software failures. Goseva–Popstojanova and Trivedi [11] developed the software reliability modeling framework which can consider the phenomena of **failure correlation** and studied its effects on the software reliability measures. A recent paper on a related topic is Dai *et al.* [12].

NHPP Models

For this type of models, the software failure process is assumed to follow a **nonhomogeneous Poisson process** (NHPP). Denote by $N(t)$ the number of

2 Software Reliability Modeling and Analysis

observed failures until time t , the main assumptions of this type of models are

1. $N(0) = 0$.
2. $\{N(t), t \geq 0\}$ has independent increments.
3. At time t , $N(t)$ follows a Poisson distribution with parameter $m(t)$, i.e.,

$$P\{N(t) = n\} = \frac{[m(t)]^n}{n!} \exp[-m(t)], \quad n = 0, 1, 2, \dots \quad (2)$$

where $m(t)$ is called the *mean value function* of the NHPP, which describes the expected cumulative number of experienced **failures** in time interval $(0, t]$.

The *failure intensity function*, $\lambda(t)$, is defined as

$$\begin{aligned} \lambda(t) &\equiv \lim_{\Delta t \rightarrow 0^+} \frac{P(N(t + \Delta t) - N(t) > 0)}{\Delta t} \\ &= \frac{dm(t)}{dt}, \quad t \geq 0 \end{aligned} \quad (3)$$

Generally, by using different mean value function of $m(t)$, we get different NHPP models (Figure 1).

The Goel–Okumoto Model. The most representative NHPP model was proposed by Goel and Okumoto [13]. In fact, many other NHPP models are modifications, or generalizations, of this model. The mean value function of the Goel–Okumoto model is

$$m(t) = a[1 - \exp(-bt)], \quad a > 0, b > 0 \quad (4)$$

where a is the expected number of faults to be eventually detected and b is the failure occurrence

rate per fault. The failure intensity function of the Goel–Okumoto model is

$$\lambda(t) = ab \exp(-bt) \quad (5)$$

The Musa–Okumoto Logarithmic Model. Musa and Okumoto [14] developed a logarithmic Poisson execution time model. The mean value function is

$$m(t) = \frac{1}{\varphi} \ln(\lambda_0 \varphi t + 1), \quad \varphi > 0, \lambda_0 > 0 \quad (6)$$

where λ_0 is the initial failure intensity and φ is the failure intensity decay parameter. Since this model allows infinite number of failures observed, it is also called an *infinite failure model*.

The Yamada S-Shaped NHPP Models. Based on experiences, it is sometimes observed that the curve of the cumulative number of faults is S-shaped. Several different S-shaped NHPP models have been proposed in the existing literature. Among them, the most interesting ones are the delayed S-shaped NHPP model [15] and the inflected S-shaped NHPP model [16]. The mean value function of the delayed S-shaped NHPP model is

$$m(t) = a[1 - (1 + bt) \exp(-bt)], \quad a > 0, b > 0 \quad (7)$$

and the mean value function of the infected S-shaped NHPP model is

$$m(t) = \frac{a[1 - \exp(-bt)]}{[1 + c \exp(-bt)]}, \quad a > 0, b > 0, c > 0 \quad (8)$$

The Log-Power Model. An interesting model called *log-power model* was proposed in Xie and Zhao [17]. It has the mean value function

$$m(t) = a \ln^b(1 + t), \quad a > 0, b > 0 \quad (9)$$

This model is a modification of the traditional Duane model [18] for general repairable system. A useful property is that if we take the logarithmic on both sides of the mean value function, we have

$$\ln m(t) = \ln a + b \ln \ln(1 + t) \quad (10)$$

Hence a graphical procedure is established. If we plot the observed cumulative number of failures versus $(t + 1)$, the plot should tend to be on a straight line

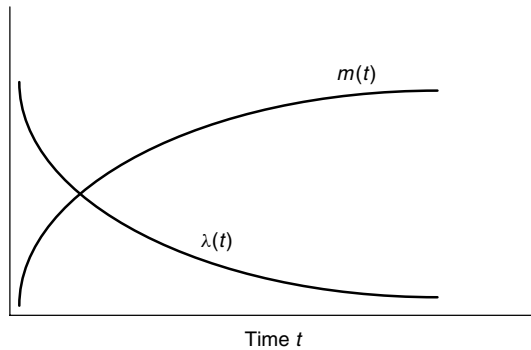


Figure 1 An example of failure intensity function $\lambda(t)$ and mean value function $m(t)$ – the Goel–Okumoto Model

on a log–loglog scale. This can be used to easily estimate the model parameters, and more importantly, to validate the model.

Other NHPP Extensions. NHPP is a very favorable alternative in software reliability modeling, which has continuously been receiving a lot of attention. Yamada *et al.* [19] incorporated the concept of test effort into the Goel–Okumoto model for a better description of the failure phenomenon. Later, Xia *et al.* [20] incorporated the concept of a learning factor into the same model. Chatterjee *et al.* [21] presented a **software reliability growth model** which incorporates the joint effect of test effort and learning factor. Gokhale and Trivedi [22] proposed an enhanced NHPP model which incorporates explicitly the time-varying test-coverage function in its analytical formulation. Teng and Pham [23] considered a model for predicting software reliability by including a random environmental factor.

Bayesian Models

For this type of models, Bayesian assumption is made or a Bayesian inference is adopted. The well-known model by Littlewood and Verrall [24] is the first **Bayesian software reliability** model assuming that times between failures are exponentially distributed with a parameter that is treated as a random variable, which is assumed to have a Gamma prior distribution. There are also many papers dealing with Bayesian treatment of the Jelinski–Moranda model, see e.g., Langberg and Singpurwalla [25], Littlewood and Sofer [26] and Csenki [27]. The most important feature of using a Bayesian model is that prior information can be incorporated into **parameter estimation** procedure.

There is also abundant research on the Bayesian inference for NHPP software reliability models. Kuo *et al.* [28] investigated the Bayesian inference for an NHPP with an *S*-shaped mean value function. Cid and Achcar [29] presented Bayesian analysis for NHPP in software reliability models assuming non-monotonic intensity functions. Pham and Pham [30] proposed a Bayesian model, for which times between failures follow Weibull distributions with stochastically decreasing ordering on the hazard functions of successive failure time intervals. Rees *et al.* [31] used Bayesian graphical models to model the uncertainties involved in testing processes.

Software Metrics–Based Quality Models

Most software reliability models make use of the failure time data during the testing. However, because of the limited amount of data available, the prediction is not accurate. Software metrics–based models take a general modeling approach that makes inference based on various measurement from software product and development processes, called *software metrics*. Normally, models of this kind are in a broader context than sole software reliability estimation as software metrics are useful for the study of the whole development process and quality attributes of the resulting software product.

Usually, software metrics are treated as independent variables which are regarded as software quality indicators, and by using different techniques, the values of dependent variables, i.e., fault-proneness of a software module (a binary Boolean value) or the number of faults in a software module can be predicted. These techniques include logistic regression [32], classification and regression trees [33], **artificial neural networks** [34], and so on.

Point and Interval Estimation for Model Parameters

In each software reliability model, there always exist some parameters that have to be estimated based on historical data. Maximum-likelihood estimation (MLE), due to its many good mathematical properties, is a commonly adopted method for parameter estimation [35]. This will be discussed here.

Point Estimation for NHPP Model Parameters

In practice, the **failure data** may arise in two different types, i.e., numbers of failures or failure times. For the first case, denote by n_i the number of failures observed in time interval $(s_{i-1}, s_i]$, where $0 \equiv s_0 < s_1 < \dots < s_k$ and $s_i (i > 0)$ is the prescribed time point in the software testing process, then the likelihood function for an NHPP model with mean value function $m(t)$ is

$$L(n_1, n_2, \dots, n_k) = \prod_{i=1}^k \frac{[m(s_i) - m(s_{i-1})]^{n_i} \exp \{-[m(s_i) - m(s_{i-1})]\}}{n_i!} \quad (11)$$

4 Software Reliability Modeling and Analysis

The parameters in $m(t)$ can be estimated by maximizing the likelihood function given above. Usually, numerical procedures have to be used.

For the second case, denote by $T_i, i = 1, 2, \dots, k$, the observed k failure times in a software testing process, where k is the prescribed number of failures to be observed, then the likelihood function is

$$L(T_1, T_2, \dots, T_k) = \exp[-m(T_k)] \prod_{i=1}^k \lambda(T_i) \quad (12)$$

Similarly, the NHPP model parameters can be estimated by maximizing the likelihood function given above.

A General Approach for Interval Estimation

In order to find asymptotic **confidence intervals** for the k model parameters, the derivation of the *Fisher information matrix* is needed, which is given by

$$I(\theta_1, \dots, \theta_k) \equiv \begin{bmatrix} -E \left[\frac{\partial^2 \ln L}{\partial \theta_1^2} \right] & -E \left[\frac{\partial^2 \ln L}{\partial \theta_2 \partial \theta_1} \right] & \dots & -E \left[\frac{\partial^2 \ln L}{\partial \theta_k \partial \theta_1} \right] \\ -E \left[\frac{\partial^2 \ln L}{\partial \theta_1 \partial \theta_2} \right] & -E \left[\frac{\partial^2 \ln L}{\partial \theta_2^2} \right] & \dots & -E \left[\frac{\partial^2 \ln L}{\partial \theta_k \partial \theta_2} \right] \\ \dots & \dots & \dots & \dots \\ -E \left[\frac{\partial^2 \ln L}{\partial \theta_1 \partial \theta_k} \right] & -E \left[\frac{\partial^2 \ln L}{\partial \theta_2 \partial \theta_k} \right] & \dots & -E \left[\frac{\partial^2 \ln L}{\partial \theta_k^2} \right] \end{bmatrix} \quad (13)$$

In (13), L is given by (11) or (12). From the asymptotic theory of MLE, when n approaches to infinity, $[\hat{\theta}_1, \dots, \hat{\theta}_k]$ converges in distribution to **k -variate normal distribution** with mean $[\theta_1, \dots, \theta_k]$ and **covariance matrix** I^{-1} . That is, the asymptotic covariance matrix of the MLEs is

$$V_e \equiv \begin{bmatrix} \text{Var}(\hat{\theta}_1) & \text{Cov}(\hat{\theta}_1, \hat{\theta}_2) & \dots & \text{Cov}(\hat{\theta}_1, \hat{\theta}_k) \\ \text{Cov}(\hat{\theta}_2, \hat{\theta}_1) & \text{Var}(\hat{\theta}_2) & \dots & \text{Cov}(\hat{\theta}_2, \hat{\theta}_k) \\ \dots & \dots & \dots & \dots \\ \text{Cov}(\hat{\theta}_k, \hat{\theta}_1) & \text{Cov}(\hat{\theta}_k, \hat{\theta}_2) & \dots & \text{Var}(\hat{\theta}_k) \end{bmatrix} \\ = I^{-1} \quad (14)$$

Therefore, the asymptotic $100(1 - \alpha)\%$ confidence interval for $\hat{\theta}_i$ is

$$\left(\hat{\theta}_i - z_{\alpha/2} \sqrt{\text{Var}(\hat{\theta}_i)}, \hat{\theta}_i + z_{\alpha/2} \sqrt{\text{Var}(\hat{\theta}_i)} \right), i = 1, \dots, k \quad (15)$$

where $z_{\alpha/2}$ is the $1 - \alpha/2$ percentile of the standard normal distribution and the quantity $\text{Var}(\hat{\theta}_i)$ can

be obtained from the covariance matrix given by equation (14).

Since the true values of θ_i 's are unknown, the *observed information matrix* is often used

$$\hat{I}(\hat{\theta}_1, \dots, \hat{\theta}_k) \\ \equiv \begin{bmatrix} -\frac{\partial^2 \ln L}{\partial \theta_1^2} & -\frac{\partial^2 \ln L}{\partial \theta_2 \partial \theta_1} & \dots & -\frac{\partial^2 \ln L}{\partial \theta_k \partial \theta_1} \\ -\frac{\partial^2 \ln L}{\partial \theta_1 \partial \theta_2} & -\frac{\partial^2 \ln L}{\partial \theta_2^2} & \dots & -\frac{\partial^2 \ln L}{\partial \theta_k \partial \theta_2} \\ \dots & \dots & \dots & \dots \\ -\frac{\partial^2 \ln L}{\partial \theta_1 \partial \theta_k} & -\frac{\partial^2 \ln L}{\partial \theta_2 \partial \theta_k} & \dots & -\frac{\partial^2 \ln L}{\partial \theta_k^2} \end{bmatrix} \begin{matrix} \theta_1 = \hat{\theta}_1 \\ \dots \dots \dots \\ \theta_k = \hat{\theta}_k \end{matrix} \quad (16)$$

and the confidence interval for $\hat{\theta}_i$ can be calculated.

If $\Phi \equiv g(\theta_1, \dots, \theta_k)$, where $g(\cdot)$ is a continuous function, then as n converges to infinity, $\hat{\Phi}$ converges

in distribution to normal distribution with mean Φ and variance

$$\text{Var}(\hat{\Phi}) = \sum_{i=1}^k \sum_{j=1}^k \frac{\partial g}{\partial \theta_i} \cdot \frac{\partial g}{\partial \theta_j} \cdot v_{ij} \quad (17)$$

where v_{ij} is the element of the i th row and j th column of the matrix in (16). The asymptotic $100(1 - \alpha)\%$ confidence interval for $\hat{\Phi}$ is

$$\left(\hat{\Phi} - z_{\alpha/2} \sqrt{\text{Var}(\hat{\Phi})}, \hat{\Phi} + z_{\alpha/2} \sqrt{\text{Var}(\hat{\Phi})} \right) \quad (18)$$

The above result can be used to obtain the confidence interval for some useful quantities such as failure intensity or reliability.

An Example

Assume that there are two model parameters, a and b , which is quite common in most of software reliability

models, the Fisher information matrix is given by

$$I \equiv \begin{bmatrix} -E \left[\frac{\partial^2 \ln L}{\partial a^2} \right] & -E \left[\frac{\partial^2 \ln L}{\partial a \partial b} \right] \\ -E \left[\frac{\partial^2 \ln L}{\partial a \partial b} \right] & -E \left[\frac{\partial^2 \ln L}{\partial b^2} \right] \end{bmatrix} \quad (19)$$

The asymptotic covariance matrix of the maximum-likelihood (ML) estimators is the inverse of (19):

$$V_e = I^{-1} = \begin{bmatrix} \text{Var}(\hat{a}) & \text{Cov}(\hat{a}, \hat{b}) \\ \text{Cov}(\hat{a}, \hat{b}) & \text{Var}(\hat{b}) \end{bmatrix} \quad (20)$$

For illustrative purposes, we take the data cited in Table 1 in Zhang and Pham [36]. The observed information matrix is

$$\begin{aligned} \hat{I} &\equiv \begin{bmatrix} -\frac{\partial^2 \ln L}{\partial a^2} & -\frac{\partial^2 \ln L}{\partial a \partial b} \\ -\frac{\partial^2 \ln L}{\partial a \partial b} & -\frac{\partial^2 \ln L}{\partial b^2} \end{bmatrix} \begin{matrix} a = \hat{a} \\ b = \hat{b} \end{matrix} \\ &= \begin{bmatrix} 0.0067 & 1.1095 \\ 1.1095 & 4801.2 \end{bmatrix} \end{aligned} \quad (21)$$

We can then obtain the asymptotic variance and covariance of the MLE as follows

$$\begin{aligned} \text{Var}(\hat{a}) &= 154.85, \text{Var}(\hat{b}) = 2.17 \times 10^{-4} \\ \text{Cov}(\hat{a}, \hat{b}) &= -0.0358 \end{aligned} \quad (22)$$

The 95% confidence intervals on parameters a and b are (117.93, 166.71) and (0.0957, 0.1535), respectively.

Optimal Testing-Resource Allocation for Modular Software Systems

In general, a software development process consists of the following four major phases: specification, design, coding and testing. The testing phase is the most costly and time-consuming one. It is thus very much important for the management to spend the limited testing resource efficiently.

Optimal Testing-Resource Allocation Problem

In this section, we directly consider the case of system optimization problem [4]. For the optimal

testing-resource allocation problem, the following assumptions are made:

1. A software system is composed of n independent modules which are developed and tested independently during the unit testing phase.
2. In the unit testing phase, each software module is subject to failures at random times caused by faults remaining in the software module.
3. The failure observation process of software module i is modeled by an NHPP with mean value function $m_i(t)$, or failure intensity function $\lambda_i(t) \equiv dm_i(t)/dt$.

A total amount of testing time T is available for the whole software system that consists of a few modules. It will be allocated to each software module in such a way that the reliability of the system after the unit testing phase will be maximized.

In the following, we study the optimal testing-resource allocation problem in a generic way. Our formulations and solutions are not restricted to any specific software reliability model but are generally applicable to NHPP models.

Optimal Testing-Resource Allocation for Serial Software Systems

If the software system fails whenever there is a failure with any of the modules that the software system is composed of, then it is called a *serial software system*. The structure for such a system is illustrated in Figure 2.

Problem Formulation and Solution. Denote by T_i the testing time allocated to module i . After T_i unit of time of testing, the failure intensity of module i is $\lambda_i(T_i)$. The reliability of module i is

$$R_i(x|T_i) = \exp\{-\lambda_i(T_i)x\}, x \geq 0, i = 1, 2, \dots, n \quad (23)$$

It should be noted that this formulation is based on operational reliability concept and is different from that based on testing reliability concept, which has

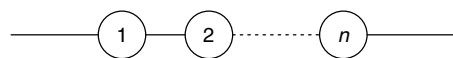


Figure 2 Structure of a serial software system

6 Software Reliability Modeling and Analysis

the formulation as follows

$$R_i(x|T_i) = \exp\{-[m_i(T_i + x) - m_i(T_i)]\},$$

$$x \geq 0, i = 1, 2, \dots, n \quad (24)$$

From a customers' point of view, operational reliability is more important and meaningful than testing reliability [37]. Thus the formulation of equation (23) is more appropriate for the optimal testing-resource allocation problem than equation (24).

The reliability of the software system is thus given by

$$R(x|T_1, \dots, T_n) = \prod_{i=1}^n R_i(x|T_i)$$

$$= \exp\left\{-x \sum_{i=1}^n \lambda_i(T_i)\right\} \quad (25)$$

The optimal testing-resource allocation problem is formulated as

$$\text{Maximize } R(x|T_1, \dots, T_n) = \exp\left\{-x \sum_{i=1}^n \lambda_i(T_i)\right\}$$

$$\text{Subject to } \sum_{i=1}^n T_i \leq T$$

$$\text{and } T_i \geq 0, \quad i = 1, 2, \dots, n \quad (26)$$

The formulation above is equivalent to minimizing the sum of failure intensities, i.e.,

$$\text{Minimize } \sum_{i=1}^n \lambda_i(T_i) \quad (27)$$

and this can also be justified by the assumption that any module failure causes system failure. However, it should be noted that this is not the same as the minimization of the total number of experienced failures, which has been commonly used:

$$\text{Minimize } \sum_{i=1}^n [m_i(T_i + x) - m_i(T_i)] \quad (28)$$

In order to derive a general solution to this problem, traditional mathematical methods can be used. The Lagrangian is constructed as

$$L = \sum_{i=1}^n \lambda_i(T_i) + \lambda \left(\sum_{i=1}^n T_i - T \right) \quad (29)$$

The necessary and sufficient conditions for the minimum are

$$\frac{\partial L}{\partial T_i} = \frac{\partial \lambda_i(T_i)}{\partial T_i} + \lambda \geq 0, i = 1, 2, \dots, n \quad (30)$$

under the condition that $T_i \cdot \frac{\partial L}{\partial T_i} = 0, i = 1, 2, \dots, n$; $\lambda (\sum_{i=1}^n T_i - T) = 0$, and $\lambda \geq 0$

The optimal solution $T_1^*, T_2^*, \dots, T_n^*$ can be obtained by solving the above equations numerically.

Example 1 Assume that the software system is composed of 10 modules, and the testing processes follow the Goel–Okumoto model [13]. An optimization algorithm can be used to obtain the solution [38]. In this case, the Lagrangian becomes

$$L = \sum_{i=1}^{10} a_i b_i \exp(-b_i t) + \lambda \left(\sum_{i=1}^n T_i - T \right) \quad (31)$$

and it can be shown that

If $a_k b_k^2 > \lambda \geq a_{k+1} b_{k+1}^2$, then

$$T_i^* = \begin{cases} \frac{1}{b_i} [\ln(a_i b_i^2) - \ln \lambda] & (i = 1, 2, \dots, k) \\ 0 & (i = k + 1, \dots, 10) \end{cases} \quad (32)$$

where λ is

$$\lambda = \exp\left(\frac{\sum_{i=1}^k \frac{1}{b_i} \ln[a_i b_i^2] - T}{\sum_{i=1}^k \frac{1}{b_i}}\right) \quad (33)$$

Suppose that the parameters of the 10 modules have been estimated and summarized in Table 1, and an

Table 1 Estimated parameters and the optimal testing-resource allocation for the example

Module	a_i	b_i (10^{-4})	T_i^*
1	89	4.1823	9198.6
2	25	5.0923	5834.4
3	27	3.9611	6426.6
4	39	2.5336	7971.3
5	45	2.2956	8561.6
6	39	1.7246	7249.8
7	59	0.8819	3661.9
8	68	0.7274	1095.8
9	14	1.5309	0
10	37	0.6824	0

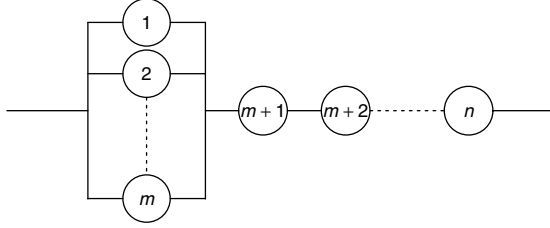


Figure 3 The structure of a mix-structured redundant software system

additional 50 000 h of testing time is to be allocated among the modules. By solving the optimization problem, the optimal allocation is obtained and shown as the last column from the right in Table 1.

Under the optimal allocation, after the unit testing phase the reliability of the software system will be:

$$R(x) = \exp(-0.01997x), \quad x \geq 0 \quad (34)$$

Optimal Testing-Resource Allocation for Redundant Software Systems

It is more appropriate to consider the optimal testing-resource allocation problem for series-parallel mix-structured systems. Figure 3 is an illustration of such a software system, which has m redundant modules and $(n-m)$ serial modules.

For a mix-structured redundant software system as shown in Figure 3, the achieved reliability of the system after unit testing phase is as follows:

$$R(x|T_1, \dots, T_m, T_{m+1}, \dots, T_n) = \left(1 - \prod_{i=1}^m [1 - R_i(x|T_i)]\right) \prod_{j=m+1}^n R_j(x|T_j), \quad x \geq 0 \quad (35)$$

where $R_k(x|T_k)$, $k = 1, 2, \dots, n$, is the operational reliability of module k which can be calculated from equation (23). Thus the optimal testing-resource allocation problem is formulated as

$$\begin{aligned} &\text{Maximize } R(x|T_1, \dots, T_m, T_{m+1}, \dots, T_n) \\ &\text{Subject to } \sum_{i=1}^n T_i \leq T, \\ &\text{and } T_i \geq 0, i = 1, 2, \dots, n \end{aligned} \quad (36)$$

Discussions

Software reliability analysis is an important yet tough problem for both computer scientists and statisticians. There are many challenging issues. Many models are based on the modeling of failure data that may not be sufficient for accurate estimation. Other data and information from product and development process could be useful in some cases. However, they are usually not collected systematically or may not be available at all, especially in the early phases in a software development process. Sometimes the increased accuracy with additional information is difficult to be established as well. On the other hand, software reliability prediction is an important issue which has to be done for risk assessment and cost analysis. Therefore, further research and developments are expected.

References

- [1] Xie, M. (1991). *Software Reliability Modelling*, World Scientific Publisher, Singapore.
- [2] Lyu, M. (ed) (1996). *Handbook of Software Reliability Engineering*, McGraw-Hill, New York.
- [3] Pham, H. (2000). *Software Reliability*, Springer-Verlag, Singapore.
- [4] Xie, M., Dai, Y.S. & Poh, K.L. (2004). *Computing Systems Reliability*, Kluwer Academic, Boston.
- [5] Jelinski, Z. & Moranda, P.B. (1972). Software reliability research, in *Statistical Computer Performance Evaluation*, W. Freiberger, ed, Academic Press, New York, pp. 465–497.
- [6] Schick, G.J. & Wolverton, R.W. (1978). An analysis competing software reliability models, *IEEE Transactions on Software Engineering* **SE-4**, 104–120.
- [7] Shanthikumar, J.G. (1981). A general software reliability model for performance prediction, *Microelectronics and Reliability* **21**, 671–682.
- [8] Xie, M. (1987). A shock model for software failures. *Microelectronics and Reliability*, **27**(4), 717–724.
- [9] Whittaker, J., Rekab, K. & Thomason, M.G. (2000). A Markov chain model for predicting the reliability of multi-build software, *Information and Software Technology* **42**, 889–894.
- [10] Boland, P.J. & Singh, H. (2003). A birth-process approach to Moranda's geometric software-reliability model, *IEEE Transactions on Reliability* **52**(2), 168–174.
- [11] Goseva-Popstojanova, K. & Trivedi, K.S. (2000). Failure correlation in software reliability models, *IEEE Transactions on Reliability* **49**, 37–48.
- [12] Dai, Y.S., Xie, M. & Poh, K.L. (2005). Modeling and analysis of correlated software failures of multiple types, *IEEE Transactions on Reliability* **R-54**(1), 100–106.

- [13] Goel, A.L. & Okumoto, K. (1979). Time dependent error-detection rate model for software reliability and other performance measures, *IEEE Transactions on Reliability* **R-28**, 206–211.
- [14] Musa, J.D. and Okumoto, K. (1984). A logarithmic Poisson execution time model for software reliability measurement. *Proceedings of the 7th International Conference on Software Engineering*, Orlando, Florida, United States, pp. 230–238.
- [15] Yamada, S., Ohba, M. & Osaki, S. (1984). S-shaped software reliability growth models and their applications, *IEEE Transactions on Reliability* **R-33**, 289–292.
- [16] Ohba, M. (1984). Software reliability analysis models, *IBM Journal of Research and Development* **28**, 428–443.
- [17] Xie, M. & Zhao, M. (1993). On some reliability growth models with graphical interpretations, *Microelectronics and Reliability* **33**, 149–167.
- [18] Duane, J.T. (1964). Learning curve approach to reliability monitoring, *IEEE Transactions on Aerospace* **AS-2**, 563–566.
- [19] Yamada, S., Hishitani, J. & Osaki, S. (1993). Software reliability growth with a Weibull test-effort: a model application, *IEEE Transactions on Reliability* **42**, 100–105.
- [20] Xia, G., Zeepongsekul, P. & Kumar, S. (1992). Optimal software release policies for models incorporating learning in testing, *Asia Pacific Journal of Operational Research* **9**, 221–234.
- [21] Chatterjee, S., Misra, R.B. & Alam, S.S. (1997). Joint effect of test effort and learning factor on software reliability and optimal release policy, *International Journal of Systems Science* **28**(4), 391–396.
- [22] Gokhale, S.S. & Trivedi, K.S. (1999). A time/structure based software reliability model, *Annals of Software Engineering* **8**, 85–121.
- [23] Teng, X.L. & Pham, H. (2006). A new methodology for predicting software reliability in the random field environments, *IEEE Transactions on Reliability* **55**(3), 458–468.
- [24] Littlewood, B. & Verrall, J.L. (1973). A Bayesian reliability growth model for computer software, *Applied Statistics* **22**, 332–346.
- [25] Langberg, N.A. & Singpurwalla, N.D. (1985). A unification of some software reliability models, *SIAM Journal on Scientific and Statistical Computing* **6**, 781–790.
- [26] Littlewood, B. & Sofer, A. (1987). A Bayesian modification to the Jelinski-Moranda software reliability growth model, *Software Engineering Journal* **2**, 30–41.
- [27] Csenki, A. (1990). Bayes predictive analysis of a fundamental software-reliability model, *IEEE Transactions on Reliability* **R-39**, 177–183.
- [28] Kuo, L., Lee, J.C., Choi, K. & Yang, T.Y. (1997). Bayes inference for S-shaped software-reliability growth models, *IEEE Transactions on Reliability* **46**, 76–81.
- [29] Cid, J.E.R. & Achcar, J.A. (1999). Bayesian inference for nonhomogeneous Poisson processes in software reliability models assuming nonmonotonic intensity functions, *Computational Statistics and Data Analysis* **32**, 147–159.
- [30] Pham, L. & Pham, H. (2001). A Bayesian predictive software reliability model with pseudo-failures, *IEEE Transactions on Systems Man and Cybernetics Part A – Systems and Humans* **31**(3), 233–238.
- [31] Rees, K., Coolen, F., Goldstein, M. & Wooff, D. (2001). Managing the uncertainties of software testing: a Bayesian approach, *Quality and Reliability Engineering International* **17**, 191–203.
- [32] Basili, V.R., Briand, L.C. & Melo, W. (1996). A validation of object-oriented design metrics as quality indicators, *IEEE Transactions on Software Engineering* **22**, 751–761.
- [33] Khoshgoftaar, T.M., Allen, E.B., Jones, W.D. & Hudspohl, J.P. (2000). Classification-tree models of software-quality over multiple releases, *IEEE Transactions on Reliability* **49**(1), 4–11.
- [34] Thwin, M.M.T. & Quah, T.S. (2005). Application of neural networks for software quality prediction using object-oriented metrics, *Journal of Systems and Software* **76**, 147–156.
- [35] Jeske, D.R. & Pham, H. (2001). On the maximum likelihood estimates for the Goel-Okumoto software reliability model, *American Statistician* **55**(3), 219–222.
- [36] Zhang, X.M. & Pham, H. (1998). A software cost model with error removal times and risk costs, *International Journal of Systems Science* **29**, 435–442.
- [37] Yang, B. & Xie, M. (2000). A study of operational and testing reliability in software reliability analysis, *Reliability Engineering and System Safety* **70**, 323–329.
- [38] Xie, M. & Yang, B. (2001). Optimal testing-time allocation for modular systems, *International Journal of Quality and Reliability Management* **18**, 854–863.

Further Reading

- Huang, C.Y. & Lyu, M.R. (2005). Optimal release time for software systems considering cost, testing-effort, and test efficiency, *IEEE Transactions on Reliability* **54**(4), 583–591.
- Pham, L. & Pham, H. (2000). Software reliability models with time-dependent hazard function based on Bayesian approach, *IEEE Transactions on Systems Man and Cybernetics Part A – Systems and Humans* **30**, 25–35.

Related Articles

Software Reliability: Bayesian Analysis; Software Failure Data Analysis.

MIN XIE AND BO YANG

Availability Models

Introduction

We study a system subject to random deterioration and failure. At failures, repairs or replacements are performed, and the system resumes operation. To characterize the performance of the system, different measures can be used including the ones enumerated below:

1. The probability that the system is functioning at a certain point in time (point availability).
2. The probability that the system is functioning during a specified interval $[u, v]$ (interval availability).
3. The average point availability in a time interval (average availability).
4. The probability distribution of the proportion of time the system is functioning (downtime distribution).
5. The distribution and mean number of system failures during a time interval $[u, v]$.

Traditionally, the focus has been on analyzing the point availability and its limit (the steady-state availability). For a single component, the steady-state formula is given by $MTTF/(MTTF + MTTR)$, where $MTTF$ and $MTTR$ represent the mean time to failure and the mean time to repair (mean repair time), respectively. The steady-state probability of a system comprising several components can then be calculated using the theory of complex (monotone) systems.

Also, performance measures related to a time interval, such as the interval availability, are frequently used. Compared to the steady-state availability, it is of course more complicated to compute the performance measures related to a time interval. Using simplifications and approximations, it is however possible to establish formulae which can be used in practice.

In this article we will look closer at the above availability measures, that is, the first four types of measures in the list above. We establish methods and formulae for computing these measures. We also present some asymptotic results for large time intervals and for systems having highly available components. We first address the one-component case, then more complex systems.

Our focus is on situations where the system is being repaired or replaced immediately at failures. Safety systems where the possible failures are hidden until the system is activated or tested will not be covered. We refer to Rausand and Høyland [1] for a review of availability models applicable in such cases.

Our main reference is Aven and Jensen [2], which also covers analysis of other performance measures apart from the availability measures, including measures related to the number of system failures. Two key references to the classical availability theory are Barlow and Proschan [3, 4].

Basic Setup

Let $X_t = 1$, if the system is functioning at time t , and 0 otherwise. Furthermore let T_k , $k \in \mathbb{N}$ (here $\mathbb{N} = \{1, 2, \dots\}$), represent the length of the k th operation period, and let R_k , $k \in \mathbb{N}$, represent the length of the k th repair/replacement time for the system, see Figure 1. We assume that (T_k) , $k \in \mathbb{N}$, and (R_k) , $k \in \mathbb{N}$, are independent and identically distributed (i.i.d.) sequences of positive random variables. We denote the probability distributions of T_k and R_k by F and G , respectively, and assume that they have finite means μ_F and μ_G , respectively.

In reliability engineering, μ_F and μ_G are referred to as the *mean time to failure* and the *mean time to repair*, respectively.

To simplify the presentation, we also assume that F is an absolutely continuous distribution, that is, F has a density function f and failure rate function λ . We do not make the same assumption for the distribution function G , since that would exclude discrete repair time distributions which are often used in practice.

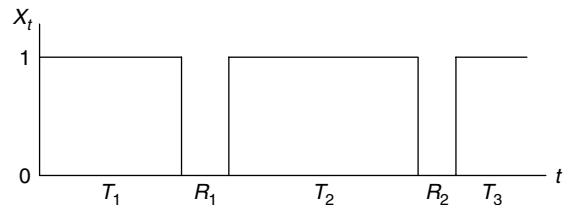


Figure 1 Time evolution of a failure and repair process for a one-component system starting at time $t = 0$ in the operating state

2 Availability Models

In some cases we also need the variances of F and G , denoted σ_F^2 and σ_G^2 , respectively. In the following, when writing the variance of a random variable, it is tacitly assumed that these are finite.

The sequence $T_1, R_1, T_2, R_2, \dots$ forms an alternating renewal process.

We introduce the following variables

$$S_n = T_1 + \sum_{k=1}^{n-1} (R_k + T_{k+1}), \quad n \in \mathbb{N} \quad (1)$$

and

$$S_n^\circ = \sum_{k=1}^n (T_k + R_k), \quad n \in \mathbb{N} \quad (2)$$

By convention, $S_0 = S_0^\circ = 0$ and sums over empty sets are zero. We see that S_n represents the n th failure time, and S_n° represents the completion time of the n th repair.

The S_n sequence generates a modified (delayed) renewal process N with renewal function M . The first interarrival time has distribution F . All other interarrival times have distribution $F \times G$ (convolution of F and G), with mean $\mu_F + \mu_G$. Let $H^{(n)}$ denote the distribution function of S_n . Then

$$H^{(n)} = F \times (F \times G)^{(n-1)} \quad (3)$$

where B^{*n} denotes the n -fold convolution of a distribution B and as usual B^{*0} equals the distribution with mass of 1 at 0. Note that we have

$$M(t) = \sum_{n=1}^{\infty} H^{(n)}(t) \quad (4)$$

as $N_t = \sum_{n=1}^{\infty} I(S_n \leq t)$, where I is the indicator function. The S_n° sequence generates an ordinary renewal process N° with renewal function M° . The interarrival times, $T_k + R_k$, have distribution $F \times G$, with mean $\mu_F + \mu_G$. Let $H^{\circ(n)}$ denote the distribution function of S_n° . Then

$$H^{\circ(n)} = (F \times G)^{*n} \quad (5)$$

Let α_t denote the forward recurrence time at time t , that is, the time from t to the next event: $\alpha_t = S_{N_t+1} - t$ on $\{X_t = 1\}$, and $\alpha_t = S_{N_t^\circ+1}^\circ - t$ on $\{X_t = 0\}$. Hence, given that the system is up at time t , the forward recurrence time α_t equals the time to the next failure time. If the system is down at time t , the forward recurrence time equals the time to complete the repair. Let F_{α_t} and G_{α_t} denote the conditional

distribution functions of α_t given that $X_t = 1$ and $X_t = 0$, respectively. Then we have for $x \in \mathbb{R}$ (here $\mathbb{R} = (-\infty, \infty)$)

$$\begin{aligned} F_{\alpha_t}(x) &= P(\alpha_t \leq x | X_t = 1) \\ &= P(S_{N_t+1} - t \leq x | X_t = 1) \end{aligned} \quad (6)$$

and

$$\begin{aligned} G_{\alpha_t}(x) &= P(\alpha_t \leq x | X_t = 0) \\ &= P(S_{N_t^\circ+1}^\circ - t \leq x | X_t = 0) \end{aligned} \quad (7)$$

Similarly for the backward recurrence time, we define β_t , F_{β_t} , and G_{β_t} . The backward recurrence time β_t equals the age of the system if the system is up at time t , and the duration of the repair if the system is down at time t , that is, $\beta_t = t - S_{N_t}^\circ$ on $\{X_t = 1\}$, and $\beta_t = t - S_{N_t}$ on $\{X_t = 0\}$.

Point Availability

The point availability $A(t)$, defined by $A(t) = P(X_t = 1)$, is given by

$$\begin{aligned} A(t) &= \bar{F}(t) + \int_0^t \bar{F}(t-x) dM^\circ(x) \\ &= \bar{F}(t) + \bar{F} \times M^\circ(t) \end{aligned} \quad (8)$$

where $\bar{F} = 1 - F$. This follows by using a standard renewal argument. Conditioning on the duration of $T_1 + R_1$, it is not difficult to see that $A(t)$ satisfies the following equation

$$A(t) = \bar{F}(t) + \int_0^t A(t-x) d(F \times G)(x) \quad (9)$$

Hence by the renewal equation (see, e.g., Aven and Jensen [2, p. 245]), Formula (8) is obtained. Alternatively, we may use a more direct approach, writing

$$X_t = I(T_1 > t) + \sum_{n=1}^{\infty} I(S_n^\circ \leq t, S_n^\circ + T_{n+1} > t) \quad (10)$$

which gives equation (8) by taking expectations. The point unavailability $\bar{A}(t)$ is given by $\bar{A}(t) = 1 - A(t) = F(t) - \bar{F} \times M^\circ(t)$.

Calculating the availability at a given point in time t , $A(t)$ is difficult except for a few special cases.

If, for example, the lifetime and repair times are independent and they have exponential distributions, it can be shown using Markov theory (see, e.g., Rausand and Høyland [1]) that the availability of the system at time t equals

$$\frac{\mu}{\lambda + \mu} + \frac{\lambda}{\lambda + \mu} e^{-(\lambda + \mu)t} \quad (11)$$

where λ is the failure rate and μ is the repair rate of the unit. It is assumed that the unit is functioning at time zero.

By the key renewal theorem (see, e.g., Aven and Jensen [2, p. 247]), it follows that

$$\lim_{t \rightarrow \infty} A(t) = \frac{\mu_F}{\mu_F + \mu_G} \quad (12)$$

The right-hand side of equation (12) is called the *limiting availability* (or steady-state availability) and is, for short, denoted by A . The *limiting unavailability* is defined as $\bar{A} = 1 - A$. Usually μ_G is small compared to μ_F , so that

$$\bar{A} \approx \frac{\mu_G}{\mu_F} \quad (13)$$

If the system has an exponential lifetime distribution with failure rate λ , it can be shown that

$$\bar{A}(t) \leq \lambda \int_0^t \bar{G}(s) ds \leq \lambda \mu_G \quad (14)$$

See Aven and Jensen [2, p. 107].

Interval Availability

The interval availability $A[u, v]$ is defined by $A[u, v] = P(X_t = 1, \forall t \in [u, v])$. Using $A[u, v] = A(u)P(\alpha_u > v - u | X_u = 1)$, we see that

$$A[u, v] = \bar{F}_{\alpha_u}(v - u)A(u) \quad (15)$$

The interval availability for the interval $[t, t + w]$, that is, the probability that the system is up at time t and the forward recurrence time at time t is greater than w is given by

$$\begin{aligned} A[t, t + w] &= P(X_t = 1, \alpha_t > w) \\ &= \bar{F}(t + w) \\ &\quad + \int_0^t \bar{F}(t - x + w) dM^\circ(x) \end{aligned} \quad (16)$$

This is shown by noting that

$$X_t I(\alpha_t > w) = \sum_{n=0}^{\infty} I(S_n^\circ \leq t, S_n^\circ + T_{n+1} > t + w) \quad (17)$$

By applying the key renewal theorem (see, e.g., Aven and Jensen [2, p. 247]), it follows that

$$\lim_{t \rightarrow \infty} A[t, t + w] = \frac{\int_0^\infty \bar{F}(x) dx}{\mu_F + \mu_G} \quad (18)$$

The limit $\int_0^\infty \bar{F}(x) dx / (\mu_F + \mu_G)$ is the asymptotic distribution of the forward recurrence time α_t and is denoted by $F_\infty(w)$. Similarly, we define the asymptotic distribution $G_\infty(w)$ of the backward recurrence time β_t .

Downtime Distribution

Let Y_u represent the downtime in the interval $[0, u]$. Then the expected downtime in $[0, u]$ is given by

$$EY_u = \int_0^u \bar{A}(t) dt \quad (19)$$

This result follows by observing that

$$\begin{aligned} EY_u &= E \int_0^u (1 - X_t) dt \\ &= \int_0^u E(1 - X_t) dt \\ &= \int_0^u \bar{A}(t) dt \end{aligned} \quad (20)$$

Using equation (19), we see that, asymptotically, the (expected) proportion of time the system is down equals the limiting unavailability, that is,

$$\lim_{u \rightarrow \infty} \frac{EY_u}{u} = \bar{A} \quad (21)$$

Furthermore, with probability one,

$$\lim_{u \rightarrow \infty} \frac{Y_u}{u} = \bar{A} \quad (22)$$

The renewal reward theorem (see, e.g., Aven and Jensen [2, p. 250]) is used to prove equation (22).

The following results can be shown for the distribution of the downtime (see Aven and Jensen [2, pp. 109–111]):

4 Availability Models

Theorem 1 *The distribution of the downtime in a time interval $[0, u]$ is given by*

$$P(Y_u \leq y) = \sum_{n=0}^{\infty} G^{*n}(y) [F^{*n}(u - y) - F^{*(n+1)}(u - y)] \quad (23)$$

Different versions of this result have been formulated and proved in the literature, for example [5, 6]. The first version was established by Takács. Theorem 1 is from Haukås and Aven [7].

In the case that the interval is rather long, the downtime will be approximately normally distributed, as is shown in Theorem 2.

Theorem 2 *The asymptotic distribution of Y_u as $u \rightarrow \infty$, is given by*

$$\sqrt{u} \left(\frac{Y_u}{u} - \bar{A} \right) \xrightarrow{D} N(0, \tau^2) \quad (24)$$

where

$$\tau^2 = \frac{\mu_F^2 \sigma_G^2 + \mu_G^2 \sigma_F^2}{(\mu_F + \mu_G)^3} \quad (25)$$

Here $\xrightarrow{D}$ means convergence in distribution and $N(0, \tau^2)$ refers to a normally distributed random variable with mean zero and variance τ^2 . The theorem follows by applying the central limit theorem for **renewal processes**, see Aven and Jensen [2, p. 150].

Steady-State Distribution

The asymptotic results established above provide good approximations for the performance measures related to a given point in time or an interval. Based on the asymptotic values, we can define a stationary (steady-state) process having these asymptotic values as their distributions and means. To define such a process in our case, we generalize the model analyzed above by allowing X_0 to be 0 or 1. Thus the time evolution of the process is as shown in Figure 1 beginning with an uptime, or similarly starting with a downtime. The process is characterized by the parameters $A(0)$, $F^*(t)$, $F(t)$, $G^*(t)$, $G(t)$, where $F^*(t)$ denotes the distribution of the first uptime provided that the system starts in state 1 at time 0 (i.e., $X_0 = 1$) and $G^*(t)$ denotes the distribution of the first downtime provided that the system starts

in state 0 at time 0 (i.e., $X_0 = 0$). Now assuming that $F^*(t)$ and $G^*(t)$ are equal to the asymptotic distributions of the recurrence times, that is, $F_\infty(t)$ and $G_\infty(t)$, respectively, and $A(0) = A$, then it can be shown that the process (X_t, α_t) is stationary, see Birolini [5]. This means that we have, for example,

$$A(t) = A, \quad \forall t \in \mathbb{R}_+ \quad (26)$$

$$A[u, u + w] = \frac{\int_w^\infty \bar{F}(x) dx}{\mu_F + \mu_G}, \quad \forall u, w \in \mathbb{R}_+ = [0, \infty] \quad (27)$$

Monotone Systems

Consider now a binary monotone system comprising n independent components. For each component, we define a model as in the section titled “Basic Setup”, indexed by “ i ”. The uptimes and downtimes of component i are thus denoted by T_{ik} and R_{ik} , with distributions F_i and G_i , respectively. The lifetime distribution F_i is absolutely continuous with a failure rate function $\lambda_i(t)$. Let $\Phi : \{0, 1\}^n \rightarrow \{0, 1\}$ be the structure function of the system, which is assumed to be monotone, that is, the structure function Φ is nondecreasing in each argument, and $\Phi(\mathbf{0}) = 0$ and $\Phi(\mathbf{1}) = 1$. The possible states of the system, “functioning” and “not functioning”, are indicated by “1” and “0”, respectively, as for the components. Then $\Phi_t = \Phi(\mathbf{X}_t)$ describes the state of the system at time t , where $\mathbf{X}_t = (X_t(1), \dots, X_t(n))$.

Let p_i denote the reliability of component i , that is, $p_i = P(X_i = 1)$. Then if the components (i.e., the X_i ’s) are independent, the reliability of the system, $P(\Phi = 1)$, is a function of the p_i ’s. This function is referred to as the *reliability function* and is denoted by $h(\mathbf{p})$, where $\mathbf{p} = (p_1, p_2, \dots, p_n)$.

Point Availability

The following results show that the point availability (limiting availability) of a monotone system is equal to the reliability function h with the component reliabilities replaced by the component availabilities $A_i(t)$ (A_i).

Theorem 3 The system availability at time t , $A(t)$, and the limiting system availability, $\lim_{t \rightarrow \infty} A(t)$, are given by

$$A(t) = h(A_1(t), A_2(t), \dots, A_n(t)) = h(\mathbf{A}(t)) \quad (28)$$

$$\lim_{t \rightarrow \infty} A(t) = h(A_1, A_2, \dots, A_n) = h(\mathbf{A}) \quad (29)$$

Equation (28) is simply an application of the reliability function formula with $A_i(t) = P(X_t(i) = 1)$. Since the reliability function $h(\mathbf{p})$ is a linear function in each p_i , and therefore a continuous function, it follows that $A(t) \rightarrow h(A_1, A_2, \dots, A_n)$ as $t \rightarrow \infty$, which proves equation (29).

The limiting system availability can also be interpreted as the expected proportion of time the system is operating in the long run, or as the long-run average availability, noting that

$$\begin{aligned} \lim_{t \rightarrow \infty} E \left[\frac{1}{t} \int_0^t \Phi_s \, ds \right] &= \lim_{t \rightarrow \infty} \frac{1}{t} \int_0^t A(s) \, ds \\ &= \lim_{t \rightarrow \infty} A(t) \end{aligned} \quad (30)$$

Interval Availability

In general, it is difficult to calculate the interval availability for a monotone system. Only in some special cases, it is possible to obtain practical computation formulae, and in the following we take a closer look at some of these.

If the repair times are small compared to the lifetimes, and the lifetimes are exponentially distributed with parameter λ_i , then clearly the number of failures of component i in the time interval $(u, u + w]$, $N_{u+w}(i) - N_u(i)$, is approximately Poisson distributed with parameter $\lambda_i w$. Hence the interval availability $A[0, u]$, which is equal to $P(N_u = 0)$, is approximately exponentially distributed.

If the system is a series system, and we make the same assumptions as above, it is also clear that the number of system failures in the interval $(u, u + w]$ is approximately Poisson distributed with parameter $\sum_{i=1}^n \lambda_i w$. The number of system failures in $[0, t]$, N_t , is approximately a **Poisson process** with intensity $\sum_{i=1}^n \lambda_i$.

If the system is highly available and the components have constant failure rates, the Poisson distribution will in fact produce good approximations also for

more general systems. To motivate this, we observe that the expected number of system failures per unit of time during the interval $(u, u + w]$, $EN_{(u, u+w]}/w$ is approximately equal to the asymptotic system failure rate λ_Φ , and $N_{(u, u+w]}$ is “nearly independent” of the history of N up to u , noting that the process $\mathbf{X}$ frequently restarts itself probabilistically (i.e., $\mathbf{X}$ reenters the state $(1, 1, \dots, 1)$). The system failure rate λ_Φ is given by

$$\lambda_\Phi = \lim_{u \rightarrow \infty} \frac{EN_u}{u} = \sum_{i=1}^n \frac{h(1_i, \mathbf{A}) - h(0_i, \mathbf{A})}{\mu_{F_i} + \mu_{G_i}} \quad (31)$$

where $h(x_i, \cdot)$ equals the reliability of the system given that the state of component i is x_i . For a proof of the second equality of equation (31), see Aven and Jensen [2, p. 114].

The detailed formulation and proof of the asymptotic result for the interval availability converting to an exponential distribution is lengthy and will not be stated here. We refer to Aven and Jensen [2, p. 119].

Downtime Distribution

In this section we study the distribution of the **system downtime** in a time interval. We immediately observe that the asymptotic expression for the expected average downtime for the one-component case also holds for monotone systems, with $A = h(\mathbf{A})$. Equation (22) requires that the process $\mathbf{X}$ is a regenerative process with finite expected cycle length.

Assume now that the components have constant failure rate and that the components are highly available, that is, the products $\lambda_i \mu_{G_i}$ are small. Then it can be heuristically argued that (Y_u) is approximately a compound Poisson process CP_u ,

$$Y_u \approx \sum_{i=1}^{N_u} Y_i \approx CP_u \quad (32)$$

Here N_u is the number of system failures in $[0, u]$ and Y_i is the downtime of the i th system failure. The dependency between N_u and the random variables Y_i is not “very strong” since N_u is mainly governed by the renewal cycles without system failures. We can ignore downtimes Y_i being the second, third, and so on system failure in a renewal cycle of the process $\mathbf{X}$. The probability of having two or more system failures in a cycle is small since we are assuming highly

available components. This means that the random variables Y_i are approximately i.i.d.

From this we can find an approximate expression for the distribution of Y_u .

A closely related approximation can be established by considering system operational time, as described in the following.

Let N_s^{op} be the number of system failures in $[0, s]$ when we consider operational time. Then, N_s^{op} is approximately a homogeneous Poisson process with intensity λ'_Φ , where λ'_Φ is given by

$$\lambda'_\Phi = \sum_{i=1}^n \frac{h(1_i, \mathbf{A}) - h(0_i, \mathbf{A})}{(\mu_{F_i} + \mu_{G_i})h(\mathbf{A})} \quad (33)$$

To motivate this result, we note that the expected number of system failures per unit of time when considering calendar time is approximately equal to the asymptotic (steady-state) system failure rate λ_Φ , given by (cf. equation (31)),

$$\lambda_\Phi = \sum_{i=1}^n \frac{h(1_i, \mathbf{A}) - h(0_i, \mathbf{A})}{\mu_{F_i} + \mu_{G_i}} \quad (34)$$

Then observing that the ratio between calendar time and operational time is approximately $1/h(\mathbf{A})$, we see that the expected number of system failures per unit of time when considering operational time, $EN^{op}(u, u+w)/w$, is approximately equal to $\lambda_\Phi/h(\mathbf{A})$.

Furthermore, $N_{[u, u+w]}^{op}$ is “nearly independent” of the history of N^{op} up to u , noting that the state process $\mathbf{X}$ frequently restarts itself probabilistically (i.e., $\mathbf{X}$ reenters the state $(1, 1, \dots, 1)$). It can be shown by the argument as for N that $N_{t/\alpha}^{op}$ has an asymptotic Poisson distribution with parameter t . The system downtimes, given system failure, are approximately identically distributed with distribution function $G_\Phi(y)$, say, independent of N^{op} , and approximately independent observing that the state process $\mathbf{X}$ with a high probability restarts itself quickly after a system failure. The distribution function $G_\Phi(y)$ is normally taken as the asymptotic (steady-state) downtime distribution, given system failure, or an approximation to this distribution. For a

parallel system of n identical components, the asymptotic (steady-state) downtime distribution, given system failure, equals

$$G_\Phi(y) = 1 - \bar{G}(y) \left[\frac{\int_y^\infty \bar{G}(x) dx}{\mu_G} \right]^{n-1} \quad (35)$$

We refer to Aven and Jensen [2, Section 4.5], for a proof of this and for derivation of distributions in the general case.

Considering the system as a one-unit system, we can now apply the exact equation (23) for the downtime distribution with the Poisson parameter λ'_Φ . It follows that

$$P(Y_u \leq y) \approx \sum_{n=0}^{\infty} G_\Phi^{*n}(y) \frac{[\lambda'_\Phi(u-y)]^n}{n!} e^{-\lambda'_\Phi(u-y)} \quad (36)$$

Equation (36) gives good approximations for “typical real life cases” with small component availabilities, see [7].

Finally, in this section, a remark on the convergence to the **normal distribution**. The theorem below is “a time average result” – it is not required that the system is highly available. The result generalizes equation (24).

Theorem 4 *If $\mathbf{X}$ is a regenerative process with cycle length S and associated downtime $Y = Y_S$, $\text{Var}[S] < \infty$ and $\text{Var}[Y] < \infty$, then as $t \rightarrow \infty$,*

$$\sqrt{t} \left[\frac{Y_t}{t} - \bar{A} \right] \xrightarrow{D} N(0, \tau_\Phi^2) \quad (37)$$

where

$$\tau_\Phi^2 = \frac{\text{Var}(Y - \bar{A}S)}{ES} \quad (38)$$

$$\bar{A} = \frac{EY}{ES} \quad (39)$$

The result follows by applying the central limit theorem for renewal reward processes.

Final Remarks

This article covers only a small number of availability models compared to the large number of models

presented in the literature, see the survey paper by Smith *et al.* [8]. Our focus has been the basic models. Examples of models not covered are models to incorporate **preventive maintenance** and standby systems. The former category of models is often developed with the purpose of optimizing maintenance (replacement) policies. We refer to Aven and Jensen [2, Chapter 5] and its Bibliographic Notes. See also Dijkhuizen and Heijden [9]. For the latter category of models, we refer to Aven and Jensen [2, Section 4.7.3]. See also the recent papers by Zhang *et al.* [10] and Montoro-Cazorla and Pérez-Ocón [11], and the references therein.

Furthermore, we have not included models where some components remain in “suspended animation” while a component is being repaired/replaced. When repair of the failed component is completed, the remaining components resume operation. At any instant, they are not “as good as new”, but rather as good as they were when the system stopped operating. Then it can be shown that the long-run (expected) proportion of time the system is functioning equals

$$\frac{1}{1 + \sum_{i=1}^n \frac{MTTR_i}{MTTF_i}} \quad (40)$$

see Barlow and Proschan [4] and Aven [12]. This expression is to be compared to

$$h(\mathbf{A}) = \prod_{i=1}^n \frac{MTTF_i}{MTTF_i + MTTR_i} \quad (41)$$

If the component availabilities are high, then

$$h(\mathbf{A}) \approx 1 - \sum_{i=1}^n \bar{A}_i \quad (42)$$

which is seen to be approximately equal to (40).

Markov models have been mentioned above – such models are often used for availability analysis. An additional book reference covering a number of Markov models for availability analysis is Misra [13]. Markov models are used to describe and analyze the movement of the system among various states, from the perfect functioning state to the complete failure state. Markov models are particularly suitable for analyzing smaller systems with complicated **maintenance policies** and possible dependencies between components. On the other

hand, the computational problems can be significant and the assumption of exponential distributions may seem to restrict applicability. However, the so-called phase-type approach can open for other distributions. The phase-type approach makes use of the fact that a distribution function can be approximated by a mixture of Erlang distributions (with the same scale parameter), for example, Asmussen [14] and Tijms [15].

Markov models are a type of multistate systems, and many of the availability results obtained for the binary case have been extended to such systems, see Aven and Jensen [2, Section 4.7.1]. See also the recent paper by Zio *et al.* [16] and the references therein. Finally, we would like to draw attention to semi-Markov models for availability analysis of multistate systems. We refer to Csenki [17] and Limnios [18].

References

- [1] Rausand, M. & Høyland, A. (2004). *System Reliability Theory*, 2nd Edition, John Wiley & Sons, New York.
- [2] Aven, T. & Jensen, U. (1999). *Stochastic Models of Reliability*, Springer, New York.
- [3] Barlow, R. & Proschan, F. (1965). *Mathematical Theory of Reliability*, John Wiley & Sons, New York.
- [4] Barlow, R. & Proschan, F. (1975). *Statistical Theory of Reliability and Life Testing*, Holt, Rinehart and Winston, New York.
- [5] Birolini, A. (1994). *Quality and Reliability of Technical Systems*, Springer, Berlin.
- [6] Takács, L. (1957). On certain sojourn time problems in the theory of stochastic processes, *Acta Mathematica Academiae Scientiarum Hungaricae* **8**, 169–191.
- [7] Haukås, H. & Aven, T. (1996). Formulae for the downtime distribution of a system observed in a time interval, *Reliability Engineering and System Safety* **52**, 19–26.
- [8] Smith, M., Aven, T., Dekker, R. & van der Duyn Schouten, F.A. (1997). A survey on the interval availability of failure prone systems, in Proceedings ESREL’97 Conference, Lisbon, 17–20 June 1997, pp. 1727–1737.
- [9] van Dijkhuizen, G. & van der Heijden, M. (1999). Preventive maintenance and the interval availability distribution of an unreliable production system, *Reliability Engineering and System Safety* **66**, 13–27.
- [10] Zhang, T., Xie, M. & Horigome, M. (2006). Availability and reliability of k-out-of-(M+N): G warm standby systems, *Reliability Engineering and System Safety* **91**, 381–387.
- [11] Montoro-Cazorla, D. & Pérez-Ocón, R. (2006). A deteriorating two-system with two repair modes and sojourn

8 Availability Models

- times phased-type distributed, *Reliability Engineering and System Safety* **91**, 1–9.
- [12] Aven, T. (1992). *Reliability and Risk Analysis*, Elsevier, Oxford.
- [13] Misra, K.B. (1992). *Reliability Analysis and Prediction*, Elsevier, Oxford.
- [14] Asmussen, S. (1987). *Applied Probability and Queues*, John Wiley & Sons, New York.
- [15] Tijms, H.C. (1994). *Stochastic Modelling and Analysis: A Computational Approach*, John Wiley & Sons, New York.
- [16] Zio, E., Marella, M. & Pedofillini, L. (2007). A Monte Carlo simulation approach to the availability assessment of multi-state systems with operational dependencies, *Reliability Engineering and System Safety*, **92**, 871–882.
- [17] Csenki, A. (1995). An integral equation approach to the interval reliability of systems modelled by finite semi-Markov processes, *Reliability Engineering and System Safety* **47**, 37–45.
- [18] Limnios, N. (1997). Dependability analysis of Semi-Markov systems, *Reliability Engineering and System Safety* **55**, 203–207.

Related Articles

System Availability; System Downtime Distributions.

TERJE AVEN

Stochastic Deterioration

Introduction

In order to ensure that an object (system or structure or one of its components) performs its required functions (including safety), maintenance should be performed on a regular basis. To optimize maintenance (*see* **Maintenance Optimization**), information about the time to failure and/or the deteriorating condition is needed. Unfortunately, there is a lot of uncertainty involved in the deterioration of an object. Mathematical models can be used to account for the stochastic nature of deterioration over time. For this purpose, we distinguish two approaches in modeling deterioration [1, Chapter 1]: the actuarial approach based on the notion of failure rate and the physical approach based on the notion of failure due to the stress exceeding the strength (*see* **Stress–Strength Model**).

Failure Rate Function

A lifetime distribution represents the uncertainty in the time to failure of an object. Let the lifetime have a cumulative probability distribution $F(t)$ with **probability density function** $f(t)$, then the failure rate (or hazard rate) function is defined as

$$r(t) = \frac{f(t)}{1 - F(t)} = \frac{f(t)}{\bar{F}(t)}, \quad t > 0 \quad (1)$$

(Barlow and Proschan [2, Chapter 2]). A probabilistic interpretation of the failure rate function is that $r(t) dt$ represents the probability that an object of age t will fail in the time interval $[t, t + dt]$. For deteriorating objects, the failure rate is increasing. By integrating both sides of equation (1) and then exponentiating, we obtain the survival function

$$\bar{F}(t) = \exp \left\{ - \int_0^t r(\tau) d\tau \right\} \quad (2)$$

Lifetime distributions and failure rate functions are especially useful in mechanical and electrical engineering. In these fields, one often considers equipment which can assume at most two states: the functioning state and the failed state. A structure, on the other hand, can be in a range of states depending on its degrading condition. A disadvantage of failure

rates is that they cannot be observed or measured for a particular object [3].

Time-Dependent Stress–Strength Model

The physical approach to deterioration modeling identifies an object's resistance R (or strength), on the one hand, and its applied stress S (or load), on the other hand. Failure is defined as the event at which the stress is greater than the resistance; that is, $S > R$. Assuming both the resistance and the strength to be a random quantity, the reliability is defined as the probability $\Pr\{R > S\}$. In a physics-based deterioration model, it is recommended to model deterioration in terms of a time-dependent stochastic process $\{X(t), t \geq 0\}$, where $X(t)$ is a random quantity for all $t \geq 0$. A resistance subject to deterioration is then defined as $R(t) = r_0 - X(t)$, where r_0 is the initial resistance. For modeling stochastic deterioration, we can use stochastic processes such as a random deterioration-rate model or a **Markov process**.

Random Deterioration-Rate Model

In structural reliability, time-dependent **degradation** functions are advocated for which the coefficients (such as the average deterioration rate) are random quantities [4, 5]. These stochastic processes are also called *degradation-path models*. An example of such a model is the cumulative amount of deterioration at time t being defined as $X(t) = At$ for all $t \geq 0$, where the average deterioration rate A has a probability distribution. However, according to Pandey *et al.* [6], the sample paths of such models are straight lines and a single inspection thus fixes the future deterioration beforehand. Although the random deterioration-rate model might be used as an approximation, one should be careful as soon as inspections are involved. For the purpose of inspection and maintenance modeling, we therefore have to rely on other stochastic processes (such as Markov processes), which properly capture the temporal variability associated with the evolution of deterioration.

Markov Process

Deterioration is usually assumed to be a Markov process [2, Chapter 5] (*see* **Markov Processes**). Roughly speaking, a Markov process is a stochastic process

2 Stochastic Deterioration

with the property that, given the value of $X(t)$, the values of $X(\tau)$, where $\tau > t$, are independent of the values of $X(u)$, $u < t$. That is, the conditional distribution of the future $X(\tau)$, given the present $X(t)$ and the past $X(u)$, $u < t$, is independent of the past. Classes of Markov processes which are useful for modeling stochastic deterioration are discrete-time Markov processes having a finite or countable state space (called *Markov chains*) and continuous-time Markov processes with independent increments such as the **Brownian motion** with drift, the compound **Poisson process**, the Lévy process, and the gamma process. The assumption of independent increments is more restrictive than the Markov property. Because the increment $X(\tau) - X(t)$ is independent of $X(t)$ and $X(\tau) = X(t) + [X(\tau) - X(t)]$, the stochastic process $\{X(t), t \geq 0\}$ is Markovian. Hence, damage accumulation with independent increments naturally leads to the Markov property.

Finite-State Markov Process. Assume that there is a finite or countable state space. This means that the condition of an object can be in any one of $n > 0$ discrete states or categories. A Markov chain is a discrete-time stochastic process $\{X_t, t = 0, 1, 2, \dots\}$ for which the Markovian property holds. To define the Markov chain X_t , it is necessary to assess the transition probabilities between all possible condition state pairs. If there are n states, then this results in an $n \times n$ -matrix $P = \|p_{ij}\|$, where the probability of a transition between states i and j during one unit of time is denoted by $p_{ij} = \Pr\{X_{t+1} = j | X_t = i\}$. If this transition probability does not depend on t (i.e., does not depend on how long the process has been running), then the process is called *stationary* in time. Obviously, the probability of moving from one state to any other state (including itself) should be 1, or $\sum_{j=1}^n p_{ij} = 1$ for $i, j = 1, \dots, n$. The probability p_{ij}^m of moving from the current state i to state j in m time steps is $p_{ij}^m = \Pr\{X_{t+m} = j | X_t = i\}$. It can be shown [7, Chapter 4] that the m -step transition matrix for a stationary Markov chain can be calculated by multiplying the transition probability matrix m times with itself, P^m .

A semi-Markov process is an extension of a Markov chain in which a random time is added between transitions. Let J_0 be the state of the process $\{X(t), t \geq 0\}$ at the beginning of the process and $J_n, n = 1, 2, \dots$ the state of $X(t)$ after n transitions. The probability of the process moving into state j

in an amount of time less than or equal to t , given that it just moved into state i , is defined as $Q_{ij}(t) = \Pr\{T \leq t, J_{n+1} = j | J_n = i\}$. This probability can be written as the product

$$Q_{ij}(t) = F_{ij}(t)p_{ij} \quad (3)$$

where $p_{ij} = \Pr\{J_{n+1} = j | J_n = i\}$ is the transition probability function of the embedded Markov chain $\{J_n, n = 0, 1, 2, \dots\}$ and $F_{ij}(t) = \Pr\{T \leq t | J_{n+1} = j, J_n = i\}$ represents the conditional probability of the random waiting time T , given that the process moves into state j after previously having moved into state i . Equation (3) shows that transitions in a semi-Markov process have two stages: if the process just moved into state i , it first selects the next state j according to p_{ij} and then waits a random time T according to $F_{ij}(t)$. The semi-Markov process $\{X(t), t \geq 0\}$ may be defined as $X(t) = J_{N(t)}$, where $N(t)$ is the total number of transitions during the interval $(0, t]$. For deterioration modeling, the waiting time distribution is often chosen to have an increasing failure rate to model the increasing probability of the occurrence of a transition the longer an object is in a state. In the case where the waiting times have an exponential distribution, which is a memoryless distribution, the semi-Markov processes are Markovian everywhere. This particular type of processes is referred to as *continuous-time Markov processes* and they are widely used due to their analytical tractability.

Finite-state Markov processes are particularly well suited for modeling the uncertain rate at which objects pass through a finite number of states, categories, or phases (see **Maintenance and Markov Decision Models**). These categories typically range from “very good” to “very bad”, or from “no damage” to “structural failure”, and are often used when the condition of an object can only be assessed by means of visual inspections. If a Markov process is to be used for deterioration modeling, the states or phases should be progressive (i.e., the probability of an improvement is zero) and in many practical applications, it is also assumed that transitions are sequential (i.e., they always proceed from one state to the next in the sequence). This resembles the physical process of deterioration more closely, where objects pass through each successive state before failing.

The parameters of the model, i.e., the transition probabilities or intensities, should be estimated based on expert judgment or based on the available data.

The method of **estimation** depends on the level of detail in the observations of the process. The most information is obtained when the process is continuously monitored, such that each transition in the sequence is observed. For a Markov chain, the maximum-likelihood estimator for each transition probability is then simply $\hat{p}_{ij} = w_{ij}(\sum_{j=1}^n w_{ij})^{-1}$, where w_{ij} is the observed number of transitions from state i to j [8]. The same estimator can be used for the embedded Markov chain in a semi-Markov process and the observed times between transitions allow for fitting the waiting time distributions using classical estimation techniques. This type of detailed information also allows for a Bayesian approach to the estimation of transition probabilities in a Markov chain, see e.g., [9, Chapter 2].

In many practical applications, the process cannot be monitored continuously. If an object is inspected periodically, the type of information that is available is often referred to as *panel data*. Typically, the time between inspections of different objects varies and some objects may be inspected more often than others. Semi-Markov processes are already quite difficult to fit to this type of data, but it becomes more straightforward if deterioration is modeled by a continuous-time Markov process. For this case, the likelihood of w observations $\mathbf{x} = (x_1, \dots, x_w)'$ is $L(\mathbf{x}|\boldsymbol{\theta}) = \prod_{k=1}^w f(x_k|\boldsymbol{\theta})$. The likelihood can be maximized as a function of the parameters $\boldsymbol{\theta}$, see [10], to obtain an estimate for the transition intensities. The same approach can be used for Markov chains using the probability $f(x|\boldsymbol{\theta}) = (P^t)_{ij}$.

In some cases, there may only be aggregate data available which contains just the number or the proportion of objects in each condition state. This is the least-detailed information as the state sequence of individual objects is not known to the modeler. The estimation of the transition parameters based on aggregate data requires special care, which is outside the scope of this article.

In the field of civil engineering, Markov chains have been used to model uncertain deterioration in a number of areas like pavement, bridge, and sewer system management. One of the first examples is the Arizona pavement management system [11], which inspired the Pontis bridge management system [12]. For a more complete review of applications in bridge management, see [5]. More recently, Markov chains have been applied to sewer system and water pipeline deterioration, see e.g., [13] and [14].

Brownian Motion with Drift. The Brownian motion with drift (see **Brownian Motion**) is a continuous-time stochastic process $\{X(t), t \geq 0\}$ for which $X(t)$ is normally distributed with mean μt and variance $\sigma^2 t$ for all $t \geq 0$. In the physical approach to reliability, a resistance modeled as a Brownian motion with drift alternately increases and decreases. For this reason, the Brownian motion is often inadequate in modeling deterioration which is monotone. On the other hand, it has certain mathematical advantages. For example, explicit expressions are available for the probability distribution of the first-passage time to zero of strength minus stress, when stress and strength are assumed to be independent Brownian motions with drift [15] (see **Stress–Strength Model**).

Compound Poisson Process. Mercer and Smith [16] may have been the first to model deterioration as a stochastic process. In particular, they propose to model deterioration as a compound Poisson process. A compound Poisson process is a monotone stochastic process with independent and identically distributed deterioration increments which occur according to a Poisson process (see **Poisson Processes**). They define failure as the event in which the deterioration reaches such a level that the object is considered to be worn out completely. Probability distributions of the time of failure (lifetime) of objects subjected to a sequence of shocks occurring randomly in time as events in a homogeneous or a nonhomogeneous compound Poisson process are further studied in [17] and [18–20], respectively.

Lévy Process. Roughly put, continuous-time stochastic processes with independent and stationary increments, and with right-continuous sample paths having left limits are Lévy processes (see **Lévy Processes**). In mathematical terms, a stochastic process has stationary increments if the probability distribution of the increments $X(t+h) - X(t)$ depends only on h for all $t, h \geq 0$. Every Lévy process may be written as a sum of a Brownian motion, a deterministic part (linear in time), and an integral of compound Poisson processes, where all the contributing processes are mutually independent. The sample paths of a Lévy process are discontinuous with probability one if the process is monotone, because such a process can be decomposed into a linear part plus an integral of compound Poisson processes [21]. The

sample paths of a Brownian motion are continuous with probability one.

Gamma Process. A gamma process is a monotone stochastic process with independent, nonnegative increments having a gamma distribution with an identical scale parameter [21, 22]. Like the compound Poisson process, the gamma process is a jump process. Compound Poisson processes have a finite number of jumps in finite time intervals, whereas gamma processes have an infinite number of jumps in finite time intervals. The former are suitable for modeling usage such as damage due to sporadic shocks and the latter are suitable for describing gradual damage by continuous use. In the presence of inspection data, the advantage of the gamma process over the compound Poisson process is evident: measurements generally consist of deterioration increments rather than of jump intensities and jump sizes.

As far as the authors know, Abdel-Hameed [23] was the first to propose the gamma process as a proper model for deterioration occurring random in time. The gamma process is suitable to model gradual damage monotonically accumulating over time in a sequence of tiny increments, such as concrete creep [24], corrosion [5, 25], wear, fatigue, crack growth, erosion, swell, degrading performance index, and so on (*see Physical Degradation Models*). The gamma deterioration process was successfully applied to model condition-based maintenance (*see Maintenance Optimization*). For an overview, we refer to [26].

In mathematical terms, the gamma process is defined as follows. A random quantity X has a gamma distribution with shape parameter $v > 0$ and scale parameter $u > 0$ if its probability density function is given by

$$\text{Ga}(x|v, u) = \frac{u^v}{\Gamma(v)} x^{v-1} \exp\{-ux\} I_{(0,\infty)}(x) \quad (4)$$

where $I_A(x) = 1$ for $x \in A$ and $I_A(x) = 0$ for $x \notin A$, and $\Gamma(a) = \int_{t=0}^{\infty} t^{a-1} e^{-t} dt$ is the gamma function for $a > 0$. Furthermore, let $v(t)$ be a nondecreasing, right-continuous, real-valued function for $t \geq 0$, with $v(0) \equiv 0$. The gamma process with shape function $v(t) > 0$ and scale parameter $u > 0$ is a continuous-time stochastic process $\{X(t), t \geq 0\}$ with the following properties: $X(0) = 0$ with probability one, $X(\tau) - X(t) \sim \text{Ga}(v(\tau) - v(t), u)$ for all $\tau > t \geq 0$,

and $X(t)$ has independent increments. If $X(t)$ is a gamma process with shape function $v(t)$ and scale parameter u , then its expectation and variance are given by

$$E(X(t)) = \frac{v(t)}{u}, \quad \text{Var}(X(t)) = \frac{v(t)}{u^2} \quad (5)$$

A component is said to fail when its deteriorating resistance, denoted by $R(t) = r_0 - X(t)$, drops below the stress s . We assume both the initial resistance r_0 and the stress s to be known. Let the time at which failure occurs be denoted by the lifetime T . Owing to the gamma-distributed deterioration, the lifetime distribution can then be written as:

$$\begin{aligned} F(t) &= \Pr\{T \leq t\} = \Pr\{X(t) \geq r_0 - s\} \\ &= \frac{\Gamma(v(t), [r_0 - s]u)}{\Gamma(v(t))} \end{aligned} \quad (6)$$

where $\Gamma(a, x) = \int_{t=x}^{\infty} t^{a-1} e^{-t} dt$ is the incomplete gamma function for $x \geq 0$ and $a > 0$. Equation (6) resembles a striking duality between space (deterioration) and time (lifetime) that makes the gamma-process model amenable to analysis. The lifetime distribution based on a gamma-process deterioration, given by equation (6), has an increasing failure rate as long as the shape function $v(t)$ is increasing in t . Abdel-Hameed [23] even took the stress–strength model in equation (6) a stage further by regarding the deterioration failure level $y = r_0 - s$ as a random quantity $Y > 0$ as well.

Under the assumption of modeling the temporal variability in the deterioration with a gamma process, the question which remains to be answered is how its expected deterioration increases over time. Empirical studies show that the expected deterioration at time t is often proportional to a power law:

$$E(X(t)) = \frac{v(t)}{u} = a t^b \quad (7)$$

for some physical constants $a > 0$ and $b > 0$ [4, 24]. There is often engineering knowledge available about the power-law exponent b in equation (7). Çinlar *et al.* [24] show how a nonstationary gamma process ($b \neq 1$) can be transformed into a stationary gamma process ($b = 1$) and how the parameters of a stationary gamma process can be estimated using the method of **moments** and the method of

maximum likelihood. Bayesian estimation of a stationary gamma deterioration process under imperfect inspection can be found in [25]. An interesting extension of the univariate gamma process is the multivariate gamma process proposed by Ebrahimi [27].

A useful property of a stationary gamma process is that the gamma density $\text{Ga}(x|aut, u)$ transforms into an exponential density if $t = 1/(au)$. When the unit time length is chosen to be $\Delta = 1/(au)$, the increments of deterioration $D_i = X(i\Delta) - X((i-1)\Delta)$, $i = 1, 2, \dots$, are exponentially distributed with mean $1/u$. For this unit time, many probabilistic properties can be expressed in analytic form [28]. For example, the probability of failure in unit time i reduces to a shifted Poisson distribution with mean $1 + [y\mu]/\sigma^2$:

$$\begin{aligned} & \Pr\{X((i-1)\Delta) \leq y, X(i\Delta) > y\} \\ &= \frac{1}{(i-1)!} \left[\frac{y\mu}{\sigma^2} \right]^{i-1} \exp\left\{-\frac{y\mu}{\sigma^2}\right\}, \\ & i = 1, 2, 3, \dots \end{aligned} \quad (8)$$

For unit time length Δ and for all $n = 1, 2, \dots$, the joint probability density function of the independent increments $D_1, \dots, D_n$ can be written as a function of the sum of the increments $X_n = \sum_{i=1}^n D_i$ (being the l_1 -norm). The infinite sequence of random quantities D_n , $n = 1, 2, \dots$ is called l_1 -norm symmetric or l_1 -isotropic. Conversely, a scale mixture of stationary gamma processes follows directly from the minimal hypothesis of l_1 -isotropy; that is, if the order in which nonnegative deterioration increments occur is irrelevant (because we are only interested in their sum), the increments are exchangeable and l_1 -isotropic [28]. Note that independence is a special case of exchangeability. Exchangeability occurs when incorporating uncertainty in the average rate of deterioration by means of a scale mixture of stationary gamma processes (see **Subjective Life Models**).

The gamma process can be regarded as a compound Poisson process of gamma-distributed increments in which the Poisson rate tends to infinity and increment sizes tend to zero in proportion. Although the gamma process can be reformulated in terms of a limit of a compound Poisson process [22], it is not efficient to simulate a gamma process in such a way. This is because there are infinitely many

jumps in each finite time interval. A better approach for simulating a gamma process is simulating independent increments with respect to very small units of time (see Ref. 30 **Simulation of the Gamma Process**).

Wenocur [29] considers a very interesting extension of a stationary gamma process for the deterioration state under two different failure modes. An object is said to fail either when its condition reaches a failure level s or when a traumatic event (such as an extreme load) destroys it. Suppose that the traumatic events occur as a Poisson process with a rate (called the *killing rate*) which depends on the object's condition. In a time-dependent reliability analysis, the killing rate can be interpreted as the frequency of the load exceeding the resistance when the variability of the random loads is modeled by a peaks-over-threshold distribution and the stochastic process of load threshold exceedances is generated by a Poisson process with intensity λ [30]. In this situation, the cumulative distribution function of the lifetime can be written as

$$\begin{aligned} F(t) = 1 - E \left(\exp \left\{ - \int_{u=0}^t k(R(u)) du \right\} \right. \\ \left. \times I_{[s, r_0]}(R(t)) \right) \end{aligned} \quad (9)$$

where $R(t) = r_0 - X(t)$ is the resistance and $k(r)$ is the “killing rate” for resistance r . In our time-dependent stress–strength model, the killing rate is $k(r) = \lambda \Pr\{L > r\}$, which is exactly the physics-based frequency of the stress or load $L = l_0 + Z$ exceeding the strength r with Z the exceedance over threshold l_0 having a peaks-over-threshold distribution.

References

- [1] Rausand, R. & Høyland, A. (2004). *System Reliability Theory; Models, Statistical Methods, and Applications*, 2nd Edition, John Wiley & Sons, Hoboken.
- [2] Barlow, R.E. & Proschan, F. (1965). *Mathematical Theory of Reliability*, John Wiley & Sons, New York.
- [3] Singpurwalla, N.D. (1995). Survival in dynamic environments, *Statistical Science* **10**(1), 86–103.
- [4] Ellingwood, B.R. & Mori, Y. (1993). Probabilistic methods for condition assessment and life prediction of concrete structures in nuclear power plants, *Nuclear Engineering and Design* **142**(2–3), 155–166.

- [5] Frangopol, D.M., Kallen, M.J. & van Noortwijk, J.M. (2004). Probabilistic models for life-cycle performance of deteriorating structures: review and future directions, *Progress in Structural Engineering and Materials* **6**(4), 197–212.
- [6] Pandey, M.D., Yuan, X.-X. & van Noortwijk, J.M. (2007). A comparison of probabilistic deterioration models for life-cycle management of structures, *Structure and Infrastructure Engineering*, in press.
- [7] Ross, S.M. (1970). *Applied Probability Models with Optimization Applications*, Dover Publications, New York.
- [8] Anderson, T.W. & Goodman, L.A. (1957). Statistical inference about Markov chains, *Annals of Mathematical Statistics* **28**, 89–110.
- [9] Lee, T.C., Judge, G.G. & Zellner, A. (1970). *Estimating the Parameters of the Markov Probability Model from Aggregate Time Series Data*, North Holland, Amsterdam.
- [10] Kalbfleisch, J.D. & Lawless, J.F. (1985). The analysis of panel data under a Markov assumption, *Journal of the American Statistical Association* **80**(392), 863–871.
- [11] Golabi, K., Kulkarni, R.B. & Way, G.B. (1982). A state-wise pavement management system, *Interfaces* **12**(6), 5–21.
- [12] Golabi, K. & Shepard, R. (1997). Pontis: a system for maintenance optimization and improvement of US bridge networks, *Interfaces* **27**(1), 71–88.
- [13] Wirahadikusumah, R., Abraham, D. & Iseley, T. (2001). Challenging issues in modeling deterioration of combined sewers, *Journal of Infrastructure Systems* **7**(2), 77–84.
- [14] Micevski, T., Kuczera, G. & Coombes, P. (2002). Markov model for storm water pipe deterioration, *Journal of Infrastructure Systems* **8**(2), 49–56.
- [15] Basu, A.P. & Ebrahimi, N. (1983). On the reliability of stochastic systems, *Statistics and Probability Letters* **1**(5), 265–267.
- [16] Mercer, A. & Smith, C.S. (1959). A random walk in which the steps occur randomly in time, *Biometrika* **46**(1–2), 30–35.
- [17] Esary, J.D., Marshall, A.W. & Proschan, F. (1973). Shock models and wear processes, *The Annals of Probability* **1**(4), 627–649.
- [18] Abdel-Hameed, M.S. & Proschan, F. (1973). Nonstationary shock models, *Stochastic Processes and their Applications* **1**, 383–404.
- [19] Zacks, S. (2004). Distributions of failure times associated with nonhomogeneous compound Poisson damage processes, in *IMS Lecture Notes, Volume 45: A Festschrift to Honor Herman Rubin*, A. Dasgupta, ed, Institute of Mathematical Statistics (IMS), Beachwood, pp. 396–407.
- [20] Zacks, S. (2006). Failure distributions associated with general compound renewal damage processes, in *Probability, Statistics and Modeling in Public Health*, M. Nikulin, D. Commenges & C. Huber, eds, Springer, New York, pp. 466–475.
- [21] Ferguson, T.S. & Klass, M.J. (1972). A representation of independent increment processes without Gaussian components, *The Annals of Mathematical Statistics* **43**(5), 1634–1643.
- [22] Dufresne, F., Gerber, H.U. & Shiu, E.S.W. (1991). Risk theory with the gamma process, *ASTIN Bulletin* **21**(2), 177–192.
- [23] Abdel-Hameed, M. (1975). A gamma wear process, *IEEE Transactions on Reliability* **24**(2), 152–153.
- [24] Çinlar, E., Bažant, Z.P. & Osman, E. (1977). Stochastic process for extrapolating concrete creep, *Journal of the Engineering Mechanics Division* **103**(EM6), 1069–1088.
- [25] Kallen, M.J. & van Noortwijk, J.M. (2005). Optimal maintenance decisions under imperfect inspection, *Reliability Engineering and System Safety* **90**(2–3), 177–185.
- [26] van Noortwijk, J.M. (2007). A survey of the application of gamma processes in maintenance, *Reliability Engineering and System Safety*, in press; DOI:10.1016/j.res.2007.03.019.
- [27] Ebrahimi, N. (2004). Indirect assessment of the bivariate survival function, *Annals of the Institute of Statistical Mathematics* **56**(3), 435–448.
- [28] van Noortwijk, J.M., Cooke, R.M. & Kok, M. (1995). A Bayesian failure model based on isotropic deterioration, *European Journal of Operational Research* **82**(2), 270–282.
- [29] Wenocur, M.L. (1989). A reliability model based on the gamma process and its analytic theory, *Advances in Applied Probability* **21**(4), 899–918.
- [30] van Noortwijk, J.M., van der Weide, J.A.M., Kallen, M.J. & Pandey, M.D. (2007). Gamma processes and peaks-over-threshold distributions for time-dependent reliability, *Reliability Engineering and System Safety* **92**(12), 1651–1658.

Related Articles

Brownian Motion; Degradation and Failure; Lévy Processes; Markov Processes; Physical Degradation Models; Poisson Processes; Stress–Strength Model; Subjective Life Models.

JAN M. VAN NOORTWIJK AND
MAARTEN-JAN KALLEN

Aging and Positive Dependence

Introduction

In the fields of reliability and quality, vectors of nonnegative random variables, say $\mathbf{T} \equiv (T_1, \dots, T_n)$, most often arise as the natural objects to be analyzed; $T_1, \dots, T_n$ may for example, have the meaning of times to failure for the n components of a reliability system or for n similar items in a quality test for some new prototype, in an industrial production.

Concerning the joint probability distribution of $\mathbf{T}$, idealized models are defined by the conditions of stochastic independence among $T_1, \dots, T_n$ and exponentiality for their marginal distributions; in such cases computations can be rather simple, and often lead to explicit and close formulas for the solution of reliability problems.

In many cases of applied interest, however, such idealized conditions must be abandoned, being not realistic at all; lack of stochastic independence is expressed by the term *dependence*, whereas lack of exponentiality leads to what we generically express by the term *aging*.

In cases of aging and dependence, not only computation complexity considerably increases but even the meaning and possible interpretations of the obtained results may sometimes become unclear or controversial. This scenario explains why the analysis of specific conditions of stochastic dependence and of probability distributions with special aging properties, and related implications have become issues of central interest in the fields of reliability and quality. The basic notions of dependence and of aging had been already singled out during 1950s to 1960s of the twentieth century see in particular [1–3]; later, in the subsequent decades there was in the literature a real explosion of other specific definitions. Parallel to that, several probabilistic models emerged from applications, which gave rise simultaneously to conditions of aging and positive dependence. A considerable part of literature about these topics concerns the possible applications of different notions of aging and dependence to special models; this is typically done to obtain useful inequalities in reliability and quality evaluations.

The relevant literature in these fields is by now really very wide and it is outside the scope and possibility of this article to provide a satisfactory review of it. The purpose here is, rather, to present a discussion about some aspects of the relations existing between aging and positive dependence, of substantial character or just of formal type. This will be essentially performed in the section titled “Multivariate Aging and Relations between Aging and Positive Dependence”, where also some definitions are recalled and some remarks are presented specifically concerning multivariate aging. In the next section we recall some basic facts and notation about stochastic orderings and univariate aging, whereas the section titled “Notions of Positive Dependence” will provide a brief overview about notions of positive dependence.

Some Basic Notions and Notation

Basic concepts of univariate aging are, in particular, IFR (*increasing failure rate*), DFR (*decreasing failure rate*), IFRA (*increasing failure rate in average*), DMRL (*decreasing mean residual life*), and NBU (*new better than used*). Several other concepts have also been defined, some among them arising as more or less direct generalizations of the former ones.

Among the basic concepts, many come out in a quite natural way when considering *stochastic comparisons* between distributions of *residual lifetimes* at different *ages* and/or monotonicity-type properties of *failure rate* or of *mean residual life* functions. First, to set notation, these notions are briefly recalled.

Let T be a *lifetime*, i.e. a nonnegative random variable with the meaning of a time to a top event (e.g., a time to a failure of a component in a system). Typically we will think of T as the lifetime of a device or a component D and assume that T admits a continuous distribution function with density function $f(t)$, i.e.,

$$P\{T \leq t\} = F(t) = \int_0^t f(u) du \quad (1)$$

We denote by $\bar{F}$ the *survival function* and by $\lambda(t)$ the *failure rate function*:

$$\bar{F}(t) = 1 - F(t) \quad (2)$$

$$\lambda(t) = \frac{f(t)}{\bar{F}(t)} \quad (3)$$

2 Aging and Positive Dependence

so that we can also write

$$\lambda(t) = \lim_{\delta \rightarrow 0} \frac{P\{T \leq t + \delta | T > t\}}{\delta} \quad (4)$$

$$\overline{F}(t) = \exp\{-\Lambda(t)\} \quad (5)$$

where

$$\Lambda(t) = \int_0^t \lambda(u) du \quad (6)$$

is the *cumulated hazard rate function*. When T admits a finite expected value $\mathbb{E}(T)$, we define the mean residual life function as

$$\mu(t) = \mathbb{E}(T - t | T > t) \quad (7)$$

where, for fixed t , $\mathbb{E}(T - t | T > t)$ is the *average remaining life* for D , if survived until time t .

A *stochastic comparison* is a notion of ordering between two probability distributions. Stochastic comparisons generally have revealed to be important tools in statistic and applied probability; they can in particular be used for giving definitions and for formalizing results concerning aging and dependence. A basic reference is [4]; see also [5–7]. Most important univariate notions for our purposes here are the *Ordinary* (or *Usual*), the *Hazard Rate*, the *likelihood ratio*, and the *mean residual life* stochastic orderings (respectively denoted by $\leq_{st}$, $\leq_{hr}$, $\leq_{lr}$, $\leq_{mrl}$). See e.g., [4] for definitions and basic properties; in particular it is shown that the following chain of implications holds:

$$\leq_{lr} \implies \leq_{hr} \implies \leq_{st}; \quad \leq_{hr} \implies \leq_{mrl} \quad (8)$$

An aspect of basic interest when studying conditions of univariate aging generally concerns the probabilistic behavior of the residual lifetimes at different *ages*. In other words attention is focused on the conditional survival probability

$$P(T > t + \tau | T > t) = \frac{\overline{F}(t + \tau)}{\overline{F}(t)} \quad (9)$$

when considered, for fixed $\tau > 0$, as a function of the age t .

In the cases when the probability distribution of T does admit a density function, and then a failure rate function $\lambda(t)$, one can also equivalently write

$$P(T > t + \tau | T > t) = \exp \left\{ - \int_t^{t+\tau} \lambda(u) du \right\} \quad (10)$$

We say that (the distribution of) a lifetime T is IFR when the function $\lambda(t)$ is nondecreasing, and that it is DFR when $\lambda(t)$ is nonincreasing. As the borderline case, we have that T is *exponentially distributed* if and only if $\lambda(t)$ is constant all over the half line $[0, +\infty)$. For fixed τ , the conditional survival probability $P(T > t + \tau | T > t)$ is a nonincreasing function of the age t , if T is increasing failure rate. Notice that, in the cases when T admits a density f , and then the failure rate λ , the two conditions are actually equivalent; but we can consider of course monotonicity conditions for $P\{T > t + \tau | T > t\}$ also for cases when T does not admit a density. A dual argument holds for the notion of decreasing failure rate.

IFR is a notion of *positive aging* whereas DFR describes *negative aging*.

Other notions of positive and negative aging studied in the literature are defined as natural generalizations, in some other sense, of the notions of IFR and DFR, respectively. A very important notion, for instance, is the one of *NBU*: T is NBU when, for any $t > 0$ and $\tau > 0$, it is $P(T > t + \tau | T > t) \leq P(T > \tau)$.

Using the language of stochastic orderings, the IFR property is expressed by the condition that the distribution of the residual lifetime $T - t'$, conditional on $\{T > t'\}$, is stochastically greater (with respect to the ordinary stochastic ordering) than that the distribution of $T - t''$, conditional on $\{T > t''\}$, for any pair $0 \leq t' < t''$. Following this line, one can characterize other notions of aging for univariate distributions, by employing other notions of stochastic ordering or by setting different conditions on t', t'' .

A detailed account is given by Shaked and Shanthikumar in [4]; an account of most recent literature is provided, among more material, in the review provided by Lai and Xie [8].

The meaning of the basic notions of univariate aging, and their interest in obtaining useful inequalities, are very clear. In particular, as basic problems in the field of reliability, one has often to decide if in specific situations it is convenient to adopt procedures of *preventive maintenance* or of **burn-in**; if so one has, furthermore, to study the problem of the optimization of such procedures; the notions of aging have a fundamental role in such a kind of analysis (see [3, 9–13] for basic references).

In the analysis of aging, a key role is obviously played by the concept of time. Thus, in a genuine probabilistic approach, it is generally fruitful, and in

some cases compulsory, to study models of aging from the point of view of the theory of stochastic processes; in particular the theory of counting processes [14] and stochastic filtering with the observations of counting processes (see [15]) is to be taken into account for a rigorous and general development of reliability problems; see in particular [16–22]. Related to this frame it is important, before concluding this section, to mention that the concept of failure rate of a lifetime T is associated in a natural way to a stochastic process as follows: consider the counting process (actually a jump process with only one jump of size 1 at the random time T) defined by the position

$$X_t = \mathbf{1}_{\{T \leq t\}}, \quad t \geq 0 \quad (11)$$

The stochastic process defined by

$$\tilde{\lambda}_t = (1 - X_t)\lambda(t), \quad t \geq 0 \quad (12)$$

is the *stochastic intensity* of the process $\{X_t\}_{t \geq 0}$, with respect its *internal history*, i.e., the intensity of failure assigned at time t on the basis of the only information about whether it is $\{T > t\}$ or $\{T \leq t\}$.

Generally, the flow of available information may, however, be richer than that; the point process approach allows for a formalized study about the changes in the failure rate when passing from one to another level of information.

It is a central point in our discussion that, along with stochastic intensity, the aging properties too do depend on the flow of information.

Another important issue concerning failure rate functions is the generalization of the univariate concept to the multivariate case. Several different definitions have been presented in the relevant literature, depending on the special type of problems or application to be dealt with; see in particular [23–26]. A very natural and useful multivariate extension leads to the definition of *multivariate conditional hazard rate* (m.c.h.r.) functions, that has been analyzed in [25, 26]. Such a definition also can be directly related to stochastic intensities of some counting processes with a finite number of jumps.

For an ordered subset $I = \{j_1, \dots, j_{|I|}\} \subset \{1, 2, \dots, n\}$, let $|I|$ denote the number of elements of I and $\bar{I}$ denote the complement of I ; $\mathbf{T}_I$ denotes the vector of lifetimes $(T_{j_1}, \dots, T_{j_{|I|}})$; $\mathbf{e}$ denotes a vector with an appropriate number of coordinates all equal to 1 so that, for $t > 0$, it is $t\mathbf{e} = (t, t, \dots, t)$; furthermore,

for $0 < t_{j_1} \leq \dots \leq t_{j_{|I|}} \leq t$ the symbol

$$\{\mathbf{T}_I = \mathbf{t}_I; \mathbf{T}_{\bar{I}} > t\mathbf{e}\} \quad (13)$$

stands for the observation

$$T_{j_1} = t_{j_1}, \dots, T_{j_{|I|}} = t_{j_{|I|}}; T_j > t, j \in \bar{I} \quad (14)$$

For $I \subset \{1, 2, \dots, n\}$, $j \notin I$, $t > 0$, $0 < t_{j_1} \leq \dots \leq t_{j_{|I|}} \leq t$, the m.c.h.r. function $\lambda_{j|I}(t_{j_1}, \dots, t_{j_{|I|}})$ is defined by

$$\begin{aligned} &\lambda_{j|I}(t_{j_1}, \dots, t_{j_{|I|}}) \\ &= \lim_{\delta \rightarrow 0} \frac{P\{T_j \leq t + \delta | \mathbf{T}_I = \mathbf{t}_I; \mathbf{T}_{\bar{I}} > t\mathbf{e}\}}{\delta} \end{aligned} \quad (15)$$

The m.c.h.r. functions can be defined when the joint distribution of lifetimes admits a joint **probability density function**. In such a case the density function can be easily obtained from the knowledge of the set of m.c.h.r. functions and *vice versa*.

In the literature, several definitions of multivariate stochastic comparison (i.e., of partial ordering between two different multivariate probability distributions) have been introduced that arise as generalizations of corresponding univariate notions. So far, we have recalled that univariate stochastic orderings can be used to define and analyze notions of univariate aging. Both univariate and multivariate orderings also provide the language for defining several notions of positive dependence and multivariate aging.

Among several existing definitions of multivariate orderings, particularly relevant are the (multivariate) notions of *ordinary* stochastic ordering, *likelihood ratio* ordering, *hazard rate* ordering, between two n -dimensional random vectors $\mathbf{X}, \mathbf{Y}$. See again [4] for a complete review with definitions, relevant properties, and further details. We only mention here that, analogously to the univariate case, the chain of implications $\mathbf{X} \leq_{lr} \mathbf{Y} \implies \mathbf{X} \leq_{hr} \mathbf{Y} \implies \mathbf{X} \leq_{st} \mathbf{Y}$ holds. Let $\leq_*$ stand for any of the above three orderings $\leq_{st}, \leq_{hr}, \leq_{lr}$. It is to be noticed that, in the special case when both the vectors $\mathbf{X} \equiv (X_1, \dots, X_n)$ and $\mathbf{Y} \equiv (Y_1, \dots, Y_n)$ have stochastically independent components, it is $\mathbf{X} \leq_* \mathbf{Y}$ (in the multivariate sense) if and only if it is $X_i \leq_* Y_i, i = 1, \dots, n$, in the univariate sense.

Notions of Positive Dependence

Mathematically, a property of positive dependence relates to one multivariate probability distribution for a vector of random variables; from a statistical viewpoint, this is a property that describes a tendency of the different random variables to assume concordant values.

Historically, the first such notion introduced in the statistical literature, as it is very well known, is the one of *positive correlation* defined by the condition that the **covariance** $Cov(X, Y)$, between two random variables X, Y , is nonnegative. Many more notions of positive dependence have been introduced along the twentieth century, but the concept of covariance continued to play an important role, at the level of giving definitions or studying implications of such notions. For reliability uses, in particular, a very important issue in this direction is the property of *association*. A vector of lifetimes $\mathbf{T} = (T_1, \dots, T_n)$ is associated if the two random variables $\varphi(\mathbf{T}), \psi(\mathbf{T})$ are positively correlated, i.e., $Cov(\varphi(\mathbf{T}), \psi(\mathbf{T})) \geq 0$, for pairs of functions $\varphi, \psi : \mathbb{R}^n \rightarrow \mathbb{R}$, with φ, ψ increasing in each coordinate and such that the covariance $Cov(\varphi(\mathbf{T}), \psi(\mathbf{T}))$ is well defined. It is interesting to mention that, in a natural way, the concept of association also has been raised in the fields of statistical mechanics and interacting particle systems (see e.g., [27]). It is well known among the reliability community that this concept has a very important role in the analysis of **coherent systems** (see in particular [3, 28, 29]). This is because of the following structural properties:

1. Any subset of a set of associated random variables is associated.
2. The union of two stochastically independent subsets of associated random variables is associated.
3. Any single random variable is associated.
4. Random variables obtained as different increasing functions of a same set of associated random variables give rise to an associate set.

A further idea that has a pervasive, and unifying, role in defining and analyzing positive dependence is, once again, that of stochastic comparisons. This fact may manifest, however, under several different forms, some of which are listed henceforth.

Take $n = 2$ random variables X and Z and consider a univariate stochastic ordering $\leq_*$ (where, this time, the symbol $\leq_*$ may stand for $\leq_{st}, \leq_{hr}, \leq_{lr}$,

mentioned in the section titled “Some Basic Notions and Notation”, or for other univariate notions). Conditions of the type

$$\bar{F}_X(\cdot | Z = z') \leq_* \bar{F}_X(\cdot | Z = z''), \quad \text{for } z' < z'' \quad (16)$$

$$\bar{F}_X(\cdot | Z > z') \leq_* \bar{F}_X(\cdot | Z > z''), \quad \text{for } z' < z'' \quad (17)$$

or

$$\bar{F}_X(\cdot | Z \leq z') \leq_* \bar{F}_X(\cdot | Z \leq z''), \quad \text{for } z' < z'' \quad (18)$$

are called *conditions of stochastic monotonicity* of X with respect to Z , and actually define properties of positive dependence for the pair (X, Z) .

As instances of that, we can cite in particular the following two concepts.

X is *stochastically increasing in Z* if, for any x , the function

$$P(X > x | Z = z) \quad (19)$$

is nondecreasing as a function of the variable z . X is *left tail decreasing in Z* if, for any x , the function

$$P(X > x | Z \leq z) \quad (20)$$

is nondecreasing as a function of the variable z . In these two cases, then, $\leq_*$ stands for the ordinary stochastic ordering.

Possibly using multivariate stochastic orderings, conditions of stochastic monotonicity can emerge in the characterization of some notions of positive dependence, or at least in providing necessary or sufficient conditions for them, also for the case with $n > 2$.

Other forms to express positive dependence are based on the comparison, with respect to some multivariate ordering, of the actual joint distribution with the distribution characterized by the same univariate marginals and stochastic independence among the variables. A very simple, and illuminating, example of this is provided by the concept of *positive quadrant dependence* (PQD) for a pair of lifetimes X, Y : this is expressed by the condition

$$\begin{aligned} \bar{F}(x, y) &= P\{X > x, Y > y\} \geq P\{X > x\} \cdot P\{Y > y\} \\ &= \bar{G}_X(x) \cdot \bar{G}_Y(y) \end{aligned} \quad (21)$$

i.e., for any choice of $x > 0, y > 0$, the two events $\{X > x\}, \{Y > y\}$, are positively dependent; it is then clear that this condition is implied by association of (X, Y) . In the multivariate case with $n > 2$, the notion of PQD can be extended to *positive upper orthant dependence* (PUOD) that, again, is a condition implied by association.

Defining concepts of dependence this way is strictly related to a more general fact: in several cases, a positive dependence property can be characterized in terms of the *connecting copula* or of the *survival copula* associated to the joint distribution at hand. From an analytical point of view, copulas are those functions $C : [0, 1] \times \dots \times [0, 1] \rightarrow [0, 1]$ that can be obtained as (restrictions to $[0, 1] \times \dots \times [0, 1]$ of) joint probability distribution functions of vectors of random variables $U_1, \dots, U_n$, having uniform marginal distributions.

Let $T_1, \dots, T_n$ be lifetimes with continuous marginal distribution functions $G_1, \dots, G_n$, marginal survival functions $\bar{G}_1, \dots, \bar{G}_n$, joint distribution function F and joint survival function $\bar{F}$. The connecting copula of the vector $\mathbf{T} = (T_1, \dots, T_n)$ is the copula C defined by the relation

$$F(x_1, \dots, x_n) = C(G_1(x_1), \dots, G_n(x_n)) \quad (22)$$

whereas the survival copula is the copula $\hat{C}$ defined by the relation

$$\bar{F}(x_1, \dots, x_n) = \hat{C}(\bar{G}_1(x_1), \dots, \bar{G}_n(x_n)) \quad (23)$$

As a simple instance of characterization of a positive dependence property in terms of copulas, it is immediate to check that PQD for a pair of lifetimes X, Y is equivalent to the condition $\hat{C}(u, v) \geq u \cdot v$ for their (bivariate) survival copula $\hat{C}$.

Many other conditions of positive dependence can be easily characterized in terms of the connecting copula or of the survival copula. For general reference books and other basic material about these issues see, in particular, [30–32] and references therein.

A different, formally very elegant, way to employ multivariate stochastic ordering concepts for the generation of positive dependence notions is as follows: it can happen that a multivariate stochastic ordering $\leq_*$ is not *reflexive* (for instance $\leq_{hr}$ and $\leq_{lr}$ are not: for a vector of lifetimes $\mathbf{T}$, it is not generally true that $\mathbf{T} \leq_{lr} \mathbf{T}$ nor that $\mathbf{T} \leq_{hr} \mathbf{T}$). For such a stochastic ordering, we can obtain a notion of

positive dependence by imposing that a vector $\mathbf{T}$ is such that $\mathbf{T} \leq_* \mathbf{T}$. In particular we say that $\mathbf{T}$ is *multivariate totally positive of order 2* (MTP_2) if $\mathbf{T} \leq_{lr} \mathbf{T}$ (see [33]) and that it is *hazard rate increasing upon failure* (HIF) if $\mathbf{T} \leq_{hr} \mathbf{T}$ (see [34]). The former notion implies the latter one; both of them imply association. Formally, the notion of MTP_2 has properties that mirror those of association mentioned above, even if it gives rise to a much more restrictive condition; in fact it is really a very strong notion, but also a very useful one, from several respects. For some uses of MTP_2 in Bayesian Statistics see Chapter 3 in [35] and some references therein.

In modern reliability theory it is often natural to look at different problems and models from the viewpoint of a *dynamic approach*. This amounts to take into consideration the evolution of the joint distribution of the residual lifetimes of surviving individuals, when time elapses and subsequent failures are progressively observed. Such an approach reveals to be necessary, for instance, in the following type of situations: we have to describe the probabilistic model for observable lifetimes $T_1, \dots, T_n$ of units that are subjected to similar environmental conditions, and such conditions change stochastically in time, giving rise to a nonobservable stochastic process.

Some dependence notions of *dynamic type* have been introduced in the literature to deal with such cases; often in these cases, the use of the concept of m.c.h.r. functions may emerge as natural.

A typical dynamic notion of dependence is the one of *weakened by failure*, introduced by Arjas and Norros in [36, 37]. This is a notion of positive dependence that implies association; it formalizes situations where, upon failure of a unit, the vector of residual lifetimes of the surviving components undergoes a decrease in the multivariate ordinary stochastic order sense.

The structure of this definition shows once again the possible role of the different notions of multivariate stochastic ordering in the description of conditions of positive dependence: in the above definition, the ordinary (multivariate) ordering can be replaced with other types of (multivariate) orderings to get different notions of positive dependence. In this respect, it is worthwhile pointing out that Shaked and Shanthikumar reobtained in this way some of the dependence properties mentioned above: when, in the above definition, the ordinary ordering is replaced

with the hazard rate ordering, in place of weakened by failure WBF we obtain HIF; when the ordinary ordering is replaced with the likelihood ratio ordering, in place of WBF we obtain MTP₂ (see [34]). This in particular shows that MTP₂ implies HIF and, in turn, HIF implies WBF; we then have the following chain of implications:

$$\begin{aligned} \text{MTP}_2 &\implies \text{HIF} \implies \text{WBF} \implies \text{Association} \\ &\implies \text{PUOD} \end{aligned} \quad (24)$$

Another notion (*Supportive Lifetimes*) can be obtained, similarly to above, in terms of the concept of (multivariate) cumulative hazard rate ordering.

The different notions of multivariate stochastic orderings and of positive dependence can take a specially expressive form, when the vector of lifetimes of concern is *exchangeable*, i.e., when the joint distribution of $T_1, \dots, T_n$ is the same of that of $T_{j_1}, \dots, T_{j_n}$ for any permutation $\{j_1, \dots, j_n\}$ of $\{1, \dots, n\}$; in particular, if $T_1, \dots, T_n$ are exchangeable, then they are identically distributed (not necessarily independent). De Finetti's theorem about exchangeable random variables says the following: if $T_1, \dots, T_n$ are exchangeable and constitute the initial segment of an exchangeable sequence $T_1, \dots, T_n, T_{n+1}, \dots$ then it is possible to construct a random variable Θ , such that $T_1, \dots, T_n$ are conditionally independent and identically distributed i.i.d., given Θ . It is easily seen that, in such a case, $T_1, \dots, T_n$ are positively correlated. Starting from this fact there has been, since a long time, a special interest in the analysis of positive dependence conditions for the case of exchangeability. In particular, e.g., characterizations of association for exchangeable variables were studied in [38]. A basic fact, concerning positive dependence and exchangeability is the study of the notion of *PDM*; a pair of random variables X, Y is PDM if they are conditionally i.i.d., given a random variable Θ (see [39, 40] and references therein). Several examples concerning aforementioned notions of multivariate orderings and positive dependence for the special case of exchangeability are discussed in [35].

Also, besides situations of exchangeability a central fact is that, more generally, positive dependence can be created by situations of conditional independence upon some nonobservable factors. In a few words, we can highlight the following: take lifetimes $T_1, \dots, T_n$ that are conditionally independent given a nonobservable scalar factor Θ ; for $j = 1, \dots, n$,

let $\bar{F}_j(\cdot|\theta)$, the conditional survival function of T_j given $\{\Theta = \theta\}$, be e.g., stochastically decreasing in θ , with respect to some fixed univariate stochastic order, then $\mathbf{T} = (T_1, \dots, T_n)$ has some property of positive dependence. A simple heuristic explanation goes as follows: observing high values for a subset of variables $\mathbf{T}_I$ induces to guess that Θ takes a low value; in turn, this induces to predict that the values taken by the remaining variables $\mathbf{T}_{\bar{I}}$ are high. In that case when nonobservable factors constitute a vector Θ , rather than being summarized by a scalar Θ , Θ must have some suitable property of positive dependence, to reach the same conclusion about positive dependence of $\mathbf{T}$.

All these are just general, roughly expressed facts; they are, however, related to fundamental aspects in the foundations of subjective probability and Bayesian statistics, as in particular very early discussed by Bruno de Finetti (see e.g., [41]).

The above heuristic argument can also be seen as the starting point to obtain several rigorous results. A basic step in such a direction is that, under suitable assumptions of stochastic monotonicity of $\mathbf{T}_I$ with respect to Θ , and for suitably ordered pairs $\mathbf{t}'_I, \mathbf{t}''_I$ the conditional distributions of Θ given the two different results $\{\mathbf{T}_I = \mathbf{t}'_I\}, \{\mathbf{T}_I = \mathbf{t}''_I\}$, are stochastic ordered (see in particular [42]).

Precise statements and specific definitions about positive dependence induced by conditional independence can be found in [43–47] and are also briefly recalled in the article **Mixture Models, Multivariate: Aging and Dependence**.

For other aspects of general interest related to positive dependence, see e.g., [43] and the volumes [48, 49]. For a very recent literature about new multivariate stochastic orderings and notions of positive dependence originated from them, see e.g., [50, 51], and references therein.

So far a few methods have been listed to “define” properties of positive dependence; one can wonder now: What is a property of positive dependence? Does there exist any axiomatic definition of what a notion of positive dependence is, in wider generality? A response to these questions has been given by Kimeldorf and Sampson in [52] for the case $n = 2$. They listed a set of axioms that a family of probability distributions should satisfy to give rise to a certain condition of positive dependence. Many important notions of dependence in fact do satisfy such axioms. In particular, related to something we

mentioned before, one of these axioms concerns the comparison of a bivariate distribution with the distribution corresponding to the case characterized by the same pair of marginals and stochastic independence. Different methods to extend that to the case $n > 2$, have been aimed more recently.

Once definitions of positive dependence have been defined and analyzed, the question arises, to understand, which are the multivariate probability models in the applications of interest that may manifest positive dependence. There is of course a very rich and interesting literature in this direction. Here, we will just limit ourselves to a brief discussion in the last part of the next section.

Multivariate Aging and Relations between Aging and Positive Dependence

Notions of *multivariate aging* concern vectors of lifetimes, like it does for notions of dependence; however, multivariate aging and dependence are different types of concepts. To understand the property of multivariate aging, we start by considering a vector of independent lifetimes. Fix first a univariate aging property $\sharp$ ($\sharp$ can stand e.g., for IFR, NBU, DFR, ...) and take n independent lifetimes $T_1, \dots, T_n$, with marginal survival functions $F_1, \dots, F_n$, respectively, so that

$$\bar{F}(t_1, \dots, t_n) = \bar{F}_1(t_1) \dots \bar{F}_n(t_n) \quad (25)$$

and let $\bar{F}_1, \dots, \bar{F}_n$ share the property $\sharp$. After that, it is analyzed which type of multivariate inequality is common to all the joint survival functions $\bar{F}$ that can be obtained in this way. We then pass to joint survival functions $\bar{F}$ of nonindependent lifetimes. We can finally assert that $\bar{F}$ satisfies a multivariate $\sharp$ property if it satisfies the previously considered multivariate inequality.

For instance, as a direct extension to the cases of dependence of a property valid for independent IFR lifetimes, the definition of multivariate increasing failure rate (MIFR) considered by Shaked and Shantikumar in [53, 54], is given as follows.

Compare two different observations

$$\{\mathbf{T}_I = \mathbf{t}_I; \mathbf{T}_{\bar{I}} > t\mathbf{e}\}, \{\mathbf{T}_I = \mathbf{t}'_I, \mathbf{T}_J = \mathbf{t}'_J; \mathbf{T}_{\bar{I} \cap \bar{J}} > t'\mathbf{e}\} \quad (26)$$

with $\mathbf{t}'_I \leq \mathbf{t}_I \leq t\mathbf{e}$, $t\mathbf{e} \leq \mathbf{t}'_J \leq t'\mathbf{e}$ and consider the vectors of the residual lifetimes $\mathbf{T}_{\bar{I} \cap \bar{J}} - t$, $\mathbf{T}_{\bar{I} \cap \bar{J}} - t'$.

Then $\mathbf{T}$ is *MIFR* if the joint distribution of $\mathbf{T}_{\bar{I} \cap \bar{J}} - t\mathbf{e}$, given the observation $\{\mathbf{T}_I = \mathbf{t}_I; \mathbf{T}_{\bar{I}} > t\mathbf{e}\}$ is stochastically larger than the joint distribution of the vector $\mathbf{T}_{\bar{I} \cap \bar{J}} - t'\mathbf{e}$, given the observation $\{\mathbf{T}_I = \mathbf{t}'_I, \mathbf{T}_J = \mathbf{t}'_J; \mathbf{T}_{\bar{I}} > t'\mathbf{e}\}$.

It is to be said that the general strategy to define properties of multivariate aging as described above, does not lead, of course, to a uniquely determined solution; in fact, still for the same notion $\sharp$, there are different possible inequalities that can be extended from the case of independence to that of stochastic dependence. As a matter of fact, for instance, several different concepts of MIFR have been presented and analyzed in the literature; see in particular [16, 55–58].

Further notions of MIFR have been considered in [59–62] in the frame of a Bayesian approach to reliability and for the special case of exchangeable lifetimes; such notions are related to the concept of Schur concavity (see [63]). A univariate survival function is IFR if and only if it is log concave; relying on this fact, it can be easily checked, in particular, that the joint survival function $\bar{F}$ of independent identically IFR distributed lifetimes is Schur concave. It was highlighted by Richard E. Barlow that this property is also maintained under mixtures: more precisely, the joint survival function of exchangeable lifetimes that are conditionally i.i.d., with respect to some nonobservable factors, is still Schur concave. This remark suggested to define a multivariate notion of IFR for exchangeable lifetimes by requiring the Schur concavity property for their joint survival function.

Let $T_1, \dots, T_n$ be the lifetimes of n different, but similar, components $D_1, \dots, D_n$ and let $T_1, \dots, T_n$ be exchangeable; for a vector of observed ages $t_1, t_2, s_3, \dots, s_n$, consider the conditional probabilities of extra survival

$$P\{T_1 > t_1 + \tau | \mathbb{D}\}, P\{T_2 > t_2 + \tau | \mathbb{D}\} \quad (27)$$

where

$$\mathbb{D} = \{T_1 > t_1, T_2 > t_2, T_3 > s_3, \dots, T_n > s_n\} \quad (28)$$

It is obvious that, for $t_1 < t_2$, it is

$$P\{T_1 > t_1 + \tau | \mathbb{D}\} \geq P\{T_2 > t_2 + \tau | \mathbb{D}\} \quad (29)$$

whenever $T_1, \dots, T_n$ are independent and obey the same IFR distribution; this amounts to say that, whatever survival data has been observed for $D_3, \dots, D_n$,

the elder, between the two components D_1, D_2 , is the one that has the smaller probability of surviving for an extra time τ . It is easy to show (see [61]) that, for exchangeable lifetimes more generally, such inequality remains equivalent to Schur concavity of the joint survival function, as in the case of i.i.d. lifetimes. Schur concavity of the joint survival function is then called a Bayesian notion of IFR (see also the discussion in [35]).

A main difference between the latter and the MIFR property of [53] stands in the fact that MIFR is defined in terms of a comparison between the residual lifetimes, for the same set of surviving components, at two different ages of them; in the inequality above, on the contrary, the residual lifetimes of two different (but exchangeable) components of different ages are compared at the same calendar time. In both the cases, two conditional distributions are compared; but, in the latter case, conditioning is performed with respect to a similar observed event; this allows one to discern the effects of aging and the effects of stochastic dependence. A further source of difference between the two types of definition is based on the fact that, in the definition of MIFR, the comparison is described in terms of a multivariate ordering.

As mentioned in the section titled “Some Basic Notions and Notations”, several natural notions of univariate aging have been defined; for some of those, corresponding multivariate extensions can be singled out by following, for instance, the logic that lead to the notion of MIFR and by using different types of multivariate stochastic orderings.

For exchangeable lifetimes (see [64]), many other notions of multivariate aging can be given in terms of univariate aging as follows. For a given univariate stochastic ordering $*$, we can define different notions of multivariate aging by comparing, in the $*$ sense, distributions of residual lifetimes $T_1 - t_1$ and $T_2 - t_2$, conditionally on different possible types of observed events; see [35] for a discussion about this approach; for a further review and further aspects of the Bayesian approach to Reliability, see also [8, 65].

A different logic that has been sometimes adopted in the definition of multivariate notions of aging is as follows: it is required that a distribution of joint lifetimes has a multivariate aging property $\sharp$ only if their univariate marginals have the corresponding (univariate) property $\sharp$ (see e.g., [66]).

In some cases, however, this requirement is not necessary; in particular when, due to the presence of nonobservable factors, dependence among observable lifetimes is explained by situations of learning about such factors; see also the discussions in [35, 62] (see **Mixture Models, Multivariate: Aging and Dependence**).

Before continuing the discussion, an important aspect that generally apply in the definition of any multivariate aging notion needs to be stressed. In any case, a multivariate aging notion $m\text{-}\sharp$ is seen as a multivariate extension of a corresponding univariate notion $\sharp$. For vectors of stochastically independent lifetimes $T_1, \dots, T_n$, the property $m\text{-}\sharp$ holds if and only if each $T_j (1 \leq j \leq n)$ satisfies the property $\sharp$ (e.g., see [54] for what specifically concerns the property of MIFR). In a few words, notions of multivariate aging for the joint distribution of several lifetimes substantially reduce to univariate **aging notions** for the marginals, when the lifetimes are independent.

We can then say that, just by definition, genuine conditions of multivariate aging manifest only in the cases of stochastic dependence. Thus multivariate aging and dependence are just two different, but concomitant, aspects of multivariate survival distributions.

On the other hand, also univariate aging of components is easily seen to be triggered by dependence in a natural way. Let $T_1, \dots, T_n$ be the lifetimes of dependent components $D_1, \dots, D_n$ and let us admit that, initially, $T_1, \dots, T_n$ have exponential (i.e., nonaging) marginal distributions. Consider then, at the calendar time $t > 0$, the conditional probability distribution of the residual lifetime $T_j - t$, given an observed history $\{\mathbf{T}_t = \mathbf{t}_t; \mathbf{T}_{\tilde{t}} > t\mathbf{e}\}$, where $j \in \tilde{t}$, i.e., j is an index such that D_j survived up to t . Generally, such conditional distribution is not exponential, anymore.

In applied models, in fact, many situations that give rise to positive dependence have been studied in detail to show how they also manifest (positive or negative) univariate aging. In turn, positive univariate aging combined with positive dependence is expected to create conditions of positive multivariate aging.

It is then an interesting and important issue, in quality and reliability, to understand possible substantial interactions existing between the phenomena of positive dependence, univariate aging, and multivariate aging. Such interactions are, however,

rather convoluted and general results are missing; this constitutes then a still open field of research. A major difficulty, of course, is associated to the fact that different approaches have been followed in defining multivariate aging and these points are further discussed briefly below.

Previously, we prefer to dwell on relations and analogies, between aging and dependence, that can be established at a conceptual or just at a formal level.

First, a conceptual aspect of fundamental importance that is common to aging and dependence is focused on. As well known, and was also recalled above, both dependence and aging properties may depend on the state of information we have on a vector of lifetimes; see articles [16, 18, 21, 22] for a general and rigorous treatment. As a very simple example we can consider here the case of the lifetimes $T_1, \dots, T_n$ of n similar components $D_1, \dots, D_n$ that work independently of each other and share a common “quality” Θ ; we suppose then that, when Θ is known to take a value θ , $T_1, \dots, T_n$ are i.i.d. with the same survival function $\bar{G}(\cdot|\theta)$; suppose also that $D_1, \dots, D_n$ are to work in a situation of wearout so that $\bar{G}(\cdot|\theta)$ is assumed to be IFR, for any possible value θ . Whenever we have a state of uncertainty about Θ and we look at Θ as a random variable with a probability density π of its own, then, in the unconditional distribution, $T_1, \dots, T_n$ are conditionally independent and, in particular, positively dependent by mixture; $T_1, \dots, T_n$ are not anymore independent, however. Simultaneously, also the IFR property of their marginal can be lost under mixture (see **Mixture Models, Multivariate: Aging and Dependence**). This example points out the following possibility: depending on the level of information, lifetimes $T_1, \dots, T_n$ can be independent, identically IFR distributed or can, alternatively, manifest positive dependence and some aspect of negative aging, in the marginal distribution.

Continuing the discussion about analogies between aging and dependence, we also must appreciate the role of concepts of stochastic ordering in the definitions of positive dependence, univariate aging, and multivariate aging. Univariate stochastic orders are used to define both univariate and (some) multivariate notions of aging. Multivariate stochastic orders can be used to define both notions of positive dependence and some of the notions of multivariate aging.

A concept that can also have formally a unifying role, at least for exchangeable lifetimes, is the one of *semicopula*. Let us limit ourselves to the case of $n = 2$ dimensions. As mentioned, several notions of dependence, for a pair of lifetimes, can be characterized in terms of their survival copula $\hat{C}$. Copulas are special cases of semicopulas. A bivariate copula, as said, is just a bivariate distribution function C with support on $[0, 1] \times [0, 1]$ and with uniform marginals; being a distribution function, C is increasing in each of the two variables and satisfies the rectangular inequality (see e.g., [31]). A semicopula is a function that satisfies all the properties of a copula except, possibly, the rectangular inequality; formally, we can extend the notions of dependence also to the field of semicopulas.

We can associate a semicopula A to a one-dimensional survival function $\bar{G}$, by formally defining, see [67],

$$A(u, v) = \bar{G} \left[\bar{G}^{-1}(u) + \bar{G}^{-1}(v) \right], \\ 0 \leq u \leq 1, 0 \leq v \leq 1 \quad (30)$$

It was an idea pointed out by Averous and Dortet-Bernadette, see [68], that (univariate) aging properties of $\bar{G}$ reflect into dependence properties of A (in [68], in particular, only cases in which A is a copula were considered).

Then univariate aging can be formally expressed as a property of dependence.

On the other hand, to an exchangeable bivariate survival function $\bar{F}$, with everywhere positive and strictly decreasing marginal survival function $\bar{G}$, one can associate a semicopula $B : [0, 1] \times [0, 1] \rightarrow [0, 1]$ defined as follows:

$$B(u, v) = \exp \left\{ -\bar{G}^{-1} \left(\bar{F}(-\log u, -\log v) \right) \right\} \quad (31)$$

so that

$$\bar{F}(x, y) = \bar{G} \left(-\log B(e^{-x}, e^{-y}) \right) \quad (32)$$

The function B is of interest in that it provides a tool to describe the family of the level curves of $\bar{F}$ by means of a semicopula, and for this reason it was introduced in [69]. It was also noticed in [69] that some of the bivariate aging properties introduced in the Bayesian approach are just properties of the

family of the level curves of $\overline{F}$ and can be translated into dependence properties of B .

For the same bivariate model $\overline{F}$, in this way, we have the following statements:

1. Some of the possible properties of dependence can be described in terms of dependence properties of the survival copula $\widehat{C}$.
2. Some of the possible properties of univariate aging can be described in terms of dependence properties of the semicopula A .
3. Some of the possible properties of bivariate aging can be described in terms of dependence properties of the semicopula B .

These facts allows one to give an axiomatic definition of “correspondence” among notions of dependence, univariate aging, and bivariate aging.

Let us, for instance, look at the condition of PQD that, as we saw in the section titled “Notions of Positive Dependence”, has the meaning of a property of positive dependence when imposed on the survival copula $\widehat{C}$. The same condition corresponds to the univariate aging property of new better than used NWU when imposed on the semicopula A ; it corresponds to a suitable bivariate aging property of bivariate NBU when imposed on the semicopula B .

Such a correspondence puts us in a position to analyze some interactions among properties of positive dependence, univariate aging, and bivariate aging. In [67], it was, e.g., shown that the combination of positive dependence and positive univariate aging imply positive bivariate aging, whereas the combination of positive dependence and negative bivariate aging imply negative univariate aging.

These relations can also be of interest in the analysis of the evolution of aging and dependence properties for residual lifetimes of surviving components, at the increase of their age. A sketch of this approach is presented in [70]; a more complete development could be the object of further research.

Obviously, these results are not generally true for other notions of multivariate aging; their interest is limited to the cases where the “Bayesian” notions can be applied and essentially they have been developed, so far, only for the cases of exchangeable lifetimes.

Other relations between positive dependence, multivariate and univariate aging have also been pointed out in the literature.

As mentioned, notions of multivariate aging defined according to the line considered e.g., in the paper [66] require corresponding properties of univariate aging for the marginal distributions, irrespective of dependence conditions.

Dynamic notions of positive multivariate aging considered in [17, 54] imply corresponding properties of positive dependence. In passing, it is to be noticed that these conditions of multivariate aging can be rather restrictive. This can happen, e.g., in the case of conditional independence; in such a case, in fact, positive dependence can likely manifest along with some aspect of negative univariate aging for the marginal distributions (*see also the discussion in Mixture Models, Multivariate: Aging and Dependence*).

Several papers in the current literature were devoted to analyzing relationships among dependence, univariate, and multivariate aging, under the specification of particular survival models, rather than from a conceptual and general point of view. For lack of room, we cannot review these types of analyses, here. In this respect, we just recall some general classes of models, that emerge for relevance of their applied interest, and that may simultaneously manifest phenomena of positive dependence and of aging.

Among models of interest, we find the multivariate mixture models, characterized by an assumption of conditional independence for the lifetimes $T_1, \dots, T_n$, given a random vector Θ of unobservable factors (*see Mixture Models, Multivariate: Aging and Dependence*). Owing to such an assumption, positive dependence is likely to be present along with some form of univariate aging.

Such models, where Θ is constant in time, can actually be looked at as very special cases of the more general situation where Θ is replaced by a stochastic process Θ_t and the observable lifetimes $T_1, \dots, T_n$ are conditionally independent given the history of Θ_t . This situation also appear in a great variety of cases of interest in the applications.

Shock models constitute relevant cases of the above situations; in shock models the process Θ_t is typically a (possibly marked) point process, that accounts for the arrival times and severity of the shocks. In **change-point models**, Θ_t can still be seen as a counting process, with only one jump corresponding to the instant of change point. Change-point models, however, are not generally seen as

special cases of shock models, in fact the impact of the change point on the failure rate of the surviving components can differ from that triggered by shocks.

All these models can be seen as special cases of the more general situation where the $T_1, \dots, T_n$ are the lifetimes of different components $D_1, \dots, D_n$ that work simultaneously and are influenced by a common environment acting on all the units; here we are thinking of situations of stochastically varying environments. It has been recognized since a long time that a similar varying environment acting on the different components induces some type of correlation among the components' lifetimes. The analysis of this type of correlation has then a long and still continuing history; some basic references are [71–75]. For more recent references, see e.g., [76, 77].

On the other hand, the interplay between aging and dependence (see e.g., [78]) and the analysis of the impact of random environmental situations and shocks on aging properties have also a long history; see for example [79–81] where it is shown how situations of PDM aging can be created simultaneously. Again, for the interplay between aging and stress shocks see [82] and some references therein.

Further interesting multivariate models that explain the simultaneous creation of aging and positive dependence phenomena are the so-called load-sharing models (see e.g., [54, 83]). In these cases the failure rates depend on the set of the surviving components. These are not really situations of varying environmental conditions; still a dynamic approach is appropriate and the probabilistic model for the vector of lifetimes is conveniently described in terms of the multivariate conditional hazard rates.

Realizing that, for multivariate survival models, there are strong interactions between aging and dependence is a very important issue when dealing with real applied problems. Situations of positive aging, for instance, may suggest to adopt procedures of preventive replacement on the different units; on the other hand, the presence of strong positive dependence among units suggests that the replacement policy cannot be decided at time 0; in fact it should be sequential and the time to the next replacement is to be reconsidered after each observed failure. In this type of problems, it is also important to remind that the properties of aging and those of dependence may be modified by a change in the information level: for the same set of components, for instance, the structure of the optimal replacement policy depends on

the state of information and on what we are able to observe; even the decision, whether it is convenient to adopt a replacement procedure or not may depend on the state of current information (see e.g., [84]).

Similar arguments are also to be taken into account when, in the presence of negative aging, we plan some politics of burn-in.

We cannot conclude this article without mentioning that, especially in the last few years, there has been an explosion of papers and books on topics related to positive dependence in the fields of insurance and mathematical finance (see e.g., [85–87]). Several models and properties studied in these frames, in particular in relation to the analysis of *interacting defaults*, are rather strictly related to the literature developed in the fields of quality and reliability.

References

- [1] Lehmann, E.L. (1966). Some concepts of dependence, *Annals of Mathematical Statistics* **37**, 1137–1153.
- [2] Barlow, R.E. & Proschan, F. (1965). *Mathematical Theory of Reliability*, SIAM, New York.
- [3] Barlow, R.E. & Proschan, F. (1975). *Statistical Theory of Reliability and Life Testing*, Holt, New York.
- [4] Shaked, M. & Shanthikumar, J.G. (1994). *Stochastic Orders and Their Applications*, Academic Press, London.
- [5] Keilson, J. & Sumita, U. (1982). Uniform stochastic orderings and related inequalities, *Canadian Journal of Statistics* **10**, 181–189.
- [6] Mosler, K. & Scarsini, M. (eds) (1991). *Stochastic Orders and Decisions Under Risk*, Institute of Mathematical Statistics, Hayward.
- [7] Szekli, R. (1995). *Stochastic Ordering and Dependence in Applied Probability, Lecture Notes in Mathematics* 97, Springer, New York.
- [8] Lai, C.-D. & Xie, M. (2006). *Stochastic Ageing and Dependence for Reliability*, Springer-Verlag, New York.
- [9] Bergman, B. (1985). On reliability and its applications, *Scandinavian Journal of Statistics* **12**, 1–41.
- [10] Zacks, S. (1992). *Introduction to Reliability Analysis*, Springer-Verlag, New York.
- [11] Block, H.W. & Savits, T.H. (1997). Burn-in, *Statistical Science* **12**, 1–19.
- [12] Aven, T. & Jensen, U. (1999). *Stochastic Models in Reliability, Applications of Mathematics*, 41, Springer-Verlag, New York.
- [13] Gertsbakh, I. (2000). *Reliability theory. With Applications to Preventive Maintenance*, Springer-Verlag, Berlin.
- [14] Cox, D. & Isham, V. (1980). *Point Processes*, Chapman & Hall, London.

- [15] Bremaud, P. (1981). *Point Processes and Queues. Martingale Dynamics*, Springer-Verlag, New York.
- [16] Arjas, E. (1981). A stochastic process approach to multivariate reliability systems: notions based on conditional stochastic order, *Mathematics of Operations Research* **6**, 263–276.
- [17] Arjas, E. (1981). The failure and hazard processes in multivariate reliability systems, *Mathematics of Operations Research* **6**, 551–562.
- [18] Norros, I. (1986). A compensator representation of multivariate life length distributions, with applications, *Scandinavian Journal of Statistics* **13**, 99–112.
- [19] Koch, G. (1986). *A Dynamical Approach to Reliability Theory: Theory of Reliability*, North-Holland, Amsterdam, pp. 215–240.
- [20] Yashin, A.I. & Arjas, E. (1988). A note on random intensities and conditional survival functions, *Journal of Applied Probability* **25**, 630–635.
- [21] Arjas, E. (1989). Survival models and martingale dynamics, *Scandinavian Journal of Statistics* **16**, 177–225.
- [22] Arjas, E. & Norros, I. (1991). Stochastic order and martingale dynamics in multivariate life length models: a review, in *Stochastic Orders and Decision Under Risk*, K. Mosler & M. Scarsini, eds, Institute of Mathematical Statistics, Hayward, pp. 7–24.
- [23] Gaver, D.P. (1963). Random hazard in reliability problems, *Technometrics* **5**, 211–226.
- [24] Johnson, N.L. & Kotz, S. (1975). A vector multivariate hazard rate, *Journal of Multivariate Analysis* **5**, 53–66.
- [25] Shaked, M. & Shantikumar, J.G. (1987). Multivariate hazard rates and stochastic ordering, *Advances in Applied Probability* **19**, 123–137.
- [26] Shaked, M. & Shantikumar, J.G. (1987). The multivariate hazard construction, *Stochastic Processes and their Applications* **24**, 85–97.
- [27] Liggett, T. (1985). *Interacting Particle Systems*, Springer-Verlag, New York.
- [28] Esary, J.D., Proschan, F. & Walkup, D.W. (1967). Association of random variables with applications, *Annals of Mathematical Statistics* **38**, 1466–1474.
- [29] Esary, J.D. & Proschan, F. (1970). A reliability bound for systems of maintained, interdependent components, *Journal of the American Statistical Association* **65**, 329–338.
- [30] Joe, H. (1997). *Multivariate Models and Dependence Concepts*, Chapman & Hall, London.
- [31] Nelsen, R.B. (2006). *An Introduction to Copulas*, 2nd Edition, Springer, New York.
- [32] Drouot Mari, D. & Kotz, S. (2001). *Correlation and Dependence*, Imperial College Press, London.
- [33] Karlin, S. & Rinott, Y. (1980). Classes of orderings measures and related correlation inequalities. I, multivariate totally positive distributions, *Journal of Multivariate Analysis* **10**, 467–498.
- [34] Shaked, M. & Shanthikumar, J.G. (1990). Multivariate stochastic orderings and positive dependence in reliability theory, *Mathematics of Operations Research* **15**, 545–552.
- [35] Spizzichino, F. (2001). *Subjective Probability Models for Lifetimes*, Chapman & Hall/CRC, Boca Raton.
- [36] Arjas, E. & Norros, I. (1984). Lifelengths and association: a dynamical approach, *Mathematics of Operations Research* **9**, 151–158.
- [37] Norros, I. (1985). Systems weakened by failures, *Stochastic Processes and their Applications* **20**, 181–196.
- [38] Burton, R.M. & Dabrowski, A.R. (1992). Positive dependence of exchangeable sequences, *Canadian Journal of Statistics* **44**, 774–783.
- [39] Shaked, M. (1977). A concept of positive dependence for exchangeable random variables, *Annals of Statistics* **5**, 505–515.
- [40] Shaked, M. (1979). Some concepts of positive dependence for bivariate interchangeable distributions, *Annals of the Institute of Statistical Mathematics* **31**, 67–84.
- [41] de Finetti, B. (1990). *Theory of Probability*, A critical introductory treatment. Translated from the Italian and with a preface by Antonio Machi and Adrian Smith. With a foreword by D.V. Lindley. John Wiley & Sons, Chichester, Vol. 2.
- [42] Fahmy, S., de, B., Pereira, C.A., Proschan, F. & Shaked, M. (1982). The influence of the sample on the posterior distribution, *Communications in Statistics A – Theory Methods* **11**, 1757–1768.
- [43] Shaked, M. (1982). A general theory of some positive dependence notions, *Journal of Multivariate Analysis* **12**, 199–218.
- [44] Jogdeo, K. (1978). On a probability bound of Marshall and Olkin, *Annals of Statistics* **6**, 232–234.
- [45] Shaked, M. & Spizzichino, F. (1998). Positive dependence properties of conditionally independent random lifetimes, *Mathematics of Operations Research* **23**, 944–959.
- [46] Lee, M.L.T. (1985). Dependence by total positivity, *Annals of Probability* **13**, 572–582.
- [47] Khaledi, B.-E. & Kochar, S. (2001). Dependence properties of multivariate mixture distributions and their applications, *Annals of the Institute of Statistical Mathematics* **53**, 620–630.
- [48] Block, H.W., Sampson, A. & Savits, T.H. (eds) (1991). *Topics in Statistical Dependence*, Institute of Mathematical Statistics, Hayward.
- [49] Tong, Y.L. (ed) (1984). *Inequalities in Statistics and Probability*, Institute of Mathematical Statistics, Hayward.
- [50] Müller, A. & Stoyan, D. (2002). *Comparison Methods for Stochastic Models and Risks*, John Wiley & Sons, Chichester.
- [51] Colangelo, A., Scarsini, M. & Shaked, M. (2005). Some notions of multivariate positive dependence, *Insurance: Mathematics and Economics* **37**, 13–26.

-
- [52] Kimeldorf, G. & Sampson, A.R. (1989). A framework for positive dependence, *Annals of the Institute of Statistical Mathematics* **41**, 31–45.
- [53] Shaked, M. & Shantikumar, J.G. (1988). Multivariate conditional hazard rates and the MIFRA and MIFR properties, *Journal of Applied Probability* **25**, 150–168.
- [54] Shaked, M. & Shantikumar, J.G. (1991). Dynamic multivariate aging notions in reliability theory, *Stochastic Processes and their Applications* **38**, 85–97.
- [55] Harris, R. (1970). A multivariate definition for increasing failure rate distribution functions, *Annals of Mathematical Statistics* **37**, 713–717.
- [56] Marshall, A.V. (1975). Multivariate distributions with monotone hazard rate, in *Reliability and Fault Tree Analysis*, SIAM, Philadelphia, pp. 259–284.
- [57] Block, H.W. & Savits, T.H. (1981). Multidimensional IFRA processes, *Annals of Probability* **9**, 162–166.
- [58] Savits, T.H. (1985). A multivariate IFR distribution, *Journal of Applied Probability* **22**, 197–204.
- [59] Barlow, R.E. & Mendel, M.B. (1992). de Finetti-type representations for life distributions, *Journal of the American Statistical Association* **87**, 1116–1122.
- [60] Barlow, R.E. & Mendel, M.B. (1993). Similarity as a characteristic of wear-out, in *Reliability and Decision Making*, R.E. Barlow, C.A. Clarotti & Spizzichino, F., eds, Chapman & Hall, London, pp. 233–245.
- [61] Spizzichino, F. (1992). Reliability decision problems under conditions of ageing, in *Bayesian Statistic*, J. Bernardo, J. Berger, A. P. Dawid & A.F.M. Smith, eds, Clarendon Press, Oxford, Vol. 4, pp. 803–811.
- [62] Barlow, R.E. & Spizzichino, F. (1993). Schur-concave survival functions and survival analysis, *Journal of Computational and Applied Mathematics* **46**, 437–447.
- [63] Marshall, A. & Olkin, I. (1979). *Inequalities: Theory of Majorization and its Applications*, Academic Press, New York.
- [64] Bassan, B. & Spizzichino, F. (1999). Stochastic comparison for residual lifetimes and Bayesian notions of multivariate ageing, *Advances in Applied Probability* **31**, 1078–1094.
- [65] Barlow, R.E. (1998). *Engineering Reliability*, SIAM, Philadelphia.
- [66] Buchanan, W.B. & Singpurwalla, N.D. (1977). Some stochastic characterizations of multivariate survival, in *The Theory and Applications of Reliability, with Emphasis on Bayesian and Nonparametric Methods*, Academic Press, New York, Vol. I, pp. 329–348.
- [67] Bassan, B. & Spizzichino, F. (2005). Relations among univariate aging, bivariate aging and dependence for exchangeable lifetimes, *Journal of Multivariate Analysis* **93**, 313–339.
- [68] Avérous, J. & Dortet-Bernadet, J. (2004). Dependence for archimedean copulas and aging properties of their generating functions, *Sankhya* **66**, 607–620.
- [69] Bassan, B. & Spizzichino, F. (2001). Dependence and multivariate aging: the role of level sets of the survival function, in *System and Bayesian Reliability*, World Scientific Publishing, River Edge, pp 229–242.
- [70] Bassan, B. & Spizzichino, F. (2003). On some properties of dependence and aging for residual lifetimes in the exchangeable case, in *Mathematical and Statistical Methods in Reliability (Trondheim, 2002)*, World Scientific Publishing, River Edge, pp. 235–249.
- [71] Lindley, D.V. & Singpurwalla, N.D. (1986). Multivariate distributions for the lifelengths of components of a system sharing a common environment, *Journal of Applied Probability* **23**, 418–431.
- [72] Cinlar, E. & Özekici, S. (1987). Reliability of complex devices in random environments, *Probability in the Engineering and Informational Sciences* **1**, 97–115.
- [73] Hougaard, P. (1987). Modelling multivariate survival, *Scandinavian Journal of Statistics* **14**, 291–304.
- [74] Cinlar, E., Shaked, M. & Shantikumar, G.J. (1989). On lifetimes influenced by a common environment, *Stochastic Processes and their Applications* **33**, 347–359.
- [75] Lee, M.L.T. & Gross, A.J. (1991). Lifetime distributions under unknown environment, *Journal of Statistical Planning and Inference* **29**, 137–143.
- [76] Belzunce, F., Lillo, R.E., Pellerrey, F. & Shaked, M. (2002). Preservation of association in multivariate shock and claim models, *Operations Research Letters* **30**(4), 223–230.
- [77] Singpurwalla, N.D., Mazzuchi, T.A., Özekici, S. & Soyer, R. (2003). Stochastic process models for reliability in dynamic environments, *Statistics in Industry, Handbook of Statist.*, North-Holland, Amsterdam, Vol. 22, pp. 1109–1129.
- [78] Brindley, E.C. & Thompson, W.A. (1972). Dependence and aging aspects of multivariate survival, *Journal of the American Statistical Association* **67**, 822–830.
- [79] Esary, J.D., Marshall, A.W. & Proschan, F. (1973). Shock models and wear process, *Annals of Probability* **1**, 627–649.
- [80] Block, H.W. & Savits, T.H. (1978). Shock models with NBUE survival, *Journal of Applied Probability* **15**, 621–628.
- [81] Cinlar, E. (1977). Shock and wear models and markov additive processes, *The Theory and Applications of Reliability, with Emphasis on Bayesian and Nonparametric Methods*, Academic Press, New York, Vol. I, pp. 193–214.
- [82] Fan, J., Ghurye, S.G. & Levine, R.A. (2000). Multicomponent lifetime distributions in the presence of ageing, *Journal of Applied Probability* **37**, 521–533.
- [83] Ross, S.M. (1984). A model in which components failure rates depend only on the working set, *Naval Research Logistics Quarterly* **31**, 297–300.
- [84] Heinrich, G. & Jensen, U. (1992). Optimal replacement based on different information levels, *Naval Research Logistics Quarterly* **39**, 937–955.
- [85] Asmussen, S. & Møller, J.R. (2003). Risk comparisons of premium rules: optimality and a life insurance study, *Insurance Mathematics and Economics* **32**, 331–344.
- [86] McNeil, A.J., Frey, R. & Embrechts, P. (2005). *Quantitative Risk Management. Concepts, Techniques and*

Tools. Princeton Series in Finance, Princeton University Press, Princeton.

- [87] Malevergne, Y. & Sornette, D. (2006). *Extreme Financial Risks. From Dependence to Risk Management*, Springer-Verlag, Berlin.

Further Reading

Whitt, W. (1982). Multivariate monotone likelihood ratio and uniform conditional stochastic order, *Journal of Applied Probability* **19**, 695–701.

Related Articles

Change-Point Models; Multivariate Stochastic Orders and Aging; Mixture Models, Multivariate: Aging and Dependence; Stochastic Orders and Aging.

FABIO SPIZZICHINO

Physical Degradation Models

Introduction

The focus within the reliability area has traditionally been on the collection and analysis of time-to-failure (TTF) data. High reliability implies few failures, so assessing and improving product/process reliability based on TTF data can be difficult. Accelerated life testing (ALT) can shorten the product life during test intervals by stressing the product beyond its normal use. However, ALT involves considerable amount of extrapolation. One should incorporate as much knowledge about the product as possible in designing the acceleration experiments.

Many physical systems degrade over time. Recent advances in sensing and measurement technologies are making it feasible to collect extensive amounts of data on physical **degradation** associated with components, systems, and manufacturing equipment. Davies [1] discusses a variety of engineering applications, types of degradation data, and recent developments in the area of condition monitoring and system maintenance. These cover data from vibration and acoustical monitoring, thermography, lubricant and wear debris analysis, and so on. Meeker and Escobar ([2], Chapters 14 and 21 and references therein) describe applications in luminosity of light bulbs, corrosion of batteries, semiconductor metal-oxide semiconductor (MOS) devices, and so on.

In practice, a number of subjects are under study. Degradation measurements of each subject are taken over time. The number of observation times and observation times themselves are allowed to vary across the subjects. These types of data are referred to as regular *degradation data*. For some applications, only one meaningful measurement can be taken on each test unit. The degradation measurement process destroys or changes the physical or mechanical characteristics of test units. This is known as *destructive degradation*. For example, the data from fatigue crack growth subject to loading cycles [3, 4] and bridge beams subject to erosion of chloride ion ingress [5] belong to regular degradation data. Data from the strength of adhesive bonds subject to environmental effects and internal erosion [6] are from destructive

degradation experiments since the test is destructive and strength can be measured only once on each unit.

When the amount of degradation reaches a pre-specified critical level D , failure occurs. Let T denote the TTF and $X(t)$ be the degradation process, then $T = \inf\{t : X(t) \geq D\}$. In many tests, the failure data is supplemented by degradation data. Degradation data are a very rich source of reliability information and offer many advantages over TTF data. The goal of this article is to present some common models for physical degradation. Specifically, we shall mainly discuss random-effect models and some Lévy process models. There is extensive literature on non-linear random-effect models (e.g., [7, 8]), and it is flexible and has been found quite useful in reliability applications. The Lévy processes such as Wiener process and gamma process have also been found useful as **degradation models**. These models lead to tractable forms for TTF distributions and thus we can make reliability inferences based on a combination of degradation and TTF data.

The paper is organized as follows. The section titled “Physical Degradation” gives some examples of physical degradation that provide useful motivation for degradation models. The section titled “Degradation Models” discusses two common physical degradation models and the reliability inferences drawn from them. In the section titled “An Illustrative Example”, we use a simulation example to illustrate that there are tremendous gains to be realized from degradation data. Finally we summarize the paper and describe some additional research topics.

Physical Degradation

In any given application, the problem context and appropriate subject matter knowledge must be used to model the underlying degradation mechanisms ([9] and [10] are good examples). For well-understood failure mechanisms, one may have a model based on physical/chemical theories that describe the underlying degradation process. Usually degradation models start with a deterministic description of the degradation process, often in the form of a differential equation. Then randomness can be incorporated into the model to describe the material parameters, component-to-component variability, initial conditions, and operating and environmental conditions.

2 Physical Degradation Models

Example 1: Automobile tire wear. If $\Lambda(t)$ is the amount of automobile tire tread wear at time t . This can be modeled simply by linear rate of degradation and wear rate is $d\Lambda(t)/dt = C$, then $\Lambda(t) = \Lambda(0) + Ct$, where $\Lambda(0)$ and C could be taken as constant for individual units but random from component-to-component.

Example 2: Fatigue crack growth. The prediction of crack growth accumulation is crucial for the reliability analysis of various materials in manufacturing. Let $\Lambda(t)$ be the crack size at time t and $\Delta K(\Lambda)$ be the stress intensity range, which is a function of the current crack size. A common used crack growth rate model is the Paris-rule model (e.g., [11]),

$$\frac{d\Lambda(t)}{dt} = C[\Delta K(\Lambda)]^m \quad (1)$$

where C and m are material properties. For example, when modeling a two-dimensional edge crack in a plate, $\Delta K(\Lambda) \propto \sqrt{\Lambda}$. Then we have

$$\Lambda(t) = \begin{cases} (C_0 + C_1 t)^{2/(2-m)} & m \neq 2 \\ C_3 e^{C_4 t} & m = 2 \end{cases} \quad (2)$$

where $C_i, i = 1, \dots, 4$ are constants which are functions of m . Note that different choices of m lead to different shapes of degradation paths.

Example 3: Current-gain degradation of bipolar transistors. In self-aligned polysilicon emitter transistors, a large electric field existing at the periphery of the emitter–base junction under reverse bias can create hot-carrier-induced degradation. Let Λ be the current gain, Burnett and Hu [12] studied a hot-electron degradation model,

$$\frac{d\Lambda(t)}{dt} = CG(\Lambda)I_R e^{-\phi/kT_e} \quad (3)$$

where ϕ is the critical energy needed for device damage, T_e is electron temperature, I_R represents the density of electrons, and $G(\Lambda)$ represents the dependence of degradation rate on existing damage. Assuming $G(\Lambda) \propto \Lambda^{-h}$, we have

$$\Lambda(t) \propto [t I_R e^{-\phi/kT_e}]^{1/(h+1)} \quad (4)$$

This relationship can be used to predict MOS device degradation dynamics.

Example 4: Photodegradation of polymeric materials. Excessive exposure of polymeric materials

leads to photodegradation failure. Let Λ be the total number of absorbed photons that effectively contribute to the photodegradation of a material during an exposure period. It has the form [13]

$$\Lambda(t) = \int_0^t \int_{\lambda_0}^{\lambda_1} E_0(\lambda, s)(1 - 10^{-A(\lambda, s)})\phi(\lambda) d\lambda ds \quad (5)$$

where λ_0 and λ_1 are the minimum and maximum photolytically effective UV–visible wavelengths, respectively, $A(\lambda, s)$ is the absorbance of sample at specified UV–visible wavelength and at time s , $E(\lambda, s)$ is the incident spectral UV–visible radiation dose to which a polymeric material is exposed at time s , and $\phi(\lambda)$ is the damage at wavelength λ relative to a reference wavelength. The overall photolytic damage to a polymeric material can be a linear, power, or exponential response of Λ .

Degradation Models

There are many factors that cause variability in the degradation curves and in the failure times. There are several approaches to modeling random degradation process in the literature. These include random-effect degradation models [4, 9, 10], stochastic process models such as Wiener process models [14–17], gamma process models [18–20], and compound **Poisson process** models [21], random shock process [22, 23], cumulative **damage models** [24], and multi-stage models [25]. For a general discussion of different degradation models, see Singpurwalla [26] and Meeker and Escobar [2]. We will mainly focus on the random-effect models and some nonhomogeneous Lévy process models.

Let $X_i(t)$ be the observed degradation path for unit i . Consider the nonlinear random-effect model

$$X_i(t) = \Lambda(t; \theta, \omega_i) + \epsilon_i(t) \quad (6)$$

where $\Lambda(t)$ is a smooth monotone function of t , θ is the population-level parameter characterizing degradation, ω_i is the random effect associated with subject i , and $\epsilon_i(t)$ is measurement error. In general, the degradation path $\Lambda(t; \theta, \omega_i)$ will be nonlinear and could be chosen by physical mechanism theory as discussed in the section titled “Physical Degradation”. For example, data from a linear random-effect model

is given by

$$X_i(t_j) = \omega_i t_j + \epsilon_i(t_j), \quad i = 1, \dots, n, \quad j = 1, \dots, m \quad (7)$$

Suppose we assume that ω_i 's are i.i.d. lognormal and $\epsilon_i(t_j)$'s are i.i.d. normally. TTF is defined as the first-passage time of the true degradation path $\omega_i t$ to the prespecified threshold D . The distribution function of TTF $T_i = D/\omega_i$ is $F_T(t) = P(\omega_i t > D) = 1 - F_{LN}(D/t)$, where F_{LN} is the lognormal distribution function. The TTF distribution function can be obtained analytically for simple models. For general nonlinear random-effect model in equation (6), Meeker and Escobar ([2], Chapter 21) also describe a simulation-based approach and the use of **bootstrap** methods for **estimation** and inference. Inference under the random-effect model (6) has been discussed in the literature, see, for example, Lu and Meeker [4] and Meeker *et al.* [27]. The common approach is to maximize the likelihood function directly. When the likelihood does not have a closed form, various approximations have been proposed [28] and are implemented in S-plus functions `lme` and `nmle`.

Next, we shall discuss two common nonhomogeneous Lévy processes models. First consider a nonhomogeneous Gaussian process $\{X(t) : t \geq 0\}$. Here, $X(t)$ represents the degradation measurement for an individual unit at time t . Let $\Lambda(t)$ be a continuous increasing function. The nonhomogeneous Gaussian process has the following properties: (a) the increments $\Delta X(t) = X(t + \Delta t) - X(t)$ are independent and (b) $\Delta X(t)$ has a **normal distribution** with mean $\Delta \Lambda(t) = \Lambda(t + \Delta t) - \Lambda(t)$ and variance $\sigma^2 \Delta \Lambda(t)$. Let $Z(t)$ be the basic Wiener process to have linear drift with $\nu = 1$. Then, we get the nonhomogeneous Gaussian process through time transformation as $Y(t) = Z(\Lambda(t))$. Thus, we can rewrite the model as

$$X(t) = \Lambda(t) + \sigma W(\Lambda(t)) \quad (8)$$

Note that this process $X(t)$ has continuous sample path and independent increments. The level crossing of the cumulative degradation threshold D by the process $X(t)$ can be obtained in terms of the inverse Gaussian (IG) distribution as $T_\Lambda = \Lambda^{-1}(T_{IG})$, where T_{IG} is $IG(D, D^2/\sigma^2)$. There is a huge literature on IG distribution and its applications ([29–31] and references therein). Assume that Λ belongs to a known parametric family from the knowledge of the degradation mechanism. For example, the family of power law curves $\Lambda(t; \nu, \beta) = \nu t^\beta$ can be used

to model increasing, constant, and decreasing degradation rates. An example with a finite asymptote is $\Lambda(t; \nu, \beta) = C(1 - \exp(-\nu t^\beta))$. These two families have been considered by Whitmore and Schenkelberg [17] in the context of accelerated degradation models. It is straightforward to get maximum-likelihood estimates of the parameters based on regular degradation data, since the differences of consecutive degradation measurements are independent normal random variables. In most cases this has to be done numerically. The Fisher information can be used, as usual, to get **confidence intervals** and test hypotheses.

Inference under nonhomogeneous gamma process is similar. The gamma process, $Y(t)$, has an independent increment $\Delta Y(t)$ which has gamma distribution $\text{Gam}(\Delta \Lambda(t), \alpha)$ with mean $\alpha \Delta \Lambda(t)$ and variance $\alpha^2 \Delta \Lambda(t)$. Unlike the Wiener process that has a non-monotone and continuous path, the gamma process has a monotone path and evolves on each time interval by means of an infinity of jumps. The TTF distribution can be obtained directly,

$$\begin{aligned} P(T < t) &= P(Y(t) > d) = 1 - G(d; \Lambda(t), \alpha) \\ &= I\left(\Lambda(t), \frac{d}{\alpha}\right) \end{aligned} \quad (9)$$

where $G(\cdot; \Lambda(t), \alpha)$ is the gamma distribution function with parameters $\Lambda(t)$ and α and $I(t, \alpha)$ is the upper incomplete gamma function, $I(t, \alpha) = \int_\alpha^\infty x^{t-1} e^{-x} dx / \Gamma(t)$. The unknown parameters are again estimated by maximizing the likelihood function.

An Illustrative Example

We will use simulated data to illustrate the parametric analysis and compare the estimators from degradation data and TTF data. We use a power law model $X(t) = \nu t^\beta + \sigma W(t^\beta)$ with $\nu = 1$, $\beta = 2$, and $\sigma^2 = 4$ be the true values of the parameters. Degradation data are simulated for $n = 40$ specimens at 10 time points $t_1 = 1, \dots, t_{10} = 10$. Figure 1(a) is the plot of this simulated data. Failure is defined when the degradation path crosses $D = 100$.

The maximum-likelihood estimators are given as: $[\hat{\nu}, \hat{\beta}, \hat{\sigma}] = [1.14, 1.96, 4.28]$ and the estimated standard errors are: 0.040, 0.096, and 0.345, respectively. Figure 1(b) shows the estimated $\Lambda(t)$ superimposed on the data (in dash-dotted line) and the pointwise

4 Physical Degradation Models

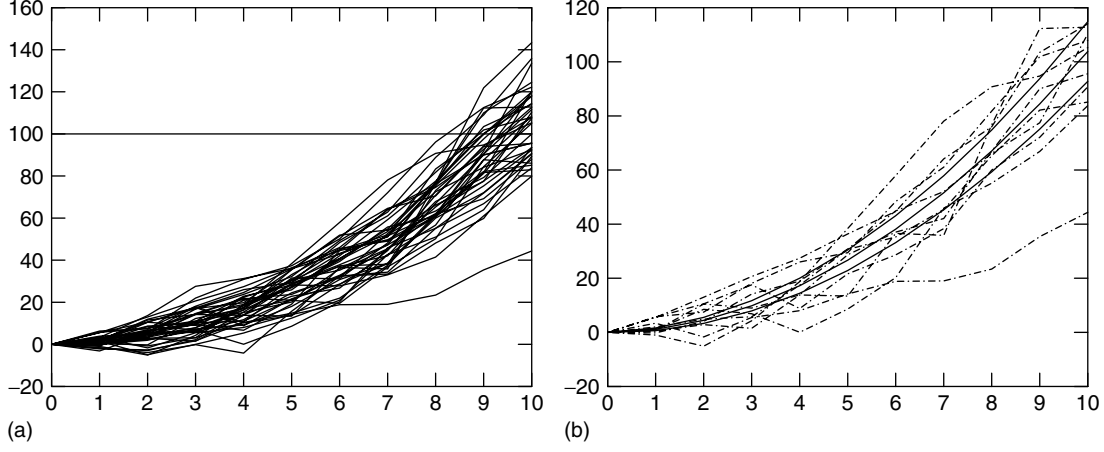


Figure 1 Simulated degradation data and estimated Λ and its pointwise 95% confidence intervals

Table 1 Maximum-likelihood estimators and their standard errors for degradation data and TTF data

Parameter	ν	β	τ
True	1.00	2.00	4.00
TTF data	0.24 (0.182)	2.62 (8.719)	6.41 (30.498)
Degradation data	1.14 (0.040)	1.96 (0.096)	4.28 (0.345)

95% confidence intervals. As we can see, the estimated $\hat{\Lambda}$ follows the data very well.

For the simulated data in Figure 1, suppose we observed only TTF data corresponding to the first-passage times above the threshold D , that is, the interval censored data. On the basis of this data, the MLEs are given by Table 1: $[\hat{\nu}, \hat{\beta}, \hat{\tau}] = [0.24, 2.62, 6.41]$, which are not good estimates compared to the estimates from degradation data. For the total 40 units, 3 units failed at interval $[t_8, t_9]$, 22 failed at $[t_9, t_{10}]$, and the other 11 failed after t_{10} . We do not have enough information to give good estimates of the parameters. We simulate the data 1000 times and the standard errors of these estimators from TTF are given in Table 1, which are very large. It is not uncommon to see situations where reliability engineers collect degradation data but convert them to “TTF” data and then do a standard Weibull or lognormal analysis with the pseudodata. Such converting is quite inefficient and there is tremendous gain to be realized from degradation data.

Discussion

An interesting extension for current degradation models is when we have information on covariates. For example, in accelerated degradation experiment, the stress level such as temperature is a **covariate**. Boulanger and Escobar [32] studied the design of accelerated degradation tests under the nonlinear random-effect model

$$X_{ij}(t) = \alpha_{ij} F(t, \xi_{ij}) + \epsilon_{ij}(t) \quad (10)$$

where $X_{ij}(t)$ is the degradation measurement at time t for the j th unit under stress condition i . F is a known continuous monotone increasing function of t with $F(0, \xi_{ij}) = 0$ and $F(\infty, \xi_{ij}) = 1$. Hence, α_{ij} measures the asymptote of the degradation level. Also, (α_{ij}, ξ_{ij}) are random coefficients with $E[\log \alpha_{ij}] = A + Bs_i$, where s_i is the i th stress level. If we have data from accelerated degradation tests, we may use model (10) to fit the data and estimate the unknown parameters. We can also use model (10) to develop the optimal degradation planning. The D -optimal plan is to maximize the determinant of Fisher information of the maximum-likelihood estimator of unknown parameters.

For nonhomogeneous Lévy processes models, Bagdonavicius and Nikulin [19] studied nonhomogeneous gamma process with covariates via an **accelerated life model** by replacing $\Lambda(t)$ by $\Lambda(te^{x^T \beta})$, where x is the covariate vector. Lawless and Crowder [20] treated scale parameter α in a gamma process as a

function of x to accommodate the covariate. Similar to **Cox** model, we may study the proportional mean model replacing $\Lambda(t)$ by $\Lambda(t)e^{x^T\beta}$ to incorporate the covariate information.

Another interesting issue to study is the optimal maintenance strategies for physical degradation systems. Good degradation models can be used as a tool to optimize maintenance decisions and to minimize the total maintenance cost of the system. Grall *et al.* [33] studied a predictive maintenance structure for a degradation system. They assumed the degradation process is a homogeneous gamma process and provided the maintenance policy to (a) determine whether the system should be replaced preventively or correctively, or whether the system should be left as is; (b) determine the time to the next inspection. There is a need for additional research to develop maintenance policy for other more complicated models.

References

- [1] Davies, A. (1998). *Handbook of Condition Monitoring: Techniques and Methodology*, Chapman & Hall, London.
- [2] Meeker, W.Q. & Escobar, L.A. (1998). *Statistical Methods for Reliability Data*, John Wiley & Sons, New York.
- [3] Hudak Jr, S.J., Saxena, A., Bucci, R.J. & Malcolm, R.C. (1978). Development of standard methods of testing and analyzing fatigue crack growth rate data, Technical Report AFML-TR-78-40, Westinghouse R&D Center, Westinghouse Electric Corporation, Pittsburgh.
- [4] Lu, C.J. & Meeker, W. (1993). Using degradation measures to estimate a time-to-failure distribution, *Technometrics* **35**, 161–174.
- [5] Elsayed, E.A. & Liao, H.T. (2004). A geometric Brownian motion model for field degradation data, *International Journal of Materials and Product Technology* **20**, 51–72.
- [6] Escobar, L.A., Meeker, W., Kugler, D. & Kramer, L. (2004). Accelerated destructive degradation tests: data, models and analysis, *Mathematical and Statistical Methods in Reliability*, B. H.Lindqvist & K. A. Doksum, eds, World Scientific Publishing Company.
- [7] Longford, N.T. (1993). *Random Coefficient Models*, Clarendon Press, Oxford.
- [8] Lindsey, J.K. (1993). *Models for Repeat Measures*, Clarendon Press, Oxford.
- [9] LuValle, M.J., Welscher, T.L. & Svoboda, K. (1988). Acceleration transforms and statistical kinetic models, *Journal of Statistical Physics* **52**, 311–320.
- [10] Meeker, W.Q. & LuValle, M.J. (1995). An Accelerated life test model based on reliability kinetics, *Technometrics* **37**, 133–146.
- [11] Dowling, N.E. (1993). *Mechanical Behavior of Materials*, Prentice Hall, Englewood Cliffs.
- [12] Burnett, J. & Hu, C. (1988). Modeling hot-carrier effects in polysilicon emitter bipolar transistors, *IEEE Transactions on Electron Devices* **35**, 2238–2244.
- [13] Signor, A.W., VanLandingham, M.R. & Chin, J.W. (2003). Effects of ultraviolet radiation exposure on vinyl ester resins: characterization of chemical, physical and mechanical damage *Polymer Degradation and Stability* **79**(2), 359–368.
- [14] Doksum, K. (1991). Degradation rate models for failure time and survival data, *CWI Quarterly* **4**, 195–203.
- [15] Doksum, K. & Hoyland, A. (1992). Models for variable stress accelerated life testing experiments based on wiener process and the inverse Gaussian distribution, *Technometrics* **34**, 74–82.
- [16] Whitmore, G.A. (1995). Estimation degradation by a Wiener diffusion process subject to measurement error. *Lifetime Data Analysis* **1**, 307–319.
- [17] Whitmore, G.A. & Schenkelberg, F. (1997). Modelling accelerated degradation data using wiener diffusion with a time scale transformation. *Lifetime Data Analysis* **3**, 27–43.
- [18] Singpurwalla, N.D. & Youngren, M.A. (1998). Multivariate distributions induced by dynamic environments. *Scandinavian Journal of Statistics* **20**, 251–261.
- [19] Bagdonavicius, V. & Nikulin, M. (2000). Estimation in degradation models with explanatory variables, *Lifetime Data Analysis* **7**, 85–103.
- [20] Lawless, J. & Crowder, M. (2004). Covariate and random effects in a Gamma process model with application to degradation and failure, *Lifetime Data Analysis* **10**, 213–227.
- [21] Zacks, S. (2004). Distributions of failure times associated with non-homogeneous compound poisson damage processes, in *A Festschrift for Herman Rubin, Lecture Notes-Monograph Series Vol. 45*, A. DasGupta, ed, Institute of Mathematical Statistics, Beachwood, pp. 396–407.
- [22] Gertsbakh, I.B. (1989). *Statistical Reliability Theory*, Marcel Dekker, New York.
- [23] Esary, J.D., Marshall, A.W. & Proschan, F. (1973). Shock models and wear processes, *Annals of Probability* **1**, 627–649.
- [24] Bogdanoff, J.L. & Kozin, F. (1985). *Probabilistic Models of Cumulative Damage*, John Wiley & Sons, New York.
- [25] Hougaard, P. (1999). Multi-state models: a review, *Lifetime Data Analysis* **5**, 239–246.
- [26] Singpurwalla, N.D. (1995). Survival in dynamic environments, *Statistical Science* **1**, 86–103.
- [27] Meeker, W.Q., Escobar, L.A. & Lu, C.J. (1998). Accelerated degradation tests: modeling and analysis, *Technometrics* **40**, 89–99.
- [28] Pinheiro, J.C. & Bates, D.M. (2000). *Mixed Effects Models in S and S-PLUS*, Springer-Verlag, New York.
- [29] Chikkara, R.S. & Folks, J.L. (1989). *The Inverse Gaussian Distribution*, Marcell Dekker, New York.

6 Physical Degradation Models

- [30] Sheshadri, V. (1999). *The Inverse Gaussian Distribution*, Springer, New York.
- [31] Jorgensen, B. (1982). *Statistical Properties of the Generalized Inverse Gaussian Distribution*, Springer-Verlag, New York.
- [32] Boulanger, M. & Escobar, L.A. (1994). Experimental design for a class of degradation tests, *Technometrics* **36**, 260–272.
- [33] Grall, A., Dieulle, L., Berenguer, C. & Roussignol, M. (2002). Continuous-time predictive-maintenance scheduling for a deteriorating system, *IEEE Transactions on Reliability* **51**(2), 141–150.

Related Articles

Degradation and Failure; Degradation Models; Degradation Processes.

XIAO WANG

Further Reading

- Dirkse, J.P. (1975). An absorption probability for the Ornstein-Uhlenbeck process, *Journal of Applied Probability* **12**, 595–599.

Degradation and Failure

Introduction

Traditionally the failure time data are usually used for product reliability **estimation**. Failures of highly reliable units are rare and other information should be used in addition to censored failure time data. One way of obtaining complementary reliability information is to use higher levels of experimental factors or covariates (such as temperature, voltage, or pressure) to increase the number of failures and, hence, to obtain reliability information quickly. The *accelerated life testing* (ALT) of biotechnical systems is meant to be a simple practical method of estimation of the reliability of new systems without having to wait through the operating life of them. It is evident that the *extrapolating reliability* from ALT always carries the risk that the accelerating stresses do not properly excite the failure mechanism which dominates at operating (*normal*) stresses. Another way of obtaining this complementary reliability information is to measure some parameters (covariates) that characterize the aging of the product in time. In analysis of *longevity* of highly reliable complex industrial or biological systems, the **degradation** processes provide additional information about the *aging*, *degradation*, and *deterioration* of systems, and from this point of view the degradation data are really a very rich source of reliability information and often offer many advantages over failure time data. Degradation is the natural response for some tests, and it is also natural that with degradation data it is possible to make useful reliability and statistical inference, even with no failures. It is evident that it may be difficult or costly to collect degradation measures from some components or materials. Sometimes it is possible to apply the expert's estimation of the level of degradation. Sometimes it is possible to construct degradation models in which degradation is measured *without errors*, but there are also situations when the degradation data are obtained with measurement errors.

Degradation Process

In reliability there is considerable interest in constructing **covariate** $Z(\cdot)$ with some properties, depending on the phenomena in consideration, which

describe the *real* process of wear or the usage history up to time t , indicate the level of fatigue, degradation, and deterioration of a system, and may influence the rate of degradation, risk of failure, and reliability of the system. Statistical modeling of *observed degradation* processes can help to understand different real physical, chemical, medical, biological, physiological, or social degradation processes of aging. Information about *real degradation* processes helps us to construct degradation models, which permit us to predict the cumulative damage, and so on. According to the principal idea of degradation models, a *soft failure* is observed if the degradation process reaches a *critical threshold* z_0 . A soft failure caused by degradation occurs when $Z(t)$ reaches the value z_0 . The *moment*, T^0 , of soft failure is defined by the relation

$$T^0 = \sup\{t : Z(t) < z_0\} = \inf\{t : Z(t) \geq z_0\} \quad (1)$$

So it is reasonable to construct the so-called degradation models, based on some suppositions about the mathematical properties of the degradation process $Z(\cdot)$, in accordance with observed *longitudinal* data in the experiment. Both considered methods may be combined to construct the so-called joint degradation models. For this, it is enough to define a failure of the system when its degradation (*internal wear*) reaches a critical value or a traumatic event occurs. The joint degradation models form the class of models with **competing risk**, since for any item we consider two competing causes of failure: degradation, reaching a *threshold* (soft failure), and occurrence of a *hard* or *traumatic failure*. Let T^0 be the moment of traumatic failure of a unit. In the considered class of models, the *failure time* τ is the *minimum* of the moments of traumatic and soft failures, $\tau = \min(T^0, T) = T^0 \wedge T$. These models are also called *degradation-threshold-shock models*. For such models the degradation process $Z(t)$ can be considered as an additional time-dependent covariable, which describes the process of wear or the usage history up to time t . The degradation models with covariates are used to estimate reliability when the environment is dynamic (see [1–9], etc.). The covariates cannot be controlled by an experimenter in such a case. For example, the tire wear rate depends on the quality of roads, the temperature, and other factors. The problem of optimization of the values of these covariates is considered in [10], see also [11].

In [1] the use of so-called path models is proposed; the authors describe many different applications

2 Degradation and Failure

and models for *accelerated degradation* and Arrhenius analysis for data involving destructive testing. Meeker and Escobar [1] used *convex degradation models* to study the growth of fatigue cracks. They also used *concave degradation models* to study the growth of failure-causing conducting filaments of chlorine-copper compounds in printed-circuit boards [1]. *Concave degradation models* are used in [12] and [1] to describe degradation of components in electronic circuits.

In [1, 10, 13] *general degradation path models* are considered, according to which

$$Z(t) = g(t, \mathbf{A}) \quad (2)$$

here $\mathbf{A} = (\mathbf{A}_1, \dots, \mathbf{A}_r)$ is a random vector with positive components and the distribution function π of $\mathbf{A}$, and g is a specified continuously differentiable function in t , which increases from 0 to $+\infty$ when t increases from 0 to $+\infty$. For example, $g(t, a) = t/a$ (linear degradation, $r = 1$) or $g(t, a) = (t/a_1)^{a_2}$ (convex or concave degradation, $r = 2$). It is also supposed that, for each $t > 0$, the degradation path $(g(s, a) | 0 < s \leq t)$ determines in a unique way the value a of the random vector $\mathbf{A}$. Degradation in these models is modeled by the process $Z(t, \mathbf{A})$, where t is time and A is some possibly multidimensional random variable. In [14] *linear degradation models* are used to study the increase in a resistance measurement over time, and in [1] and [12] *convex degradation models* and *concave degradation models* are used to study the growth of fatigue cracks.

In [1, 2, 4, 9, 15] many different applications and models for *accelerated degradation* for data involving a destructive testing are described. The influence of covariates on degradation is also modeled in [1, 4, 5] to estimate reliability when the environment is dynamic. In [16] the parametric and semiparametric estimates in the model described in [4] are compared. The semiparametric analysis of several new degradation and failure time regression models without and with covariables is described in [6–8]. More information can be found in [1, 15, 17, 18].

Example: Degradation model with noise

Suppose that the lifetime of a unit is determined by the degradation process $Z(t)$ and the moment of its potential traumatic failure is T . Denote by T^0 the moment at which the degradation attains some critical

value z_0 . Then the moment of the unit's failure is $\tau = T^0 \wedge T$. To model the degradation–failure time process, we suppose that the real degradation process is modeled by the *general path model*:

$$Z_r(t) = g(t, \mathbf{A}) \quad (3)$$

where $\mathbf{A} = (A_1, \dots, A_r)$ is a random positive vector, and g is continuously differentiable function that increased in t . As remarked earlier, the real degradation process $Z_r(t)$ is often not observed, and we have to measure (to estimate) it. In this case, the *observed degradation process* $Z(t)$ is different from the real degradation process $Z_r(t)$. We supposed that we have the *degradation model with noise* according to which the *observed degradation process* is

$$Z(t) = Z_r(t)U(t), \quad t > 0 \quad (4)$$

where $\ln U(t) = V(t) = \sigma W(c(t))$, W is the standard *Wiener process* independent of A , and c is a specified continuous increasing function, with $c(0) = 0$. For any $t > 0$ the *median* of $U(t)$ is 1.

Using this model it is easy to construct a *degradation model with noise*. An important class of models based on degradation processes was developed recently in [19] and [4–6]. Wulfssohn and Tsiatis [19] proposed the so-called joint model for survival and longitudinal data measured with error, given by

$$\lambda_T(t|A) = \lambda_0(t)e^{\beta(A_1 + A_2 t)} \quad (5)$$

where $A = (A_1, A_2)$ follows bivariate **normal distribution**. On the other hand Bagdonavicius and Nikulin [7] proposed the model in terms of a conditional survival function of T given the real degradation process:

$$\begin{aligned} S_T(t|A) &= P\{T > t | g(s, A), 0 \leq s \leq t\} \\ &= \exp \left\{ - \int_0^t \lambda_0(s, \theta) \lambda(g(s, A)) ds \right\} \end{aligned} \quad (6)$$

where λ is the unknown **intensity function**, $\lambda_0(s, \theta)$ is from a parametric family of hazard functions. The distribution of A is not specified. This model states that the conditional hazard rate $\lambda_T(t|A)$ at the moment t given the degradation $g(s, A)$, $0 \leq s \leq t$, has the multiplicative form as in the famous **Cox model**:

$$\lambda_T(t|A) = \lambda_0(s, \theta) \lambda(g(s, A)) \quad (7)$$

If for example, $\lambda_0(s, \theta) = (1 + t)^\theta$ or $\lambda_0(s, \theta) = e^{t\theta}$, then $\theta = 0$ corresponds to the case when the

hazard rate at any moment t is a function of the degradation level at this moment. One can note that in the second model the function λ , characterizing the influence of degradation on the hazard rate, is nonparametric.

References

- [1] Meeker, W. & Escobar, L. (1995). *Statistical Methods for Reliability Data*, John Wiley & Sons, New York.
- [2] Nelson, W. (1990). *Accelerated Testing: Statistical Models, Test Plans, and Data Analysis*, John Wiley & Sons, New York.
- [3] Singpurwalla, N. (1995). Survival in dynamic environments, *Statistical Science* **10**, 86–103.
- [4] Bagdonavicius, V. & Nikulin, M. (2001). Estimation in degradation models with explanatory variables, *Lifetime Data Analysis* **7**, 85–103.
- [5] Bagdonavicius, V. & Nikulin, M. (2002). *Accelerated Life Models*, Chapman & Hall/CRC, Boca Raton.
- [6] Bagdonavicius, V., Bikelis, A., Kazakevicius, V. & Nikulin, M. (2003). Estimation from simultaneous degradation and failure data, in *Mathematical and Statistical Methods in Reliability*, B. Linquist & K.A. Doksum, eds, World Scientific, Vol. 7, pp. 301–318.
- [7] Bagdonavicius, V. & Nikulin, M. (2004). Semiparametric analysis of degradation and failure time data with covariate, in *Parametric and Semiparametric Models with Applications to Reliability, Survival Analysis, and Quality of Life*, M. Nikulin, N. Balakrishnan, M. Mesbah & N. Limnios, eds, Birkhauser, Boston, pp. 41–65.
- [8] Bagdonavicius, V., Haghighi, F. & Nikulin, M. (2005). Statistical analysis of general degradation path model and failure time data with multiple failure modes, *Communications in Statistics: Theory and Methods* **34**, 1643–1662.
- [9] Lawless, J. (2003). *Statistical Models and Methods for Lifetime Data*, John Wiley & Sons, New York.
- [10] Chiao, C.H. & Hamada, M. (1996). Using degradation data from an experiment to achieve robust reliability for light emitting diodes, *Quality and Reliability Engineering International* **12**, 89–94.
- [11] Ceci, C. & Mazliak, L. (2004). Optimal design in nonparametric life testing, *Statistical Inference for Stochastic Processes* **7**, 305–325.
- [12] Carey, M. & Koenig, R. (1991). Reliability assessment based on accelerated degradation: a case study, *IEEE Transactions on Reliability* **40**, 499–506.
- [13] Lu, C. & Meeker, W. (1993). Using degradation measures to estimate a time-to-failure distribution, *Technometrics* **35**, 161–174.
- [14] Suzuki, K., Maki, K. & Yokogawa, K. (1993). An analysis of degradation data of a carbon film and properties of the estimators, in *Statistical Sciences and Data Analysis*, K. Matusita, M. Puri & T. Hayakawa, eds, VSP, Utrecht.
- [15] Lehmann, A. (2004). On a degradation-failure model for repairable items, in *Parametric and Semiparametric Models with Applications to Reliability, Survival Analysis, and Quality of Life*, M. Nikulin, N. Balakrishnan, M. Mesbah & N. Limnios, eds, Birkhauser, Boston, pp. 65–80.
- [16] Couallier, V. (2004). Comparison of parametric and semi-parametric in joint models with covariables and traumatic censoring, in *Parametric and Semiparametric Models with Applications to Reliability, Survival Analysis, and Quality of Life*, M. Nikulin, N. Balakrishnan, M. Mesbah & N. Limnios, eds, Birkhauser, Boston, pp. 81–99.
- [17] Wendt, H. & Kahle, W. (2004). On parametric estimation for a position-dependent marking of a doubly stochastic Poisson process, in *Parametric and Semiparametric Models with Applications to Reliability, Survival Analysis, and Quality of Life*, M. Nikulin, N. Balakrishnan, M. Mesbah & N. Limnios, eds, Birkhauser, Boston, pp. 473–486.
- [18] Kahle, W. & Wendt, H. (2006). Statistical analysis of some parametric degradation models, in *Probability, Statistics and Modelling in Public Health*, M. Nikulin, D. Commenges & C. Huber, eds, Springer, pp. 266–279.
- [19] Wulfsohn, M. & Tsiatis, A. (1997). A joint model for survival and longitudinal data measured with error, *Biometrics* **53**, 330–339.

Further Reading

- Doksum, K.A. & Hoyland, A. (1992). Models for variable-stress accelerated life testing experiment based on Wiener processes and the inverse Gaussian distribution, *Technometrics* **34**, 74–82.
- Doksum, K.A. & Normand, S.L.T. (1995). Gaussian models for degradation processes - part I: methods for the analysis of biomarker data, *Lifetime Data Analysis* **1**, 131–144.
- Harlamov, B. (2004). Inverse gamma-process as a model of wear, in *Longevity, Aging and Degradation Models in Reliability, Health, Medicine and Biology*, V. Antonov, C. Huber, M. Nikulin & V. Politschook, eds, St. Petersburg State Polytechnical University, Saint Petersburg, Vol. 2, pp. 180–190, ISBN 5-7422-0638-0.
- Singpurwalla, N. (1997). Gamma processes and their generalizations: an overview, in *Engineering Probabilistic Design and Maintenance for Flood Protection*, R. Cook, M. Mendel & H. Vrijling, eds, Kluwer Academic Publishers, Dordrecht, pp. 67–73.
- Meeker, W., Escobar, L. & Lu, C. (1998). Accelerated degradation tests: modeling and analysis, *Technometrics* **40**, 89–99.
- Whitmore, G.A. & Schenkelberg, F. (1997). Modelling accelerated degradation data using Wiener diffusion with a time scale transformation, *Lifetime Data Analysis* **3**, 27–45.

MIKHAIL NIKULIN

Accelerated Life Models

Introduction

In industry the failure time data are used for product reliability **estimation**. Failures of highly reliable units are rare and information other than traditional censored failure time data should be used. One way of obtaining this complementary reliability information is to apply the methods of the so-called *accelerated life testing* (ALT). The purpose of ALT is to give estimators of the main reliability characteristics under *usual* (normal, standard) stress (**design**) using data of *accelerated experiments* when units are tested at higher than usual stress conditions, that is, ALT is supposed to use higher levels of experimental time-varying factors or covariates (such as temperature, voltage, or pressure) to increase the number of failures and, hence, obtain reliability information more quickly. The so-called *accelerated life models* (ALMs) with time depending explanatory variables (stresses) are used to treat these data in ALT. The ALMs with dynamic environment give the possibility to include the effect of the usage history on the lifetime distributions of various units, subjects, items, populations, and so on. They are also well adapted to obtain the statistical inference in accelerated experiments with dynamic environment, when data are censored. All famous models such as **Cox** (proportional hazards) model, linear transformation model, frailty model, generalized proportional model, additive hazards model, accelerated failure time (AFT) model, and so on, are in the class of ALMs.

Suppose at first that the explanatory variable $x(\cdot) \in E$ is a deterministic time function:

$$x(\cdot) = (x_1(\cdot), \dots, x_m(\cdot))^T : [0, \infty) \rightarrow B \in \mathbf{R}^m \quad (1)$$

where E is a set of all possible stresses, $x_i(\cdot)$ is a univariate explanatory variable. If $x(\cdot)$ is constant in time, $x(\cdot) \equiv x = \text{const}$, we shall write x instead of $x(\cdot)$ in all formulas. We note E_1 a set of all constant in time stresses, $E_1 \subset E$.

Denote informally by $T_{x(\cdot)}$ the failure time under $x(\cdot)$ and by

$$S_{x(\cdot)}(t) = \mathbf{P}\{T_{x(\cdot)} \geq t\}, \quad t > 0, \quad x(\cdot) \in E \quad (2)$$

the *survival* or *reliability function* of $T_{x(\cdot)}$. The *hazard rate function* of $T_{x(\cdot)}$ under $x(\cdot)$ is

$$\begin{aligned} \alpha_{x(\cdot)}(t) &= \lim_{h \downarrow 0} \frac{1}{h} \mathbf{P}\{T_{x(\cdot)} \in [t, t+h) | T_{x(\cdot)} \geq t\} \\ &= -\frac{S'_{x(\cdot)}(t)}{S_{x(\cdot)}(t)} \end{aligned} \quad (3)$$

Denote by

$$A_{x(\cdot)}(t) = \int_0^t \alpha_{x(\cdot)}(u) du = -\ln\{S_{x(\cdot)}(t)\} \quad (4)$$

the *cumulative hazard function* of $T_{x(\cdot)}$.

We say that a stress $x(\cdot)$ is higher than a stress $y(\cdot)$ and we write $x(\cdot) > y(\cdot)$, if for any $t \geq 0$ the inequality $S_{y(\cdot)}(t) \geq S_{x(\cdot)}(t)$ holds and there exists $t_0 > 0$ such that $S_{y(\cdot)}(t_0) > S_{x(\cdot)}(t_0)$. In ALT the most used types of stresses are constant in time stresses, **step stresses**, progressive (monotone) stresses, cyclic stresses, and random stresses (see, e.g., Singpurwalla [1], Viertl [2], Nelson [3, 4], Meeker and Escobar [5], Lawless [6]). The most common case is when the stress is unidimensional, for example, pressure, temperature, voltage, but then more than one accelerated stresses may be used, and so on. The most frequently used time-varying stresses in ALT are step stresses: units are placed on test at an initial low stress and if they do not fail in a predetermined time t_1 , the stress is increased. If they do not fail in a predetermined time $t_2 > t_1$, the stress is increased once more, and so on. Thus step stresses have the form

$$x(u) = \begin{cases} x_1, & 0 \leq u < t_1, \\ x_2, & t_1 \leq u < t_2, \\ \dots & \dots \\ x_k, & t_{k-1} \leq u < t_k \end{cases} \quad (5)$$

where $x_1, \dots, x_k$ are constant stresses, $x_i \in E_1$. Sets of step stresses of the form equation (5) will be denoted by E_k , $E_k \subset E$. Let E_2 , $E_2 \subset E$, be a set of step stresses of the form

$$x(u) = \begin{cases} x_1, & 0 \leq u < t_1 \\ x_2, & u \geq t_1 \end{cases} \quad (6)$$

Sedyakin's Principle

In 1966, Sedyakin formulated his famous *physical principle in reliability*, which states that for two identical populations of units functioning under different

2 Accelerated Life Models

constant stresses x_1 and x_2 , two moments t_1 and t_2 are equivalent if the probabilities of survival until these moments are equal:

$$\begin{aligned} \mathbf{P}\{T_{x_1} \geq t_1\} &= S_{x_1}(t_1) = S_{x_2}(t_2) \\ &= \mathbf{P}\{T_{x_2} \geq t_2\}, \quad x_1, x_2 \in E_1 \end{aligned} \quad (7)$$

If after these equivalent moments the units of both groups are observed under the same stress x_2 , that is, the first population is observed under the step stress $x(\cdot) \in E_2$ of the form equation (6) and the second all the time under the constant stress x_2 , then for all $s > 0$

$$\alpha_{x(\cdot)}(t_1 + s) = \alpha_{x_2}(t_2 + s) \quad (8)$$

Using this idea of Sedyakin, it was proposed (see [7]) to consider one generalization of the Sedyakin's model to the case of any time-varying stresses by supposing that the hazard rate $\alpha_{x(\cdot)}(t)$ at any moment t is a function of the value of the stress at this moment and of the probability of survival until this moment:

$$\alpha_{x(\cdot)}(t) = g(x(t), A_{x(\cdot)}(t)), \quad x(\cdot) \in E \quad (9)$$

for some positive function g on $E \times \mathbf{R}^+$. This model is called the *generalized Sedyakin model* (GSM). Let us consider the meaning of the *rule of time-shift* in equation (9) in the case of the GSM when one applies the step stresses in equation (5) or equation (6). It is easy to show that if the GSM holds on E_2 then the hazard rate under the stress $x(\cdot) \in E_2$ satisfies the equality

$$\alpha_{x(\cdot)}(t) = \begin{cases} \alpha_{x_1}(t), & 0 \leq t < t_1 \\ \alpha_{x_2}(t - t_1 + t_1^*), & t \geq t_1 \end{cases} \quad (10)$$

where the moment t_1^* is determined by the equality $S_{x_1}(t_1) = S_{x_2}(t_1^*)$.

From this proposition it follows that for any $x(\cdot) \in E_2$ and for all $s \geq 0$

$$\alpha_{x(\cdot)}(t_1 + s) = \alpha_{x_2}(t_1^* + s) \quad (11)$$

This is the model on E_2 , proposed by Sedyakin [8].

If one takes a set E_k of stepwise stresses of the form equation (5), the GSM represents the famous *basic cumulative exposure model*, see Bhattacharyya and Stoejoeti [9], Nelson [3], Meeker and Escobar [5]. We note here that the Sedyakin model (SM) can be not appropriate in situations of periodic and quick

change of the stress level or when switch-up's of the stress from one level to the another can imply failures or shorten the life.

Accelerated Failure Time Model

The famous AFT model on E will be considered here. The AFT model holds on E if there exist on E a positive function r and on $[0, \infty)$ a survival function S_0 such that

$$S_{x(\cdot)}(t) = S_0\left(\int_0^t r\{x(u)\} du\right) \quad \text{for any } x(\cdot) \in E \quad (12)$$

where the survival function S_0 , called the *baseline survival function*, does not depend on x . It is the Sedyakin's prolongation of the AFT model for constant stresses, given by the formula:

$$S_x(t) = S_0\{r(x)t\} \quad \text{for any } x \in E_1 \quad (13)$$

For any fixed $x \in E_1$ the value $r(x)$ can be interpreted as the *timescale change constant* or the acceleration constant of reliability function. More about one AFT model can see, for example, Cox and Oakes [10], Lawless [6, 11], LuValle [12], Meekers and Escobar [5], Nelson [3, 4], Viertl [2], Singpurwalla [13], Duchesne [14], and so on.

Natural generalization of the AFT model, known as the *changing shape and scale* (CHSS) model (see [7]) holds on E if there exist two positive functions r and ν on E such that

$$S_{x(\cdot)}(t) = S_{x_0}\left(\int_0^t r\{x(\tau)\} \tau^{\nu(x(\tau))-1} d\tau\right), \quad x(\cdot) \in E \quad (14)$$

where x_0 is fixed stress, for example, design (usual) stress. Note that the variation of stress changes locally not only the scale but also the shape of distribution. This model is not in the class of the GSMs because the hazard rate

$$\alpha_{x(\cdot)}(t) = r\{x(t)\} q(A_{x(\cdot)}(t)) t^{\nu(x(t))-1} \quad (15)$$

depends not only on $x(t)$ and $A_{x(\cdot)}(t)$ but also on t . Generally it is recommended to choose $\nu(x) = e^{\gamma T x}$. Statistical analysis of this model is done in Bagdonavičius *et al.* [15]. If the functions $r(\cdot)$ and $S_0(\cdot)$ are unknown, we have a nonparametric

AFT (or CHSS) model. **Nonparametric analysis** of AFT and CHSS models can be found, for example, in Bagdonavičius and Nikulin [7], Bagdonavičius *et al.* [15]. The function $r(\cdot)$ can be parameterized. In this case we have a parametric AFT (or CHSS) model and the maximum-likelihood estimators of the parameters are obtained by almost standard way for any plans. If the baseline function S_0 is completely unknown, in this case we have a semiparametric model. In practice, often a baseline survival function S_0 is taken from some class of simple parametric distributions, such as Weibull, lognormal, log-logistic, and so on, and in this case we obtain the parametric AFT (or CHSS) model. An interesting family is obtained when S_0 belongs to the power generalized Weibull (PGW) family of distributions (see [7]). In terms of the survival functions, the PGW family is given by the next formula:

$$S(t, \sigma, \nu, \gamma) = \exp \left\{ 1 - \left[1 + \left(\frac{t}{\sigma} \right)^\nu \right]^{\frac{1}{\gamma}} \right\},$$

$$t > 0, \gamma > 0, \nu > 0, \sigma > 0 \quad (16)$$

If $\gamma = 1$ we have the Weibull family of distributions. If $\gamma = 1$ and $\nu = 1$, we have the exponential family of distributions. This class of distributions has very nice probability properties. All moments of this distribution are finite. In dependence of parameter values the hazard rate can be constant, monotone (increasing or decreasing), unimodal or $\cap$ -shaped, and bathtub or $\cup$ -shaped. Parametric AFT model was studied by many people, see, for example, Bagdonavičius *et al.* [16], Bagdonavičius *et al.* [17], LuValle [12], Nelson [3, 4], Meeker and Escobar [5], Lawless [6], Singpurwalla [1], and Shaked and Singpurwalla [18].

References

- [1] Singpurwalla, N. (1971). Inference from accelerated life tests when observations are obtained from censored data. *Technometrics* **13**, 161–170.
- [2] Viertl, R. (1988). *Statistical Methods in Accelerated Life Testing*, Vandenhoeck and Ruprecht, Göttingen.
- [3] Nelson, W. (1990). *Accelerated Testing: Statistical Models, Test Plans, and Data Analysis*, John Wiley & Sons, New York.
- [4] Nelson, W. (2005). A bibliography of accelerated test plan, Part II, *IEEE Transaction on Reliability* **54**(3), 370–373.
- [5] Meeker, W. & Escobar, L. (1998). *Statistical Methods for Reliability Data*, John Wiley & Sons, New York.
- [6] Lawless, J.F. (2003). *Statistical Models and Methods for Lifetime Data*, John Wiley & Sons, New York.
- [7] Bagdonavičius, V. & Nikulin, M. (2002). *Accelerated Life Models*, Chapman and Hall/CRC, Boca Raton.
- [8] Sed'yakin, N. (1966). On one physical principle in reliability theory, *Technical Cybernetics* **3**, 80–87.
- [9] Bhattacharyya, G.K. & Stoejoeti, Z. (1989). A tampered failure rate model for step-stress accelerated life test, *Communications in Statistics: Theory and Methods* **18**, 1627–1643.
- [10] Cox, D. & Oakes, D. (1984). *Analysis of Survival Data*, Methuen, New York.
- [11] Lawless, J.F. (2000). Dynamic analysis of failures in repairable systems and software, in *Recent Advances in Reliability Theory*, N. Limnios, M. Nikulin, eds, Birkhauser, Boston, pp. 341–354.
- [12] LuValle, M. (2000). A theoretical framework for accelerated testing, in *Recent Advances in Reliability Theory*, N. Limnios & M. Nikulin, eds, Birkhauser, Boston, pp. 419–434.
- [13] Singpurwalla, N. (1995). Survival in dynamic environments, *Statistical Sciences* **1**, 86–103.
- [14] Duchesne, T. (2000). Methods common to reliability and survival analysis, in *Recent Advances in Reliability Theory*, N. Limnios & M. Nikulin, eds, Birkhauser, Boston, pp. 279–290.
- [15] Bagdonavičius, V., Cheminade, O. & Nikulin, M. (2004). Statistical planning and inference in accelerated life testing using the CHSS model, *Journal of Statistical Planning and Inference* **2**, 535–551.
- [16] Bagdonavičius, V., Gerville-Réache, L., Nikoulina, V. & Nikulin, M. (2000). Expériences accélérées: analyse statistique du modèle standard de vie accélérée, *Revue de Statistique Appliquée* **3**, 5–38.
- [17] Bagdonavičius, V., Gerville-Réache, L. & Nikulin, M. (2002). On parametric inference for step-stresses models, *IEEE Transaction on Reliability* **1**, 27–31.
- [18] Shaked, M. & Singpurwalla, N. (1983). Inference for step-stress accelerated life tests, *Journal of Statistical Planning and Inference* **7**, 295–306.

Further Reading

- Bagdonavičius, V. (1978). Testing hypothesis of the additive accumulation of damage. *Theory of Probability and its Applications* **2**, 403–408.
- Nelson, W. (2002). A bibliography of accelerated test plan, *IEEE Transaction on Reliability* **2**, 194–197.

Related Articles

Accelerated Life Tests: Classical Methods for Design and Analysis; Accelerated Life Tests:

Bayesian Design; Accelerated Life Tests: Stress Step (Classical Methods); Accelerated Life Tests: Designs Comparison within a Bayesian Framework; Accelerated Life Tests: Nonparametric

Approach; Accelerated Life Tests: Analysis with Competing Failure Modes; Prediction of Expected Fatigue Lives of Fiber Reinforced Plastic Joints.

MIKHAIL NIKULIN

Competing Risks

Introduction

Suppose that units under study can experience any one of several distinct failure types, and that for each unit we observe both the time to failure and the type of failure. Here failure may, for example, correspond to breakdown of a mechanical component where there are several possible root causes for the failure, such as vibration and corrosion. While this is a typical case of a “competing risks” situation in reliability, the theory of competing risks does not originate from reliability theory. In fact, it can be traced back to David Bernoulli’s attempts in 1760 to disentangle the risk of dying from smallpox from other causes. This is indeed a classical example of competing risks, where individuals are subject to multiple causes of death. Similar applications occur in demography and actuarial science, usually under the name of multiple-decrement analysis.

Formal Definition of Competing Risks

Formally one observes the pair (T, C) where $T > 0$ is the time of failure and $C \in \{1, 2, \dots, k\}$ represents the type of failure. It is thus assumed that there are k different failure types, and that each failure can be classified as belonging to exactly one of the k types. Note that “failure” is used as a generic term and may in practice correspond to any event of interest depending on the application at hand. Also, “time” need not mean calendar time, but can in principle be any suitable measurement which is nondecreasing with calendar time, such as operation time, number of cycles, number of kilometers run, or length of a crack.

An intuitive way of describing a competing risks situation with k risks, is to assume that to each risk is associated a failure time T_j , $j = 1, \dots, k$. These k times are thought of as the hypothetical failure times if the other risks were not present, and they are referred to as *latent failure times*. When all the risks are present, the observed time to failure of the system is the smallest of these failure times along with the actual cause of failure. Thus by letting $T = \min\{T_1, \dots, T_k\}$ and $C = c$ if $T = T_c$, we observe

the pair (T, C) and we are back to the formulation of the previous paragraph. Note that it is assumed that the T_c for which minimum is attained is uniquely given.

Uses of Competing Risks in Reliability and Maintenance Studies

Crowder [1, Chapter 1] gives some simple examples of the uses of competing risks in reliability studies. One of these is taken from King [2] who studied data of breaking strengths of certain wire connections. Two types of failure were defined (so $k = 2$): breakage at the bonded end and breakage along the wire itself.

Modern reliability data banks usually distinguish between a large number of failure modes, which suggests the use of methods from the theory of competing risks. Cooke [3, 4] reviews some main styles in the design of *reliability databases*, as well as models and methods for the analysis.

Traditionally, competing risks were analyzed as if they were independent of each other. This assumption appears to be dubious in many reliability applications, however. Even if assumptions of stochastic independence may often be justified by the physically independent functioning of components, a dependence between risks may be introduced by many sources, for example, *load sharing* between components, shared working environment, manufacture, and maintenance.

A simple case of dependent risks occur in the case when a potential component failure at some time T_1 may be avoided by a **preventive maintenance** (PM) at time T_2 (see [4]). The assumption that T_1 and T_2 are independent is clearly unreasonable in this application, since the maintenance crew is likely to have some information regarding the component’s state during operation, and this insight is used to perform maintenance with the aim of avoiding component failure. Note that the observable result is the pair (T, C) , rather than the latent times T_1 and T_2 themselves, which are usually the times of primary interest. For example, knowing the distribution of T_1 , the true failure time distribution, could be the basis for *maintenance optimization*. However, as will be discussed later, the marginal distributions of the T_j are not identifiable from the observation of (T, C) alone, unless specific assumptions are made on the dependence.

2 Competing Risks

Basic Literature on Competing Risks

The recent book by Crowder [1] gives a comprehensive review of the theory and methods of competing risks. An older book devoted to the subject is David and Moeschberger [5]. Several standard books on reliability and survival analysis contain chapters on competing risks, for example, Lawless [6], Kalbfleisch and Prentice [7], Nelson [8], Bedford and Cooke [9], and Andersen *et al.* [10]. Among several review papers written on the subject we mention Gail [11] and Moeschberger and Klein [12].

Model Specification

The joint distribution of the pair (T, C) from an individual is completely specified by the subdistribution functions

$$F_j(t) = P(T \leq t, C = j); \quad t > 0, \quad j = 1, \dots, k \quad (1)$$

The corresponding subdensity functions, when they exist, are given by differentiation, $f_j(t) = F_j'(t)$. In the formulas given in the following, the ranges of t and j are clear and will be mostly suppressed.

The marginal distribution of T is given by the *cumulative distribution function*

$$F(t) = P(T \leq t) = \sum_{j=1}^k F_j(t) \quad (2)$$

or by the survival function $\bar{F}(t) = 1 - F(t)$. Note that here and in the sequel, a bar above a capital letter means that this is the (sub)survival function corresponding to the (sub)distribution function given without a bar. The marginal distribution of C is given by $\pi_j = P(C = j) = F_j(\infty)$.

The distribution of (T, C) can alternatively be specified by the subhazard functions, which when they exist are given by

$$\begin{aligned} \lambda_j(t) &= \lim_{\Delta t \rightarrow 0} \frac{P(T \leq t + \Delta t, C = j | T > t)}{\Delta t} \\ &= \frac{f_j(t)}{\bar{F}(t)} \end{aligned} \quad (3)$$

It follows that $\lambda(t) = \sum_{j=1}^k \lambda_j(t)$ is the hazard function of T . Moreover, from equation (3) it follows

that $f_j(t) = \lambda_j(t) \bar{F}(t)$ which by integration gives the useful connection

$$F_j(t) = \int_0^t \lambda_j(u) \bar{F}(u) du. \quad (4)$$

Next, defining the cumulative subhazard functions as $\Lambda_j(t) = \int_0^t \lambda_j(u) du$ it is seen that $\Lambda(t) = \sum_{j=1}^k \Lambda_j(t)$ is the cumulative hazard function of T . Thus we have

$$\bar{F}(t) = e^{-\Lambda(t)} = e^{-\sum_{j=1}^k \Lambda_j(t)} = \prod_{j=1}^k \bar{G}_j^*(t) \quad (5)$$

where

$$\bar{G}_j^*(t) = e^{-\Lambda_j(t)} \quad (6)$$

Note that $\bar{G}_j^*(t)$ is a survival function (possibly with an atom at infinity), but that it is not in general the distribution of any observable random variable (see, however, section title “Assuming Independent Risks”).

The subhazard functions $\lambda_j(t)$ have the intuitive interpretation as the failure rate from a specific cause conditional on survival up to time t . They are also known under the names of mode-specific or cause-specific hazard functions, and have in older literature been called *crude hazard rates*.

Latent Failure Time Representation

Consider again the representation where $T = \min\{T_1, \dots, T_k\}$ and $C = c$ if $T = T_c$, with c assumed uniquely given. Let the joint survival function of $T_1, \dots, T_k$ be $\bar{K}(t_1, \dots, t_k) = P(T_1 > t_1, \dots, T_k > t_k)$. Then the survival function of T can be evaluated as $\bar{F}(t) = \bar{K}(t, t, \dots, t)$. The subdensity functions can also be calculated directly from the joint survival function as

$$f_j(t) = - \left(\frac{\partial \bar{K}(t_1, \dots, t_k)}{\partial t_j} \right)_{t_1 = \dots = t_k = t} \quad (7)$$

and it further follows that

$$\lambda_j(t) = \frac{f_j(t)}{\bar{F}(t)} = - \left(\frac{\partial \log \bar{K}(t_1, \dots, t_k)}{\partial t_j} \right)_{t_1 = \dots = t_k = t} \quad (8)$$

Let the marginal survival function of T_j be denoted $\bar{G}_j(t) = P(T_j > t)$ and let the corresponding hazard rate function be $h_j(t) = -\bar{G}'_j(t)/\bar{G}_j(t)$. It is noted that in general $h_j(t)$ and $\lambda_j(t)$ are different and have different interpretations. While the former has traditionally been called the *net rate*, the latter is the *crude rate* as mentioned before. It will be seen later, however, that $h_j(t) = \lambda_j(t)$ for all $t > 0$ when the T_j are independent.

The Identifiability Problem

As already mentioned, the main interest in a competing risks analysis is often in the joint and marginal distributions of the latent failure times $T_1, \dots, T_k$. The classical problem of competing risks is, however, that the distribution of the observable pair (T, C) does not in general determine the distribution of the latent failure times. This means that there are several different joint distributions of $T_1, \dots, T_k$ which give rise to the same distribution of (T, C) . This fact, called *nonidentifiability* of the distributions of $T_1, \dots, T_k$, was noted by Cox [13] for the case of two failure causes, while Tsiatis [14] studied the general case.

The main result of Tsiatis [14] states that if the set of subdistribution functions $F_j(t)$ is given for some model with dependent risks, then there exists a unique model with independent risks yielding the same $F_j(t)$. This model is defined by the joint survival function

$$\bar{K}(t_1, \dots, t_n) = \prod_{j=1}^k \bar{G}_j^*(t_j) \quad (9)$$

where the $\bar{G}_j^*(t)$ are given by equation (6). Thus, one cannot know, from observations of (T, C) alone, which of the two models is correct, since they will both fit the data equally well.

How to Deal with the Identifiability Problem

In this section we consider ways of overcoming the identifiability problem under the latent failure time representation of competing risks. It should be stressed that this can only be done by imposing additional assumptions in the model. These may be of various kinds, but one should always have in

mind that under observation of the pair (T, C) , the assumptions will always be nontestable.

Assuming Independent Risks

The classical assumption is that the risks act independently, so that the latent failure times T_j are independent. It then follows from Tsiatis [14] (see above) that we have identifiability of the distributions in question (under regularity assumptions), meaning that the marginal distributions of the T_j now can be computed from the subdistribution functions $F_j(t)$. In practice this means that the marginal distributions can be estimated in a consistent manner from competing risks data. Furthermore, in the case of independence we can write $\bar{K}(t_1, \dots, t_k) = \prod_{j=1}^k \bar{G}_j(t_j)$ and hence it follows from equation (9) that $\bar{G}_j^*(t) = \bar{G}_j(t)$ and from equation (8) that $\lambda_j(t) = h_j(t)$. Thus in the independent risks case the $\bar{G}_j^*(t)$ are in fact the marginal survival distributions of the T_j .

Zheng and Klein [15] generalized the result on identifiability in the independent risks case by proving that the marginal distributions are identifiable when the dependence of the T_j is given by a known copula.

Computing Bounds for the Marginal Survival Functions

Peterson [16] gave bounds for the marginal distribution functions $G_j(t) = P(T_j \leq t)$ in terms of the observable subdistribution functions F_j . The bounds are given by

$$F_j(t) \leq G_j(t) \leq F(t) \quad (10)$$

which are easily verified. Peterson [16] showed, moreover, that these bounds are pointwise sharp. They are not, however, functionally sharp. In fact, Crowder [1] found that the functions $G_j(t) - F_j(t)$ need to be nondecreasing in addition to being nonnegative. Subsequently, Bedford and Meilijson [17] obtained the complete characterization of the feasible marginal distribution functions G_j for a given set of subdistribution functions F_j . They showed in particular that the condition that the functions $G_j(t) - F_j(t)$ are nonnegative and nondecreasing, is also sufficient, provided a subtle additional measure theoretic assumption is satisfied. Note that when F_j and G_j are differentiable, this leads to the inequality

$$f_j(t) \leq g_j(t) \text{ for all } t > 0 \quad (11)$$

4 Competing Risks

where $g_j(t)$ is the marginal density function of T_j . We use this inequality in the following example.

Example (adapted from Bedford and Meilijson [17]) Consider the model with $k = 2$ given by constant subhazard functions $\lambda_j(t) = \lambda_j$, $j = 1, 2$. In this case $F_j(t) = (\lambda_j/\lambda_+) (1 - e^{-\lambda_+ t})$ and $F(t) = 1 - e^{-\lambda_+ t}$, where $\lambda_+ = \lambda_1 + \lambda_2$. Assume now that T_1 has an exponential marginal distribution, so that $G_1(t) = 1 - e^{-\lambda t}$ for some λ . The upper bound of inequality (10) easily gives $\lambda \leq \lambda_+$. Further, the inequality (11) with $j = 1$ gives $\lambda_1 e^{-\lambda_+ t} \leq \lambda e^{-\lambda t}$ for all $t > 0$, which implies $\lambda \geq \lambda_1$ by letting $t \rightarrow 0$. Thus we have shown that any feasible value of λ must satisfy the inequality

$$\lambda_1 \leq \lambda \leq \lambda_+ \quad (12)$$

Bedford and Meilijson [17] in fact showed that the set of possible values is the half-open interval $[\lambda_1, \lambda_1 + \lambda_2)$. Note that an assumption that T_1 and T_2 are independent leads to $\lambda = \lambda_1$. The example therefore shows that the independence assumption leads to the most optimistic value of the failure rate λ among the ones that are possible when the dependence is not specified. Williams and Lagakos [18] proved a corresponding result in a more general setting.

Parametric Identifiability

Note that the meaning of identifiability of marginal distributions as discussed above has been in the nonparametric sense that the marginal distribution functions G_j can be derived in terms of the subdistribution functions F_j . If a parametric model is specified for the latent failure times, then the identifiability problem is a completely different one since it now has to do with identification of a finite set of parameters. Crowder [1, Chapter 7.7] and Moeschberger and Klein [12] review models for which identifiability holds (see section title “Modeling of Latent Failure Times”).

Modeling of Competing Risks

Modeling of subdistributions

Mixture models. These are models given by specifying subdistribution functions of the form $F_j(t) =$

$\pi_j Q_j(t)$ for given (parametric) distribution functions $Q_j(t)$ and weights π_j with $\sum_{j=1}^k \pi_j = 1$. An example is the Weibull version where $\bar{Q}_j(t) = \exp\{-(t/\theta_j)^{\alpha_j}\}$.

Modeling subhazard functions. Another common approach is to assume parametric models for the subhazard functions, for example using Weibull hazards, $\lambda_j(t; \alpha_j, \theta_j) = (\alpha_j/\theta_j) (t/\theta_j)^{\alpha_j-1}$.

Regression models. Let $\mathbf{x}$ be a vector of covariates for the unit under study. Two main approaches for regression modeling in survival analysis are *proportional hazards models* and *accelerated life models*. The versions for competing risks can be given as follows:

Proportional hazards

The subhazard functions are given by $\lambda_j(t; \mathbf{x}) = \psi_j(\mathbf{x}) \lambda_{0j}(t)$ where $\psi_j(\mathbf{x})$ is a positive function of the covariates, for example, $\psi_j(\mathbf{x}) = \exp\{\boldsymbol{\beta}' \mathbf{x}\}$ for a parameter vector $\boldsymbol{\beta}$. If $\lambda_{0j}(t) = \lambda_0(t)$ does not depend on j , then T and C are stochastically independent.

Accelerated life model

The dependence of $\mathbf{x}$ is here through factors $\phi_j(\mathbf{x})$ which accelerate time in such a manner that $F_j(t; \mathbf{x}) = F_{0j}(\phi_j(\mathbf{x})t)$.

Modeling of Latent Failure Times

The traditional way of modeling dependent risks has been through specification of the joint distribution function $K(t_1, \dots, t_k)$ or the joint survival function $\bar{K}(t_1, \dots, t_k)$. Classical examples are the bivariate normal and Weibull distributions considered by Moeschberger [19] and for which there are no identifiability problems. Other simple examples are given by the following.

Example: bivariate exponential (Gumbel [20])

Let $k = 2$ and define

$$\bar{K}(t_1, t_2) = \exp\{-\lambda_1 t_1 - \lambda_2 t_2 - \nu t_1 t_2\} \quad (13)$$

Then $\lambda_j(t) = \lambda_j + \nu t$, so that the independent risks model giving the same subdistribution functions would have marginal survival functions $\bar{G}_j^*(t) =$

$\exp\{-\lambda_j t - \nu t^2/2\}$. However, the true marginal survival functions corresponding to the given $\bar{K}(t_1, t_2)$ are $\bar{G}_j(t) = \exp\{-\lambda_j t\}$. The dilemma is that it is not possible to distinguish between these two models from observation of (T, C) . Note, however, that for both models we have identifiability of all the parameters $\lambda_1, \lambda_2, \nu$.

Example: PM modeling In the section title “Uses of Competing Risks in Reliability and Maintenance Studies” we considered the case with $k = 2$ where T_1, T_2 are, respectively, the time of critical failure of a component and the potential time of a PM. Cooke [21, 4] introduced the notion of random signs censoring which is tailored for such cases. In our notation the PM time T_2 is called a *random signs* censoring of the failure time T_1 if the event $\{T_2 < T_1\}$ is stochastically independent of T_1 . The idea is that the component emits some kind of signal before failure, and that this signal is discovered with a probability which does not depend on the age of the component. The interesting fact is that random signs censoring implies identifiability of the distribution of T_1 . Lindqvist *et al.* [22] suggested a special case of random signs censoring obtained by introducing a so-called repair alert function which describes the “alertness” of the maintenance crew as a function of time.

Statistical Inference

Consider the case where we have data of the form (T, C) for n independent observation units. In practice such data are often right censored by some source independent of the k given risks. For example, this could be a censoring due to a time limit of the experiment (type I censoring). If the i th observation is noncensored, we observe both the lifetime t_i and the cause c_i . On the other hand, if the i th observation is right censored at time t_i , then we do not observe c_i and all we know is that $T > t_i$. If we let $\delta_i = 0$ if the i th observation is censored and $\delta_i = 1$ otherwise, we get the likelihood function

$$L = \prod_{i=1}^n f_{c_i}(t_i)^{\delta_i} \bar{F}(t_i)^{1-\delta_i} = \prod_{i=1}^n \lambda_{c_i}(t_i)^{\delta_i} \bar{F}(t_i) \quad (14)$$

where the last equality follows by using equation (3). Further, defining $\delta_{ij} = I(c_i = j)$ when $\delta_i = 1$ and

$\delta_{ij} = 0$ when $\delta_i = 0$ and using equation (5), we arrive at (see Lawless [6, Chapter 9.1])

$$L = \prod_{j=1}^k \left[\prod_{i=1}^n g_j^*(t_i)^{\delta_{ij}} \bar{G}_j^*(t_i)^{1-\delta_{ij}} \right] \equiv \prod_{j=1}^k L_j \quad (15)$$

where $g_j^*(t) = \lambda_j(t) \bar{G}_j^*(t)$ is the density function corresponding to $\bar{G}_j^*(t)$.

Thus we can write L as a product of the likelihoods L_j , where L_j has the same form as the likelihood function of a censored sample associated with the j th failure cause where all observations where C is not observed to equal j are treated as censorings.

Any of the above expressions for the likelihood L may be used for parametric **maximum likelihood estimation**, depending on the way the model is defined. Equation (15) has important implications for nonparametric estimation. In fact, the likelihood L_j is of the same form as the one leading to the usual Kaplan–Meier estimator (see [10]) of a survival function under independent right censoring. This means, in particular, that the “artificial” survival functions $\bar{G}_j^*(t)$ and hence $\Lambda_j(t) = -\log \bar{G}_j^*(t)$, can be estimated nonparametrically by the Kaplan–Meier estimator. Alternatively, one may from the same reasoning estimate the Λ_j for each j using the Nelson–Aalen estimator [10]. In this case, denoting the resulting estimator by $\hat{\Lambda}_j(t)$, it follows from equation (4) that one may estimate the F_j nonparametrically by

$$\hat{F}_j(t) = \int_0^t \hat{\bar{F}}(u) d\hat{\Lambda}_j(u), \quad j = 1, \dots, k \quad (16)$$

where $\hat{\bar{F}}(t)$ is the ordinary Kaplan–Meier estimate of the survival function $\bar{F}(t)$ of T . Note that this estimator uses the censored data obtained by collapsing the k failure causes into one single cause.

Beyond Classical Competing Risks

Markov process models Aalen [23] modeled the classical competing risks problem as a continuous-time *Markov process* with one working state (“0”) and k absorbing failure states corresponding to the k risks. More generally one may consider **Markov processes** with more complex state spaces and again k absorbing states corresponding to the k failure types. This leads to the consideration of multivariate phase

type distributions, see Assaf *et al.* [24]. Hokstad and Frøvig [25] considered failure models for periodically tested items where the failures are either degraded or critical. Whitmore [26] discussed the use of first-passage-time distributions connected with multidimensional *Brownian motion* as models for competing risks.

Competing risks in repairable systems Consider a system where failures are classified in k different types. Until now we have considered the case of non-repairable units observed until failure or PM. Now assume that the unit is repaired after failure, then is put into operation again, and that the process continues in this way (see **Repairable Systems: Statistical Inference**). Suppose that, at each failure, the type of event is recorded. Assume also that time durations of repair and maintenance can be disregarded, so that the system is always restarted immediately after failure or maintenance action. This leads to the observation of a marked point process $(S_1, C_1), (S_2, C_2), \dots$ with successive failure times $0 < S_1 < S_2 < \dots$ and marks C_i in $\{1, 2, \dots, k\}$. The properties of such a process depends on the repair strategy. For example, a perfect repair (renewal) of the system at failures leads to independent and identically distributed pairs (T_i, C_i) where $T_i = S_i - S_{i-1}$ ($i = 1, 2, \dots$) and $S_0 = 0$. We are hence back to the basic case considered in this article. In general, however, the situation could be more complicated. Doyen and Gaudoin [27] present a point process approach for modeling of **imperfect repair** in competing risks situations between failure and PM. Bedford and Lindqvist [28] considered a series system of k repairable components where only the failing component is repaired at failures. A general setup for this kind of processes is suggested by Lindqvist [29].

Discussion

As argued in the article, the latent failure time approach to competing risks is particularly useful in reliability applications. However, as risks will usually be dependent, the identifiability problem occurs. Recall that although additional assumptions on the dependence may lead to identifiability, these assumptions are nontestable from data of the form (T, C) . Prentice *et al.* [30] therefore rejected the latent failure time approach and instead advocated the use of the observable subhazard functions in analyses of

competing risks. One may argue against this, however, that in practical applications it often makes good sense to use ones information and prior beliefs in order to model the underlying probability mechanisms beyond what is actually observable.

On the other hand, an erroneous assumption of independent risks may lead to seriously misleading conclusions. As noted by Cooke [3], such assumptions are often made in the assessments of competing failure rates in reliability databases. One then assumes that for each of the k failure causes, failures occur as independent *Poisson processes*. This implies, however, that the rate of occurrence of each of the competing risks would be unaffected by removing the others. For competing risks corresponding to critical failure and PM this means that the rate of occurrence of critical failures would be unaffected by stopping PM activity, an assumption which is completely unreasonable. The appropriate method would be to invoke a more careful modeling by competing risks, using if possible all available information. More sophisticated methods for analysis of competing risks in connection with repairable systems, as briefly mentioned in the previous section, may be needed here.

References

- [1] Crowder, M.J. (2001). *Classical Competing Risks*, Chapman & Hall/CRC, Boca Raton.
- [2] King, J.R. (1971). *Probability Charts for Decision Making*, Industrial Press, New York.
- [3] Cooke, R.M. (1996a). The design of reliability databases, Part I, *Reliability Engineering and System Safety* **51**, 137–146.
- [4] Cooke, R.M. (1996b). The design of reliability databases, Part II, *Reliability Engineering and System Safety* **51**, 209–223.
- [5] David, H.A. & Moeschberger, M.L. (1978). *The Theory of Competing Risks*, Griffin, London.
- [6] Lawless, J.F. (2003). *Statistical Models and Methods for Lifetime Data*, 2nd Edition, Wiley-Interscience, Hoboken.
- [7] Kalbfleisch, J.D. & Prentice, R.L. (1980). *The Statistical Analysis of Failure Time Data*, John Wiley & Sons, New York.
- [8] Nelson, W. (1982). *Applied Life Data Analysis*, John Wiley & Sons, New York.
- [9] Bedford, T. & Cooke, R.M. (2001). *Probabilistic Risk Analysis: Foundations and Methods*, Cambridge University Press, Cambridge.
- [10] Andersen, P., Borgan, O., Gill, R. & Keiding, N. (1992). *Statistical Models Based on Counting Processes*, Springer, New York.

- [11] Gail, M. (1975). A review and critique of some models used in competing risks analysis, *Biometrics* **31**, 209–222.
- [12] Moeschberger, M.L. & Klein, J.P. (1995). Statistical methods for dependent competing risks, *Lifetime Data Analysis* **1**, 195–204.
- [13] Cox, D.R. (1959). The analysis of exponentially distributed lifetimes with two types of failure, *Journal of Royal Statistical Society Series B* **21**, 411–421.
- [14] Tsiatis, A. (1975). A nonidentifiability aspect of the problem of competing risks, *Proceedings of the National Academy of Sciences of the United States of America* **72**, 20–22.
- [15] Zheng, M. & Klein, J.P. (1995). Estimates of marginal survival for dependent competing risks based on an assumed copula, *Biometrika* **82**, 127–138.
- [16] Peterson, A.V. (1976). Bounds for a joint distribution function with fixed sub-distribution functions: application to competing risks, *Proceedings of the National Academy of Sciences of the United States of America* **73**, 11–13.
- [17] Bedford, T. & Meilijson, I. (1997). A characterization of marginal distributions of (possibly dependent) lifetime variables which right censor each other, *Annals of Statistics* **25**, 1622–1645.
- [18] Williams, J.S. & Lagakos, S.W. (1977). Models for censored survival analysis: constant- sum and variable-sum models, *Biometrika* **64**, 215–224.
- [19] Moeschberger, M.L. (1974). Life tests under dependent competing causes of failure, *Technometrics* **16**, 39–47.
- [20] Gumbel, E.J. (1960). Bivariate exponential distributions, *Journal of the American Statistical Association* **55**, 698–707.
- [21] Cooke, R.M. (1993). The total time on test statistics and age-dependent censoring, *Statistics and Probability Letters* **18**, 307–312.
- [22] Lindqvist, B.H., Støve, B. & Langseth, H. (2006). Modelling of dependence between critical failure and preventive maintenance: the repair alert model, *Journal of Statistical Planning and Inference* **136**, 1701–1717.
- [23] Aalen, O.O. (1976). Nonparametric inference in connection with multiple decrement models, *Scandinavian Journal of Statistics* **3**, 15–27.
- [24] Assaf, D., Langberg, N.A., Savits, T.H. & Shaked, M. (1984). Multivariate phase-type distributions, *Operations Research* **32**, 688–702.
- [25] Hokstad, P. & Frøvig, A.T. (1996). The modelling of degraded and critical failures for components with dormant failures, *Reliability Engineering and System Safety* **51**, 189–199.
- [26] Whitmore, G.A. (1986). First-passage time models for duration data: regression structures and competing risks, *The Statistician* **35**, 207–219.
- [27] Doyen, L. & Gaudoin, O. (2006). Imperfect maintenance in a generalized competing risks framework, *Journal of Applied Probability* **43**, 825–839.
- [28] Bedford, T. & Lindqvist, B.H. (2004). The identifiability problem for repairable systems subject to competing risks, *Advances in Applied Probability* **36**, 774–790.
- [29] Lindqvist, B.H. (2006). On the statistical modelling and analysis of repairable systems, *Statistical Science* **21**, 532–551.
- [30] Prentice, R.L., Kalbfleisch, J.D., Peterson, A.V., Flournoy, N., Farewell, V.T. & Breslow, N.E. (1978). The analysis of failure times in the presence of competing risks, *Biometrics* **34**, 541–554.

Related Articles

Accelerated Life Tests: Analysis with Competing Failure Modes; Masked Failure Data: Competing Risks; Masked Failure Data.

BO HENRY LINDQVIST

Subjective Life Models

Introduction

Subjective life models form a part of the decision-theoretic approach to equipment maintenance. In decision theory, the optimal maintenance policy is the policy that maximizes the decision maker's expected utility. This requires a utility function expressing the benefits and costs and a probability distribution expressing the decision maker's beliefs.

Models arise when the distribution is specified in two assessments. The first assessment is a general assumption concerning the nature of the aging process. This leads to a parametric model for the lifetimes of the items. The second assessment concerns the parameters themselves and leads to a Bayesian framework for inference.

The first five sections lay the groundwork for the models by defining **maintenance policies** that formalize the optimization procedure. It shows how the concepts of hazard and their differentials replace the classical failure rates and hazard gradients (see Barlow and Proschan [1] for the classical definitions). The remaining sections derive subjective life models from first principles and show under what conditions we can expect to obtain the classical failure models.

Lifetimes

Consider a sequence of N components with lifetimes $\mathbf{X} = X_1, \dots, X_N$ having outcomes $\mathbf{x} = x_1, \dots, x_N$. It will not be necessary to consider an infinite sequence from the outset because the distribution on this sequence is subjective and need not be estimated. To show the connections with the classical theory, we will consider limits when $N \rightarrow \infty$, but in practice there are many problems concerning even a small number of components and this is a definite advantage.

The sequence $\mathbf{X}$ is an element of $\mathbb{R}_+^N$ and so the beliefs are a subjective probability distribution P on this sample space. It is important to see that this sample space is not a Euclidean N -space. There is no metric defining distance or angles and so it immediately follows that Marshall's hazard gradient (see Marshall [2]) cannot be defined since a gradient gives the *direction* of greatest

increase and the concept of direction cannot be defined.

However, lifetime space has structure that Euclidean space does not have and this can be used to define the survival distribution and the hazard, *viz.*

$$\bar{F}(\mathbf{X}) = P(\mathbf{X} \geq \mathbf{x}) \quad \text{and} \quad H(\mathbf{X}) = -\log \bar{F}(\mathbf{X}) \quad (1)$$

Unlike coordinates in Euclidean space, items are physical entities and they endow the lifetime space with N physical coordinates at any point. Such a space is also known as a *fiber bundle*. The survival distribution is useful for a subjectivist because they depend only on the fibers and are invariant under changes of units for measuring lifetimes. It would not be useful for a subjectivist in Euclidean space since it is not invariant under its invariance group, i.e., rotations and translations.

Maintenance Policies

A maintenance policy is a path $\mu : \mathbb{R} \rightarrow \mathbb{R}_+^N$ mapping each value t of time into a sequence of ages for the items. For instance, a replace-at-failure policy would start at $\mathbf{0}$, first run to $x_1, 0, \dots, 0$, then to $x_1, x_2, 0, \dots, 0$ and so on until $x_1, \dots, x_N$. In general, the path would run from $\mathbf{0}$ (unless the items are used) in any positive direction (unless the items rejuvenate). Time is the parameter on this path. This is a very general mathematical representation and covers any type of operating policy as well.

The tangent vector to the maintenance policy at any point is the aging rate:

$$\vec{a} = \sum_{i=1}^N \dot{x}_i \frac{\partial}{\partial x_i} \quad (2)$$

Here $\partial/\partial x_i$ is the unit vector in the i direction, i.e., a unit of aging of the i th item. The aging rate specifies the rate at which each item is being used. Typically, we assume that one of the $\dot{x}_i$'s is unity and the rest zero, and so measure the lifetime of the current item in time t , but it is significant that we do not have to do this and we can accelerate the aging. Think of the aging-rate vector as an engineer or a physicist would: as arrow tangent to the path that exists independent of its coordinate representation

2 Subjective Life Models

in lifetime space $(\dot{x}_1, \dots, \dot{x}_N)$ which depends on the units with which we measure lifetime.

Hazard Differential and Failure Rate

A maintenance policy together with the hazard provides a function $H \circ \mu : \mathbb{R} \rightarrow \mathbb{R}$ that maps each value of time into a hazard. The time derivative of this composite function is the failure rate

$$r(t) = \dot{H} = \frac{d(\mu \circ H)}{dt} = \sum_{i=1}^N \dot{x}_i \frac{\partial H}{\partial x_i} \quad (3)$$

Notice that in the subjective framework, the failure rate depends on the maintenance policy and we have used Newton's fluxion notation to bear this out. If the items are used differently, the failure rate will change.

To find a differential quantity that is independent of the path, we need to consider the differential of the hazard function

$$\mathbf{h} = dH = \sum_{i=1}^N \frac{\partial H}{\partial x_i} dx_i = \sum_{i=1}^N H'_i dx_i \quad (4)$$

Here we are using Leibniz's derivative notation for the partial hazards. The hazard differential is a 1-form which we will call the hazard differential and it is the correct replacement for Marshall's hazard gradient. As a 1-form, it acts naturally on the aging vector to produce the failure rate

$$r = \mathbf{h} \cdot \vec{a} = \sum_{i=1}^N H'_i \dot{x}_i \quad (5)$$

Here $\cdot$ is the tensor saturation operator.

The failure rate of the i th item can be defined by considering an aging-rate $\vec{1}_i$ vector that is 0 on all except the $\dot{x}_i$ component:

$$r_i(t|\mathbf{x}) = \mathbf{h} \cdot \vec{1}_i = H'_i(t|\mathbf{x}) \quad (6)$$

This corresponds most closely to the notion of a failure rate in the classical theory. However, notice that the failure rate depends where on the path, i.e., at which $\mathbf{x} = \mu(t)$ the aging rate is anchored, which is emphasized here by conditioning on $\mathbf{x}$; knowing that items have survived until $\mathbf{x}$ changes the decision maker's beliefs concerning the failure rate.

Geometrically, think of the hazard as a mountain on the lifetime space. As you move along the maintenance path, you climb up the hazard mountain. Locally the mountain is approximated by the form $\mathbf{h}$ in the same way the path is approximated by the vector $\vec{a}$ and the result of a small movement is the failure rate. The hazard differential can be pictured geometrically by a set of straightened-out level lines of the hazard function. Its coordinate representation is the row $(H'_1, \dots, H'_N)^T$ consisting of the partial hazards and the failure rate can be calculated as the inner-product-like contraction of that with the coordinate representation of the aging vector as in equation (3). This inner product does not depend on the particular coordinates and the failure rate emerges as an invariant. This is also expressed by saying that the coordinate representation of the hazard differential is contravariant and that of the failure rate is covariant.

Optimal Policies

A maintenance policy maps every value of the lifetimes into a consequence that includes all the benefits and costs. The utility of such a consequence is measured by a function $U : \mathbb{R}_+^N \rightarrow \mathbb{R}$. The optimal policy is then the policy whose expected utility

$$EU = \int_{\mathbb{R}_+^N} U(\mathbf{x}) P(d\mathbf{x}) \quad (7)$$

is higher than all other available policies. The optimal policy is not necessarily unique, but the decision maker is indifferent which of the optimal policies is followed.

This optimization is easier to do incrementally. Differentiate the expression for the expected utility with respect to time to obtain the flow of expected utility at any time t :

$$\dot{EU}(t) = \int_{\mathbb{R}_+^N} \dot{U}(\mathbf{x}) r(t|\mathbf{x}) d\mathbf{x} \quad (8)$$

This expression highlights the importance of the failure rate. The expected utility flow is a weighted average of the utility flows, and the weights are represented by the failure rate.

The picture that emerges is that of gathering utility per unit of time as we move along the path of the

maintenance policy. The expected utility of the policy is found by integrating along the path μ

$$EU = \int_{\mu} \dot{E}U(t) dt \quad (9)$$

In nonsubjective approaches, concepts such as “reliability” and “availability” are used to judge maintenance policies. Here, the utility is the place to encapsulate any such requirements. The utility flow $\dot{U}$ should be formulated to reflect any benefit for uptime, cost for downtime, or any other benefits and costs associated with the operation of the equipment and purchasing and replacing the parts. If the utility flow is correctly formulated, no other requirements are needed. While the utility of uptimes and downtimes may be modeled this way, there are elements such as replacement costs that typically occur instantly, at least theoretically, and they should be added on during the integration.

Lifetime Models

In the classical theory, a parametric probability model such as the exponential or Weibull model is assumed with the parameters estimated. The first problem in applying this approach here is that the assumption of the model needs to be justified on a subjective basis, i.e., motivated from first principles so that the decision maker can corroborate that the model corresponds to his beliefs. The second problem concerns the parameters. Since the number of items is allowed to be very small, standard **estimation** procedures do not apply. Therefore, the parameters need to be defined in such a way that they relate to the finiteness of the sequence.

In the following subsections, we will describe the procedure for establishing subjective models from first principles, for several popular life models. We then show how the classical equivalents emerge as limiting cases for $N \rightarrow \infty$.

No-Aging Property and the Exponential Model

Consider a maintenance policy at time t and suppose none of the items has died, having lived until $\mathbf{x} = \mu(t)$. For the next interval $h > 0$ of operation, the decision maker can employ any of the N items. If the decision maker considers that the chances of the items dying in this interval are equal for all $\mathbf{x}$ and

h , this decision maker considers that these items do not age. Specifically, the no-aging property requires that

$$\bar{F}(x_i + h | \mathbf{X} \geq \mathbf{x}) = \bar{F}(x_j + h | \mathbf{X} \geq \mathbf{x}) \quad (10)$$

for all $h, \mathbf{x}$ and $i, j \in \{1, \dots, N\}$. The items are all the same, no matter how much they are used.

The no-aging property is a first principle from which we can derive a parametric model. Let the average lifetime be

$$\theta = \frac{1}{N} \sum_{i=1}^N x_i \quad (11)$$

Then it follows that a subsequence $X_1, \dots, X_n$ of $\mathbf{X}$ has a survival distribution that can be written as a mixture as follows:

$$\bar{F}(\mathbf{x}_n) = \int_0^\infty \left[1 - \frac{\sum_{i=1}^n x_i}{N\theta} \right]^{N-1} P(d\theta) \quad (12)$$

This is the class of Schur-constant distributions. The mixture represents a parametric model for a finite population; the parameter is a function of only the N lifetimes. The model is a Bayesian model in the sense that, if a subsequence of lifetimes has been observed, the decision maker can employ a posterior distribution for the parameter θ that is calculated using Bayes' theorem as follows

$$P(d\theta | \mathbf{x}_n) \propto \frac{1}{\theta^n} \left[1 - \frac{\sum_{i=1}^n x_i}{N\theta} \right]^{N-1-n} P(d\theta) \quad (13)$$

rather than using classical estimators.

When we require the no-aging property to hold for all N , no matter how large, the no-aging property converges to the classical memoryless property. This is perhaps best seen by letting $N \rightarrow \infty$ in the above expression for the survival distribution to obtain a mixture of independent and identically distributed (i.i.d.) exponentials:

$$\bar{F}(\mathbf{x}_n) = \int_0^\infty \prod_{i=1}^n e^{-x_i/\theta} P(d\theta) \quad (14)$$

4 Subjective Life Models

The parameter can now be interpreted as the limiting average lifetime, and Bayesian procedures can still be used to infer the posterior distribution as before, although the likelihood is now the exponential.

Notice that in the subjective approach, the finite population models are the natural models to consider. This is an advantage since oftentimes there are only a small number of items. These models are not the usual i.i.d. models since, even conditional on the average lifetime, they cannot be independent unless the number of items increases indefinitely. Also notice that the parameter θ emerges as a rigorously defined random variable, i.e., as a function on the sample space $\mathbb{R}_N^+$. This gives it an immediate physical meaning, obviates the need for a separate parameter space, and makes a Bayesian approach to inference unavoidable.

Generalized Weibull Model

The no-aging property forms the basis for a large class of models that can be characterized as generalized Weibull models. Although the no-aging property (equation 10) might not apply to the lifetimes directly, the decision maker may judge that there is some continuous increasing function ψ of the lifetimes such that the no-aging property applies to the transformed values instead. The power function

$$\psi(x_i) = x_i^\alpha \quad (15)$$

with $\alpha, \beta > 0$ is an example. In particular, when items to which the no-aging property applies are accelerated at a rate ψ^{-1} , then the no-aging property will apply to $\psi(x_i)$.

The parametric form for the distribution is easy to find:

$$\bar{F}(\mathbf{x}_n) = \int_0^\infty \left[1 - \frac{\sum_{i=1}^n \psi(x_i)}{N\theta} \right]^{N-1} P(d\theta) \quad (16)$$

where the parameter now is

$$\theta = \frac{1}{N} \sum_{i=1}^N \psi(x_i) \quad (17)$$

The limiting form is

$$\bar{F}(\mathbf{x}_n) = \int_0^\infty \prod_{i=1}^n e^{-\psi(x_i)/\theta} P(d\theta) \quad (18)$$

Substituting x_i^α for $\psi(x_i)$ gives the familiar i.i.d. Weibull model.

Notice that the parameter α , or, in general, any of the parameters describing the functional form of ψ , is different in nature from the parameter θ . This distinction does not appear in nonsubjective approaches. Here θ is a statistical parameter for which the decision maker has a probability distribution, since it is a function on the sample space, and Bayesian methods can be used to estimate it. However, α is not a function of the sample space, the decision maker has no probability distribution for it, and it makes no sense to attempt to estimate it unless we enlarge the sample space in some manner so that it can be defined as a function. Instead, α describes the acceleration of the items or some other physical law that transforms lifetimes, so that the no-aging property applies.

Aging, Schur Convexity, and Increasing Failure Rate (IFR) Models

Univariate life distributions representing aging items have been described as having increasing failure rate and the corresponding distributions as IFR. The problem with the IFR concept from a subjective point of view is that it is not preserved under mixing, i.e., if $\bar{F}(x|\theta)$ is IFR then $\bar{F}$ is not necessarily IFR and this latter survival function is what is used for decision making.

The appropriate subjective definition of aging was first given in Barlow and Mendel [3]: items are considered to exhibit the aging property if

$$\bar{F}(x_i + h|\mathbf{X} \geq \mathbf{x}) \geq \bar{F}(x_j + h|\mathbf{X} \geq \mathbf{x}) \quad (19)$$

for all $h > 0$ and $\mathbf{x}$ whenever $x_i < x_j$. The joint distribution for N items represents aging components if and only if the joint survival distribution is Schur concave. Also, any mixture of IFR distribution is Schur concave, and so Schur concavity is the appropriate generalization of the IFR property.

For the failure rate itself, we have the following. Let $x_{[1]} \leq x_{[2]} \leq \dots \leq x_{[N]}$ denote an increasing rearrangement of the lifetimes. The failure $r_{[i]}(\mathbf{x})$ rate of

the $[i]$ th component is found by saturating the hazard differential by the unit aging vector in the $[i]$ th direction. The aging property is then found to be equivalent to

$$r_{[1]} \leq r_{[2]} \leq \cdots \leq r_{[N]} \quad (20)$$

The subjectivist failure rates are indeed increasing.

Summary

The subjective approach sees lifetime modeling in the context of optimal decision making. A utility flow expresses all the benefits and costs associated with the various outcomes under a maintenance policy and the optimum policy is the one exhibiting the highest expected utility.

The expectation is taken over an unconditional joint distribution. A decision maker has to make some assessment about this distribution and this is made easier through the concepts of no-aging and aging. These concepts are generalizations of classical concepts such as the memoryless property and the IFR property.

In this way, the subjective life models provide a consistent and general framework for reliability analysis.

References

- [1] Barlow, R.E. & Proschan, F. (1965). *Mathematical Theory of Reliability*, John Wiley & Sons, New York.
- [2] Marshall, A.W. (1975). Some comments on the hazard gradient, *Stochastic Processes and their Applications* **3**, 293–300.
- [3] Barlow, R.E. & Mendel, M.B. (1992). De Finetti-type representations for life distributions, *Journal of the American Statistical Association*, **87**(420), 1116–1122.

Further Reading

Proschan, F. (1975). Applications of majorization and schur functions in reliability and life testing, *Reliability and Fault Tree Analysis*, SIAM, pp. 237–258.

MAX MENDEL

Stochastic Orders and Aging

Introduction

Many reliability systems or components of such systems have random lifetimes that indicate aging properties of these systems or components. Most of these aging properties can be characterized by various stochastic orders. Such characterizations are useful for the identification of the corresponding aging properties and for a better understanding of the meaning of these properties. In this article we describe such characterizations. The characterizations can be used in practice to develop various inequalities that yield bounds on various probabilistic quantities of interest. In this article “increasing” and “decreasing” stand for “nondecreasing” and “nonincreasing”, respectively. Expectations are assumed to exist whenever they are mentioned.

Some Stochastic Orders

In this section we give the definitions of some stochastic orders that will be used in the sequel. Useful references that cover the area of stochastic orders are Shaked and Shanthikumar [1] and Müller and Stoyan [2]. Let X and Y be two absolutely continuous nonnegative random variables with distribution functions F and G , survival functions $\bar{F} \equiv 1 - F$ and $\bar{G} \equiv 1 - G$, and density functions f and g . For the sake of simplicity we assume that the supports of X and Y are $(0, \infty)$. Let F^{-1} and G^{-1} denote the quantile functions of X and Y . The random variable X is said to be smaller than the random variable Y in the

- *Ordinary stochastic order* (denoted as $X \leq_{\text{st}} Y$) if $\bar{F}(x) \leq \bar{G}(x)$ for all x .
- *Hazard rate order* (denoted as $X \leq_{\text{hr}} Y$) if $\bar{F}(x)\bar{G}(y) \geq \bar{F}(y)\bar{G}(x)$ for all $x \leq y$.
- *Mean residual life order* (denoted as $X \leq_{\text{mrl}} Y$) if $\bar{G}(x) \int_x^\infty \bar{F}(y) dy \leq \bar{F}(x) \int_x^\infty \bar{G}(y) dy$ for all $x \geq 0$.
- *Harmonic mean residual life order* (denoted as $X \leq_{\text{hmrl}} Y$) if $(\int_x^\infty \bar{F}(y) dy)/EX \leq (\int_x^\infty \bar{G}(y) dy)/EY$ for all $x \geq 0$.
- *Dispersive order* (denoted as $X \leq_{\text{disp}} Y$) if $G^{-1}(\alpha) - F^{-1}(\alpha)$ is increasing in $\alpha \in (0, 1)$.

- *Excess wealth order* (denoted as $X \leq_{\text{ew}} Y$) if $\int_{F^{-1}(\alpha)}^\infty \bar{F}(x) dx \leq \int_{G^{-1}(\alpha)}^\infty \bar{G}(x) dx$ for all $\alpha \in (0, 1)$.
- *Increasing convex order* (denoted as $X \leq_{\text{icx}} Y$) if $E[\phi(X)] \leq E[\phi(Y)]$ for all increasing convex functions ϕ for which the expectations are defined.
- *Increasing concave order* (denoted as $X \leq_{\text{icv}} Y$) if $E[\phi(X)] \leq E[\phi(Y)]$ for all increasing concave functions ϕ for which the expectations are defined.
- *Dilation order* (denoted as $X \leq_{\text{dil}} Y$) if $X - EX \leq_{\text{icx}} Y - EY$.
- *Convex transform order* (denoted as $X \leq_c Y$) if $G^{-1}F(x)$ is convex in $x \geq 0$.
- *Star order* (denoted as $X \leq_* Y$) if $G^{-1}F(x)$ is starshaped in x ; that is, if $G^{-1}F(x)/x$ increases for $x \geq 0$.
- *Superadditive order* (denoted as $X \leq_{\text{su}} Y$) if $G^{-1}F(x)$ is superadditive in x ; that is, if $G^{-1}F(x+y) \geq G^{-1}F(x) + G^{-1}F(y)$ for all $x \geq 0$ and $y \geq 0$.
- *Laplace transform order* (denoted as $X \leq_{\text{Lt}} Y$) if $E[\exp\{-sX\}] \geq E[\exp\{-sY\}]$ for all $s > 0$.

Characterizations of Aging Properties

Throughout this section, X denotes a random lifetime, that is, a nonnegative random variable. For simplicity of exposition we assume that X is absolutely continuous with distribution function F , survival function $\bar{F} \equiv 1 - F$, and density function f . Unless stated otherwise, it will also be assumed that the support of X is $(0, \infty)$. For any event A , we use the notation $[X|A]$ to denote any random variable that is distributed according to the conditional distribution of X given A . The results below can be found in Shaked and Shanthikumar [1], unless stated otherwise.

Most of the characterizations in this section are based on stochastic comparisons of residual lives at different times. Another kind of characterization that we describe is based on comparison of the “aging” random variable with an exponential random variable, which in the context of reliability theory, has the “nonaging” property.

Increasing and Decreasing Failure Rate (IFR and DFR)

The random variable X is said to have the aging property of increasing failure rate (IFR) if $\bar{F}$ is

2 Stochastic Orders and Aging

logconcave. It has the property of decreasing failure rate (DFR) if $\bar{F}$ is logconvex. The ordinary stochastic order can be used to characterize the IFR and DFR notions as follows.

Theorem 1 *The random variable X is IFR [DFR] if, and only if, $[X - t | X > t] \geq_{st} [\leq_{st}] [X - t' | X > t']$ whenever $t \leq t'$.*

The hazard rate order can be used to characterize the IFR and DFR notions as follows.

Theorem 2 *The random variable X is IFR [DFR] if, and only if, one of the following equivalent conditions holds (when the support of X is bounded, condition 3 does not have a simple DFR analog):*

1. $[X - t | X > t] \geq_{hr} [\leq_{hr}] [X - t' | X > t']$ whenever $t \leq t'$.
2. $X \geq_{hr} [\leq_{hr}] [X - t | X > t]$ for all $t \geq 0$.
3. $X + t \leq_{hr} X + t'$ whenever $t \leq t'$.

Similarly, the dispersive order can characterize the IFR and DFR notions as follows (see Belzunce *et al.* [3] and Pellerey and Shaked [4]).

Theorem 3 *The random variable X is IFR [DFR] if, and only if, one of the following equivalent conditions holds:*

1. $[X - t | X > t] \geq_{disp} [\leq_{disp}] [X - t' | X > t']$ whenever $t \leq t'$.
2. $X \geq_{disp} [\leq_{disp}] [X - t | X > t]$ for all $t \geq 0$.

Another characterization of the IFR and DFR notions, given in Belzunce *et al.* [5], is the following.

Theorem 4 *The random variable X is IFR [DFR] if, and only if, $[X - t | X > t] \geq_{icv} [\leq_{icv}] [X - t' | X > t']$ whenever $t \leq t'$.*

In some places in the literature, if the condition of Theorem 4 holds, that is, if $[X - t | X > t] \geq_{icv} [\leq_{icv}] [X - t' | X > t']$ whenever $t \leq t'$, then X is said to have the property of IFR [DFR] in the second order stochastic dominance (IFR(2) [DFR(2)]). Theorem 4 shows that the IFR(2) [DFR(2)] property is equivalent to the IFR [DFR] property.

Still another characterization, similar to the characterizations in Theorems 1–4, is given in the next theorem (see Belzunce *et al.* [6]).

Theorem 5 *The random variable X is IFR if, and only if, $[X - t | X > t] \geq_{Lt} [X - t' | X > t']$ whenever $t \leq t'$.*

The convex transform order characterizes the IFR notion as follows. We denote any exponential random variable by Exp (no matter what the mean is).

Theorem 6 *The random variable X is IFR if, and only if, $X \leq_c \text{Exp}$.*

Increasing Failure Rate Average (IFRA)

The random variable X is said to have the aging property of increasing failure rate average (IFRA) if $-\log \bar{F}$ is starshaped; that is, if $-\log \bar{F}(t)/t$ is increasing in $t \geq 0$. The star order can be used to characterize the IFRA notion as follows. As in Theorem 6, Exp denotes any exponential random variable.

Theorem 7 *The random variable X is IFRA if, and only if, $X \leq_* \text{Exp}$.*

New Better (Worse) than Used (NBU and NWU)

The random variable X is said to have the aging property of new better than used (NBU) if $\bar{F}(s)\bar{F}(t) \geq \bar{F}(s+t)$ for all $s \geq 0$ and $t \geq 0$. It has the property of new worse than used (NWU) if $\bar{F}(s)\bar{F}(t) \leq \bar{F}(s+t)$ for all $s \geq 0$ and $t \geq 0$. The ordinary stochastic order can be used to characterize the NBU and NWU notions as follows.

Theorem 8 *The random variable X is NBU [NWU] if, and only if, $X \geq_{st} [\leq_{st}] [X - t | X > t]$ for all $t > 0$.*

Similar to Theorems 6 and 7, the superadditive order can be used to characterize the NBU notion as follows.

Theorem 9 *The random variable X is NBU if, and only if, $X \leq_{su} \text{Exp}$.*

The order $\leq_{icx}$ was used in Cao and Wang [7] to define the property of new better than used in convex order (NBUC) as follows. The random variable X is said to have the aging property of NBUC if $X \geq_{icx} [X - t | X > t]$ for all $t > 0$. The order $\leq_{icv}$ was used in Deshpande *et al.* [8] to define the property of NBU in second degree stochastic dominance (NBU(2)) as follows. The random variable X is said to have the

aging property of NBU(2) if $X \geq_{\text{icv}} [X - t | X > t]$ for all $t > 0$.

Decreasing and Increasing Mean Residual Life (DMRL and IMRL)

The nonnegative random variable X with finite positive mean is said to have the aging property of decreasing mean residual life (DMRL) if $E[X - t | X > t]$ is decreasing for $t \geq 0$. It has the property of increasing mean residual life (IMRL) if $E[X - t | X > t]$ is increasing for $t \geq 0$. The mean residual life order can be used to characterize the DMRL notion as follows.

Theorem 10 *The random variable X is DMRL if, and only if, one of the following equivalent conditions holds:*

1. $[X - t | X > t] \geq_{\text{mrl}} [X - t' | X > t']$ whenever $t \leq t'$.
2. $X \geq_{\text{mrl}} [X - t | X > t]$ for all $t \geq 0$.
3. $X + t \leq_{\text{mrl}} X + t'$ whenever $t \leq t'$.

In a similar manner the order $\leq_{\text{hmrl}}$ can be used to characterize the DMRL notion.

Theorem 11 *The random variable X is DMRL if, and only if, $[X - t | X > t] \geq_{\text{hmrl}} [X - t' | X > t']$ whenever $t \leq t'$.*

Cao and Wang [7] observed that the order $\leq_{\text{icx}}$ can be used to characterize the DMRL and IMRL notions.

Theorem 12 *The random variable X is DMRL [IMRL] if, and only if, $[X - t | X > t] \geq_{\text{icx}} [\leq_{\text{icx}}] [X - t' | X > t']$ whenever $t \leq t'$.*

Belzunce *et al.* [3] similarly observed that the dilation order can be used to characterize these **aging notions**. The result below follows from their Theorem 4 and the fact that the dilation order is location invariant.

Theorem 13 *The random variable X is DMRL [IMRL] if, and only if, $[X - t | X > t] \geq_{\text{dil}} [\leq_{\text{dil}}] [X - t' | X > t']$ whenever $t \leq t'$.*

Belzunce [9] noticed that the excess wealth order can also be used to characterize the DMRL and IMRL notions.

Theorem 14 *Let X be a continuous random variable with support $(0, \infty)$. Then X is DMRL [IMRL] if, and only if, any one of the following equivalent conditions holds:*

1. $[X - t | X > t] \geq_{\text{ew}} [\leq_{\text{ew}}] [X - t' | X > t']$ whenever $t' \geq t \geq 0$.
2. $X \geq_{\text{ew}} [\leq_{\text{ew}}] [X - t | X > t]$ for all $t \geq 0$.

There exist analogs of Theorems 6, 7, and 9 for the DMRL aging notion. According to Kocher and Wiens [10] the random variable X is said to be *smaller than Y in the DMRL order* (denoted by $X \leq_{\text{dmrl}} Y$) if

$$\frac{\int_{G^{-1}(\alpha)}^{\infty} \bar{G}(x) dx}{\int_{F^{-1}(\alpha)}^{\infty} \bar{F}(x) dx} \text{ is increasing for } \alpha \in [0, 1]$$

Theorem 15 *The random variable X is DMRL if, and only if, $X \leq_{\text{dmrl}} \text{Exp}$.*

Li and Li [11] introduced another order, denoted by $\leq_{\text{drlc}}$, such that X is DMRL if, and only if, $X \leq_{\text{drlc}} \text{Exp}$.

New Better (Worse) than Used in Expectation (NBUE and NWUE)

The nonnegative random variable X with finite mean is said to have the aging property of new better than used in expectation (NBUE) if $E[X] \geq E[X - t | X > t]$ for all $t \geq 0$. It has the property of new worse than used in expectation (NWUE) if $E[X] \leq E[X - t | X > t]$ for all $t \geq 0$. Lefèvre and Utev [12] proved that the harmonic mean residual life order can be used to characterize the NBUE notion as follows.

Theorem 16 *Let X be a nonnegative random variable with positive mean. Then X is NBUE if, and only if, any one of the following equivalent conditions holds:*

1. $X \leq_{\text{hmrl}} X + Y$ for any nonnegative random variable Y with a finite positive mean, which is independent of X .
2. $X + Y_1 \leq_{\text{hmrl}} X + Y_2$ whenever Y_1 and Y_2 are almost surely positive random variables with finite means, which are independent of X , such that $Y_1 \leq_{\text{hmrl}} Y_2$.

There exist analogs of Theorems 6, 7, 9, and 15 for the NBUE aging notion. According to Kochar and Wiens [10] the random variable X is said to be *smaller than Y in the NBUE order* (denoted by $X \leq_{\text{nbue}} Y$) if

$$\frac{1}{EX} \int_{F^{-1}(\alpha)}^{\infty} \bar{F}(x) dx \leq \frac{1}{EY} \int_{G^{-1}(\alpha)}^{\infty} \bar{G}(x) dx$$

for all $\alpha \in [0, 1]$

Theorem 17 *The random variable X is NBUE if, and only if, $X \leq_{\text{nbue}} \text{Exp}$.*

Note that for any two nonnegative random variables Z and W with positive expectations we have $Z \leq_{\text{nbue}} W \iff (Z/EX) \leq_{\text{ew}} (W/EX)$. Thus, Theorem 17 can be rewritten as follows.

Theorem 18 *The random variable X is NBUE if, and only if, $(X/EX) \leq_{\text{ew}} \text{Exp}(1)$, where $\text{Exp}(1)$ denotes an exponential random variable with mean 1.*

An aging notion that is closely related to the NBUE notion is the notion of the harmonic new better than used in expectation (HNBUE). Formally, the random variable X with mean $\mu > 0$ is said to have the aging property of HNBUE if $X \leq_{\text{icx}} \text{Exp}(\mu)$, where $\text{Exp}(\mu)$ denotes an exponential random variable with mean μ . The random variable X with mean $\mu > 0$ is said to have the property of harmonic new worse than used in expectation (HNWUE) if $X \geq_{\text{icx}} \text{Exp}(\mu)$. Note that by known basic properties of the orders $\leq_{\text{icx}}$ and $\leq_{\text{dil}}$ it follows that the random variable X with mean $\mu > 0$ is HNBUE [HNWUE] if, and only if, $X \leq_{\text{dil}} [\geq_{\text{dil}}] \text{Exp}(\mu)$. The harmonic mean residual life order can be used to characterize the HNBUE and HNWUE notions as follows.

Theorem 19 *The random variable X with mean $\mu > 0$ is HNBUE [HNWUE] if, and only if, $X \leq_{\text{hmrl}} [\geq_{\text{hmrl}}] \text{Exp}(\mu)$.*

References

- [1] Shaked, M. & Shanthikumar, J.G. (2007). *Stochastic Orders*, Springer, New York.
- [2] Müller, A. & Stoyan, D. (2002). *Comparison Methods for Stochastic Models and Risks*, John Wiley & Sons, New York.
- [3] Belzunce, F., Candel, J. & Ruiz, J.M. (1996). Dispersive ordering and characterizations of aging classes, *Statistics and Probability Letters* **28**, 321–327.
- [4] Pellerey, F. & Shaked, M. (1997). Characterizations of the IFR and DFR aging notions by means of the dispersive order, *Statistics and Probability Letters* **33**, 389–393.
- [5] Belzunce, F., Hu, T. & Khaledi, B.-E. (2003). Dispersion-type variability orders, *Probability in the Engineering and Informational Sciences* **17**, 305–334.
- [6] Belzunce, F., Gao, X., Hu, T. & Pellerey, F. (2004). Characterizations of the hazard rate order and IFR aging notion, *Statistics and Probability Letters* **70**, 235–242.
- [7] Cao, J. & Wang, Y. (1991). The NBUC and NWUC classes of life distributions, *Journal of Applied Probability* **28**, 473–479.
- [8] Deshpande, J.V., Kochar, S.C. & Singh, H. (1986). Aspects of positive ageing, *Journal of Applied Probability* **23**, 748–758.
- [9] Belzunce, F. (1999). On a characterization of right spread order by the increasing convex order, *Statistics and Probability Letters* **45**, 103–110.
- [10] Kochar, S.C. & Wiens, D.P. (1987). Partial orderings of life distributions with respect to their aging properties, *Naval Research Logistics* **34**, 823–829.
- [11] Li, B. & Li, X. (2005). New partial orderings of life distributions with respect to the residual life function, *Journal of Lanzhou University (Natural Sciences)* **41**, 134–138.
- [12] Lefèvre, C. & Utev, S. (2001). Comparison of individual risk models, *Insurance Mathematics and Economics* **28**, 21–30.

Related Articles

Aging and Positive Dependence; Multivariate Stochastic Orders and Aging.

FÉLIX BELZUNCE AND MOSHE SHAKED

Characterization in Reliability

What is the Problem of Characterization of Probability Distributions?

The problem of characterization of distributions is today a section of **probability theory** and mathematical statistics that is rather substantial in its size and variety of accumulated facts. The aim of characterization theory is to describe a class of all probability distributions possessing some desirable property $\mathcal{P}$. The property $\mathcal{P}$ is usually the main property we would like to use for the construction of mathematical models of real world phenomenon. Characterization theorems utilize numerous classical tools of mathematical analysis, such as, advanced theory of functions of complex variables, differential equations of different types, theory of functional equations, and other tools. It is clear that $\mathcal{P}$ may be connected, say, with the property of **aging**, and therefore the connection between characterization theory and reliability is obvious. It is possible to have classification of corresponding problems of reliability theory based on the types of property $\mathcal{P}$ in use or the analytical tools utilized to obtain the solution. Both methods are popular among the specialists.

Characterizations Connected to the Lack of Memory Property

A considerable amount of results dealing with characterization of exponential distribution, the Weibull distribution, and some other laws either have their origin in mathematical theory of reliability or can be interpreted in terms of this theory. Some of them are of certain practical interest, while others have only a theoretical importance. However, all these results are important since they give an idea of a mechanism of appearance of distributions either in problems associated with the aging of technical devices or in connection with various methods of control of these devices (e.g., with verification or durability). One of the best analyzed characterizations of this kind is the exponential distribution by the lack of memory property, or, in other words, by the lack of aging property.

Lack of memory (aging) property is defined by the following. Let T be a positive continuous random variable (rv) having interpretation as the lifetime of a device. The probability of having a failure in time interval $(t, t + \Delta t]$ given that there was no failure on time interval $(0, T]$ is

$$\mathbb{P}\{t < T \leq t + \Delta t | T > t\} = \frac{\mathbb{P}\{t < T \leq t + \Delta t\}}{\mathbb{P}\{T > t\}} \quad (1)$$

On the other hand, for a new device of the same type, the probability of having a failure in a time interval of the same length Δt but starting from zero that is, in the interval $(0, \Delta t)$ is $\mathbb{P}\{0 < T \leq \Delta t\}$. If there is no aging for the device, the probabilities have to be identical:

$$\frac{\mathbb{P}\{t < T \leq t + \Delta t\}}{\mathbb{P}\{T > t\}} = \mathbb{P}\{0 < T \leq \Delta t\} \quad (2)$$

for all $t > 0, \Delta t > 0$. This relation can be expressed in terms of a survival function (sf) $S(x) = \mathbb{P}\{T > x\}$, and coincides with the well-known Cauchy functional equation for exponential function [1]. It means *the lack of aging property is valued for the exponential distribution with arbitrary scale parameter and only for it*. An equivalent form of the lack of aging property is

$$\mathbb{P}\{T > t + x | T > t\} = \mathbb{P}\{T > x\} \quad (3)$$

A detailed study of the lack of memory property and its generalizations can be found in [2].

One of the most interesting generalizations of the lack of aging property is the *integrated (or averaged) lack of memory property*. This represents the case when x in equation (3) is a value of rv independent of T . One of typical formulations of the averaged lack of memory property is the following:

Let $\{X_t, t > 0\}$ be a family of rv with an sf $H_t(x)$ and T be a positive rv with sf $S(t)$. Suppose that $\min\{x \in \text{supp} H_t\} > t$. Then the averaged lack of memory property is defined by the relation

$$\mathbb{P}\{T > X_t | T > t\} = \mathbb{P}\{T > X_t - t\} \quad (4)$$

Under some additional conditions, the averaged lack of memory property is equivalent to the exponentiality of rv T . The averaged lack of memory property can be analytically reformulated as a convolution-type equation. Either the method of

2 Characterization in Reliability

intensively monotone operators [3] or the Choquet–Deny theorem [4] is used for the solution.

The aforementioned methods are applicable to another type of generalization of the lack of memory property. Namely, let h be a real-valued continuous function. Consider the condition of *constancy of conditional expected values*:

$$\mathbb{E}(h(T - z) | T > z) = \mathbb{E}(h(T)) \quad (5)$$

Under some restriction imposed on the function h this property is characteristic of exponential distribution. The method based on Choquet–Deny theorem [4] is applicable to the case of a monotone h function, while the method of intensively monotone operators [3] can be applied to nonmonotone h . There exist some examples, showing that the characterization is not true for some continuous nonmonotonic h [5].

Functioning of Redundant Systems

The *relevation-type equation* appeared in reliability theory in the following way. Let there be two types of elements working simultaneously. The first element is the servicing one and the second is an element in storage. If an element of the first type fails, then its role is fulfilled by a survival in storage (moreover, it is assumed that the distribution of failures of the element of the second type does not change). The failure of the device occurs after the failure of an element of the second type which has replaced the first one [6]. If the sf of the first (second) element is $S_1(t)$ ($S_2(t)$) then the sf of the system is $S_1(t) - S_2(t) \int_0^t dS_1(x)/S_2(x)$. This expression is called the *relevation of S_1 and S_2* . If an element of the second type would not work before the failure of a servicing element, then the sf of the device would be found with the help of a convolution of S_1 and S_2 that is, it would have the form of $S_1(t) - \int_0^t S_2(t - x) dS_1(x)$. Intuitively the coincidence of the relevation and the convolution of corresponding sf's signifies the lack of aging of an element of the second type. Really, under some mild conditions, the equation

$$\begin{aligned} S_1(t) - S_2(t) \int_0^t \frac{dS_1(x)}{S_2(x)} \\ = S_1(t) - \int_0^t S_2(t - x) dS_1(x) \end{aligned} \quad (6)$$

has only exponential solution for S_2 for any given S_1 [3].

Characterization of the sf's by a connection between the reliabilities of two systems

Consider technical systems $\mathcal{A}$ and $\mathcal{B}$, consisting of identical and mutually independent (from the point of failure probabilities) elements. The system $\mathcal{A}$ consists of n elements and the system $\mathcal{B}$ consists of k elements, $n > k \geq 1$. Assume that each system goes out of service under failure of at least one of its elements. If the time in service up to the failure of each element is distributed according to a nondegenerate law $F(t)$, which is continuous and strictly monotone for $t > 0$, then at any given moment of time $t > 0$ the reliability of the element is equal to $S(t) = 1 - F(t)$. Hence the reliability of the system $\mathcal{A}$ is equal to $h_1(t) = (S(t))^n$, and for the system $\mathcal{B}$ is equal to $h_2(t) = (S(t))^k$. Connect with each moment of time $t > 0$ a smallest moment $T(t) = T(t; F)$ such that the reliability of the system $\mathcal{A}$ at the moment t is equal to the reliability of the system $\mathcal{B}$ at the moment $T(t)$. It appears that the assignment of the function $T(t)$ characterizes the distribution F . It implies, for instance, the following fact. Let $b > 1$ be a constant. Assume that the limit $\lim_{t \rightarrow 0} S(t)/t^\alpha$ exists for $\alpha = \log(n/k)/\log b$. The equality $T(t; F) = bt$ holds for all $t > 0$ if and only if $F(t)$ is a distribution function of the Weibull law that is, $F(t) = 1 - \exp(-at^\alpha)$ for $t > 0$.

Among the distributions with nondecreasing hazard rate the function $T(t)$ takes maximal possible value at each moment t for exponential distribution.

The result of characterization of the Weibull distribution generalizes the following fact. Let $T_1, \dots, T_n$ be a sequence of independent identically distributed (i.i.d.) positive rv's with cumulative distribution function (cdf) F , and $T_{\min} = \min(T_1, \dots, T_n)$, $n > 2$. Suppose that for some $\alpha > 0$ there exists a positive limit $\lim_{t \rightarrow 0} (1 - F(t)/t^\alpha)$. Then T_1 and $n^{1/\alpha} T_{\min}$ are identically distributed if and only if F is the Weibull distribution. The existence of the limit above can be replaced by the fact of identical distribution of T_1 and $n^{1/\alpha} T_{\min}$ for two values of n with irrational ratio of their logarithms. Some generalizations for the case of non-i.i.d. can be founded in [3], where some generalizations for the case of random n for the characterization of the Pareto distribution are also given.

On the Reliability of Hierarchical Structures

Let $N = N_1$ be an arbitrary network called the *first generation*. The second generation N_2 is obtained from N_1 by replacing each component of N_1 by N . Similarly, if we replace each component of N_2 by N , then we get the third generation N_3 , and so on. Suppose that the components of N function independently and each component has the same probability p of proper functioning, then the chance that N functions properly is a polynomial of p , called the *reliability polynomial*, and will be denoted by $h(p)$ [7]. The n th function iteration of h is the reliability polynomial of N_n . Denote this n th iteration of h by $h^{(n)}(p)$. The theory of iterative functions is a good tool to study these hierarchical networks. The limiting distribution $\bar{F}$ of the sf's $\bar{F}_n$ of N_n with linear normalization (as n tends to infinity) can be reduced to the solution of the equation

$$\bar{F}(t) = h(\bar{F}(\tau t)) \quad (7)$$

for some $\tau \in (0, 1)$. This equation is equivalent to the property of identical distribution of T_1 and the maximum of a random number (having the probability generating function h) of i.i.d. moments $T_1, T_2, \dots$ taking in random number with probability generating function h . The description of all possible F can be given in terms of solution of Poincare equation, and depends on the number of fixed points for iterations of the function h (see [8]).

Characterizations in Special Classes of Distributions

Let T be a positive rv having sf $S(t)$. Remember that S is *NABU* (new better than used in average) if $\int_t^\infty S(x) dx \leq \mu_1 S(t)$, where $\mu_1 = \mathbb{E}(T)$. Denote $S_0(t) = S(t)$ and introduce $S_j = (j\mu_{j-1}/\mu_j) \int_t^\infty S_{j-1}(x) dx$, $\mu_j = \mathbb{E}(T^j)$. If S is *NABU* then $\mu_j \leq j!\mu_1^j$ with the equality sign if and only if S is an sf of the exponential distribution. The stability of this characterization can be expressed in the following way. Let S be an sf having the moments μ_j , $j = 1, \dots, k-1$. For any $1 \leq j \leq k$ the following holds

$$\sup_{t \geq 0} |S_j(t) - \exp\left(\frac{-t}{\mu_1}\right)| \leq 1 - \frac{\mu_{j+1}}{((j+1)\mu_1\mu_j)} \quad (8)$$

There are numerous publications on the use of characterizations for statistical test. One of such test belongs to Koul [9]. The Bahadur efficiency of this test was studied by Tchirina [10], who calculated the asymptotic of the large deviations for the test statistic, and found the class of alternatives for which the test is efficient.

References

- [1] Aczel, J. (1961). *Vorlesungen Über Funktionalgleichung und Ihre Anwendungen*, Birkhäuser, Basel and Stuttgart.
- [2] Galambos, J. & Kotz, S. (1978). *Characterizations of Probability Distributions*, Springer-Verlag, Berlin, Heidelberg, New York.
- [3] Kakosyan, A.V., Klebanov, L.B. & Melamed, J.A. (1984). *Characterization of Distributions by the Method of Intensively Monotone Operators*, Springer-Verlag, Berlin, Heidelberg, New York.
- [4] Radhakrishna Rao, C. & Shanbhag, D.N. (1989). Further extensions of the Choquet-Deny and Deny theorems with applications in characterization theory, *The Quarterly Journal of Mathematics* **40**, 333–350.
- [5] Klebanov, L.B. (1980). Several results connected with characterization of the exponential distribution, *Teoriya Veroyatnostei i Yeye Primeniya* **25**, 628–633.
- [6] Krakowski, M. (1973). The revelation transform and generalization of the gamma distribution function, *Revue Francaise d'Automatique, Informatique, Recherche Operationnelle* **7**, 107–120.
- [7] Barlow, R.E. & Proschan, F. (1975). *Statistical Theory of Reliability and Life Testing Probability Models*, Holt, Rinehart and Winston.
- [8] Klebanov, L.B. & Szekely, G.J. (2000). Characterization of distributions in reliability, in *Recent Advances in Reliability Theory: Methodology, Practice and Inference*, N. Limnios & M. Nikulin, eds, Birkhäuser, Boston, pp. 105–115.
- [9] Koul, H.L. (1977). A test for new better than used, *Communications in Statistics: Theory and Methods* **A6**, 563–574.
- [10] Tchirina, A.V. (1997). Bahadur efficiency of a test of exponentiality based on the loss-of-memory property, *Zapiski Nauchnyh Seminarov POMI* **244**, 315–329.

Related Article

Probability Density Functions.

LEV B. KLEBANOV

Random Evolutions toward Applications

Introduction

Random evolution or stochastic evolution were introduced by Griego and Hersh [1]. They observed that certain abstract Cauchy problems are related to generators of **Markov processes**. The name of *Random evolution* is due to P. Lax. (see [2]).

An evolutionary system that changes its mode of evolution following a stochastic mode is a random evolution. For example, a particle moving on the real line with constant velocity, say V , changes it at a random time, given by the arrival of a **Poisson process** to $-V$, that is, the velocity of the particle at time t is $(-1)^{N(t)}V$, where $N(t), t \geq 0$, is a (homogeneous) Poisson process on $[0, \infty)$.

Random evolutions are models to study stochastic systems in random media (environment). In the above example, the Poisson process represent the random medium.

From a mathematical point of view, random evolutions are operator-valued stochastic processes in a Banach space. They represent an abstract form of stochastic evolutionary systems in random media. Discrete-time random evolution studied in connection with Markov chains (see [3]), and as embedded random evolution in continuous-time semi-Markov process (see [4, 5]). In Skorokhod *et al.* [6] several random evolution problems are studied.

Applications of random evolutions include population dynamics in **random environment** [7, 8]; in insurance [9]; reliability analysis [5] and so on.

In this paper we present random evolutions *via* dynamical systems, the so-called piecewise deterministic Markov processes (see [10]) in ergodic and nonergodic media.

In the next section a particular dynamical system in Markov media and its abstract model in terms of random evolutions is presented. In the section on “Asymptotic Results” two types of limit theorems for Markov random evolution are presented. Finally, in the “Concluding Remarks”, some remarks on the literature and further topics connected to random evolution are presented.

Evolutionary Systems in Markov Random Medium

Let us consider the dynamical system

$$\begin{cases} \frac{d}{dt}U(t) &= C(U(t); x(t)), \quad t \geq 0 \\ U(0) &= u \end{cases} \quad (1)$$

where the process $x(t), t \geq 0$, is a Markov process with state space E , and generator $\mathbf{Q}$, and $U(t) \in \mathbf{R}^d$. This system usually describes dynamic reliability where the Markov process $x(t), t \geq 0$, describes the structure of the system and the process $U(t)$ the physical throughput in the system, for example, temperature, pressure, velocity, and so on (see [11]).

The coupled stochastic process $(U(t), x(t)), t \geq 0$, is a Markov process, with state space $\mathbf{R}^d \times E$ and generator $\mathbf{L}$:

$$\mathbf{L}\varphi(u, x) = \mathbf{\Gamma}(x)\varphi(u, \cdot) + \mathbf{Q}\varphi(\cdot, x) \quad (2)$$

where the operator $\mathbf{\Gamma}(x)$ acts on functions $\varphi(u) \in C^1(\mathbf{R}^d)$ as follows

$$\mathbf{\Gamma}(x)\varphi(u) = C(u; x) \left(\frac{\partial}{\partial u_1}, \dots, \frac{\partial}{\partial u_d} \right) \varphi(u_1, \dots, u_d) \quad (3)$$

For example, in a two-state Markov process, $x(t)$, say $E = \{0, 1\}$ and $d = 1$, the generator $\mathbf{L}$ can be written in the following matrix form

$$\mathbf{L} = \begin{pmatrix} C_0(u) \frac{\partial}{\partial u} & 0 \\ 0 & C_1(u) \frac{\partial}{\partial u} \end{pmatrix} + \begin{pmatrix} -\lambda & \lambda \\ \mu & -\mu \end{pmatrix} \quad (4)$$

where $C_x(u) := C(u; x), x = 0, 1$.

As the trajectory of $U(t)$ at time $t + s$, with initial value $U(0) = u$, is an extension of the trajectory at time s of the trajectory with initial value $U(t; u, x)$ it is easy to see that the solution of equation (1) for fixed state $x \in E$, verifies the semigroup property

$$U(t + s; u, x) = U(s; U(t; u, x), x) \quad (5)$$

where $U(t; u, x)$ is the value of $U(t)$ on $\{U(0) = u, x(0) = x\}$.

2 Random Evolutions toward Applications

An abstract formulation of the above property can be represented by the following continuous Markov random evolution

$$\Phi(t)\varphi(u, x(t)) := \varphi(U(t); x(t)) \quad (6)$$

with the family of semigroups

$$\Gamma_x(t)\varphi(u) := \varphi(U(t; u, x)), \quad x \in E, t \geq 0 \quad (7)$$

and generators

$$\Gamma(x)\varphi(u) := C(u, x) \frac{d}{du} \varphi(u), \quad x \in E \quad (8)$$

Asymptotic Results

Let us consider the dynamical system in equation (1) in two different series schemes. The first one is the ergodic case for the switching process $x(t), t \geq 0$, and in the second one a split with absorption is considered. In both cases, average diffusion approximation limit theorems will be presented.

Let us consider the switching jump Markov processes $x^\varepsilon(t), t \geq 0, \varepsilon > 0$, in series scheme with series parameter $\varepsilon > 0$ ($\varepsilon \rightarrow 0$), with state space $(E, \mathcal{E})$, where E is a Polish space and $\mathcal{E}$ its Borel σ -algebra, and generators

$$Q^\varepsilon \varphi(x) = q(x) \int_E P^\varepsilon(x, dy) [\varphi(y) - \varphi(x)] \quad (9)$$

where $P^\varepsilon(x, dy)$ is the transition kernel of the embedded Markov chain $x_n^\varepsilon = x^\varepsilon(\tau_n^\varepsilon)$ and τ_n^ε are the jump times of $x^\varepsilon(t), t \geq 0$. The function $q(x)$ is the intensity of jumps.

Ergodic Case

Let us first present the result on averaging. The following assumptions are needed.

A1. The switching Markov process $x(t), t \geq 0$, is uniformly ergodic with stationary probability measure $\pi(B), B \in \mathcal{E}$.

A2. The function $C(u; x)$ is globally Lipschitz continuous on $u \in \mathbf{R}^d$ with common constant L for all $x \in E$. So, the system

$$\frac{d}{dt} U_x(t) = C(U_x(t); x) \quad (10)$$

has a global solution.

Theorem 1 [Average approximation] Under Assumptions A1 and A2, the solution $U^\varepsilon(t)$ of the following system

$$\begin{cases} \frac{d}{dt} U^\varepsilon(t) = C(U^\varepsilon(t); x(t/\varepsilon)) \\ U^\varepsilon(0) = u \end{cases} \quad (11)$$

converges weakly to the solution of the average equation

$$\begin{cases} \frac{d}{dt} \widehat{U}(t) = \widehat{C}(\widehat{U}(t)) \\ \widehat{U}(0) = u \end{cases} \quad (12)$$

It is worth noticing that this is a deterministic system and this averaging scheme is the so-called Bugolyubov principle.

For the *diffusion approximation* we consider the following time-scaling of the switching process.

$$\begin{cases} \frac{d}{dt} U^\varepsilon(t) = C^\varepsilon(U^\varepsilon(t); x(t/\varepsilon^2)) \\ U^\varepsilon(0) = u \end{cases} \quad (13)$$

where $C^\varepsilon(u; x) := \varepsilon^{-1} C(u; x) + C_1(u; x)$.

We suppose here that the following balance condition is fulfilled

$$\widehat{C}(u) := \int_E \pi(dx) C(u; x) = 0 \quad (14)$$

Theorem 2 [Diffusion approximation] Under Assumption A1, and the balance condition (14), the solution of the system (13) converges to a diffusion process, provided that the diffusion coefficient B is > 0 , the limit diffusion process is defined by the following infinitesimal generator

$$\mathbf{L}\varphi(u) = \widehat{C}_u \varphi'(u) + \frac{1}{2} B(u) \varphi''(u) \quad (15)$$

where the drift coefficient is

$$\widehat{C}_u = \int_E \pi(dx) [C_1(u; x) + C(u; x) R_0 C'_u(u; x)] \quad (16)$$

and the diffusion coefficient is

$$B(u) = 2 \int_E \pi(dx) C_0(u; x) R_0 C(u; x) \quad (17)$$

Operator R_0 is the potential operator of Q , that is

$$QR_0 = R_0Q = \Pi - I \quad (18)$$

and Π is the projector operator defined by $\Pi(x, dy) = \pi(dy)$, for every $x \in E$.

Notation φ', φ'' stand for the first and second derivative of function φ , and $C'_u(u; x)$ stands for the partial derivative of $C(u; x)$ with respect to the variable u .

Nonergodic Case

While the ergodic case is more adapted to availability modeling, the nonergodic case, presented here, is more adapted to reliability modeling.

Let us now consider the following split to the state space of the Markov switching processes $x^\varepsilon(t)$, $t \geq 0$, $\varepsilon > 0$,

$$E^0 = E \cup \{0\}, \quad E = \bigcup_{k=1}^N E_k, \quad E_k \cap E_{k'} = \emptyset, \quad k \neq k' \quad (19)$$

with absorbing state $\{0\}$. An interpretation of this split in reliability problems is that $E_1, \dots, E_N$ are subsets of different performance levels and 0 is the failure state of the system.

Let us suppose that the transition kernel P^ε has the following representation

$$P^\varepsilon(x, B) = P(x, B) + \varepsilon P_1(x, B) \quad (20)$$

where $P(x, B)$ is a transition kernel, coordinated with the above split (19) as follows

$$P(x, E_k) = \mathbf{1}_k(x) := \begin{cases} 1, & x \in E_k \\ 0, & x \notin E_k \end{cases} \quad (21)$$

and $P_1(x, B)$ a perturbing kernel.

The Markov supporting process $x(t)$, $t \geq 0$, on the state space $(E, \mathcal{E})$, determined by the generator

$$Q\varphi(x) = q(x) \int_E P(x, dy) [\varphi(y) - \varphi(x)] \quad (22)$$

is supposed to be uniformly ergodic in every class E_k , $1 \leq k \leq N$, with the stationary distribution $\pi_k(dx)$, $1 \leq k \leq N$, satisfying the following relations

$$\pi_k(dx)q(x) = q_k\rho_k(dx), \quad q_k = \int_{E_k} \pi_k(dx)q(x) \quad (23)$$

$$\rho_k(B) = \int_{E_k} \rho_k(dx)P(x, B), \quad \rho_k(E_k) = 1 \quad (24)$$

where ρ_k is the stationary probability of the embedded Markov chain $x_n, n \geq 0$, of the supporting Markov process $x(t)$, $t \geq 0$, on the class E_k , $1 \leq k \leq N$.

Let v be the merging function, defined by

$$v(x) = k \quad \text{if } x \in E_k, \quad 1 \leq k \leq N \quad (25)$$

Then we have

$$v\left(x^\varepsilon\left(\frac{t}{\varepsilon}\right)\right) \Longrightarrow \widehat{x}(t), \quad \varepsilon \rightarrow 0 \quad (26)$$

$\Longrightarrow$ stands for weak convergence of probability measures.

The process $\widehat{x}(t)$, $t \geq 0$ is a Markov process on the state space $E^0 = E \cup \{0\}$, with $\widehat{E} = \{1, \dots, N\}$.

The random variable ζ^ε is the absorption time or lifetime or failure time for the system, and is defined by

$$\zeta^\varepsilon := \inf \left\{ t \geq 0 : x^\varepsilon\left(\frac{t}{\varepsilon}\right) = 0 \right\} \quad (27)$$

We have $\zeta^\varepsilon \Longrightarrow \widehat{\zeta}$, as $\varepsilon \rightarrow 0$, where $\widehat{\zeta}$ is the absorption time of the merged process $\widehat{x}(t)$.

Theorem 3 [Average approximation with split and merging] Under the above assumptions the stochastic system $U^\varepsilon(t)$, $t \geq 0$, defined by

$$\begin{cases} \frac{d}{dt} U^\varepsilon(t) &= C(U^\varepsilon(t); x^\varepsilon(t/\varepsilon)) \\ U^\varepsilon(0) &= u \end{cases} \quad (28)$$

converges weakly to the averaged stochastic system $\widehat{U}(t \wedge \widehat{\zeta})$:

$$U^\varepsilon(t \wedge \zeta^\varepsilon) \Longrightarrow \widehat{U}(t \wedge \widehat{\zeta}), \quad \text{as } \varepsilon \rightarrow 0 \quad (29)$$

The limit process $\widehat{U}(t)$, $t \geq 0$, is defined by a solution of the evolutionary equation

$$\frac{d}{dt} \widehat{U}(t) = \widehat{C}(\widehat{U}(t); \widehat{x}(t)), \quad \widehat{U}(0) = 0 \quad (30)$$

on the time interval $0 \leq t \leq \widehat{\zeta}$, ($\widehat{\zeta}$ is the stoppage time of the merged Markov process $\widehat{x}(t)$, $t \geq 0$).

The averaged velocity is determined by

$$\begin{aligned}\widehat{C}(u; k) &= \int_{E_k} \pi_k(dx) C(u; x), \quad 1 \leq k \leq N, \\ \widehat{C}(u; 0) &= 0\end{aligned}\quad (31)$$

For the diffusion approximation scheme, we suppose that $N = 1$ for simplicity, and that the transition kernel has the following representation

$$P^\varepsilon(x, B) = P(x, B) + \varepsilon^2 P_1(x, B) \quad (32)$$

The random variable $\widehat{\xi}$ is exponentially distributed with parameter $\widehat{\Lambda} = pq$, where $p := -\int_E \rho(dx) P_1(x, E)$ and $q := \int_E \rho(dx)/q(x)$.

The considered system here is in the following time-scaling scheme.

$$\begin{cases} \frac{d}{dt} U^\varepsilon(t) &= C^\varepsilon(U^\varepsilon(t); x^\varepsilon(t/\varepsilon^2)) \\ U^\varepsilon(0) &= u \end{cases} \quad (33)$$

where $C^\varepsilon(u; x) := \varepsilon^{-1} C(u; x) + C_1(u; x)$.

Theorem 4 [Diffusion approximation with split and merging] Under the balance condition (14), the stochastic system (33) converges weakly, as $\varepsilon \rightarrow 0$,

$$U^\varepsilon(t \wedge \xi^\varepsilon) \Longrightarrow \xi(t \wedge \widehat{\xi}) \text{ as } \varepsilon \rightarrow 0 \quad (34)$$

The limit diffusion process $\xi(t)$, $t \geq 0$, is defined by the generator

$$\mathbf{L}\varphi(u) = b(u)\varphi'(u) + \frac{1}{2}B(u)\varphi''(u) - \widehat{\Lambda}\varphi(u) \quad (35)$$

The drift coefficient is defined by

$$b(u) = \int_E \pi(dx) [C_1(u; x) + C(u; x) R_0 C'_u(u; x)] \quad (36)$$

The **covariance** function is defined by

$$B(u) = 2 \int_E \pi(dx) C(u; x) R_0 C(u; x) \quad (37)$$

The parameter $\widehat{\Lambda}$ in (35) is the parameter of the exponential distribution of the lifetime of the merged Markov process and of the limit diffusion process too.

Concluding Remarks

Detailed definitions of random evolutions, in semi-Markov and Markov media, can be found in [5]. Average and diffusion approximation scheme of semi-Markov random evolutions with split and merging of the phase space as well as nonergodic switching processes with applications in reliability can be found in [5]. For more references on random evolutions see books [4, 5, 12–14], and review paper [15] and references thereby.

Apart from the average and diffusion approximation considered here, random evolutions can be used also in Poisson and Lévy approximation schemes, where they are naturally associated to the predictable characteristics in additive semimartingale approach, see, for example, [5, Chapter 7].

References

- [1] Griego, R. & Hersh, R. (1969). Random evolutions, Markov chains, and systems of partial differential equations, *Proceedings of the National Academy of Sciences of the United States of America* **62**, 305–308.
- [2] Hersh, R. (2003). The birth of random evolutions, *Mathematical Intelligencer* **25**(1), 53–60.
- [3] Keepler, M. (1998). Random evolutions processes induced by discrete time Markov chains, *Portugaliae Mathematica* **55**(4), 391–400.
- [4] Korolyuk, V.S. & Swishchuk, A. (1995). *Random Evolution for Semi-Markov Systems*, Kluwer Academic Publishers, Dordrecht.
- [5] Korolyuk, V.S. & Limnios, N. (2005). *Stochastic Systems in Merging Phase Space*, World Scientific, Singapore.
- [6] Skorokhod, A.V., Hoppensteadt, F.C. & Salehi, H. (2002). *Random Perturbation Methods with Applications in Science and Engineering*, Springer-Verlag, New York.
- [7] Cohen, J.E. (1984). Eigenvalue inequalities for random evolutions: origins and open problems, *Inequalities in Statistics and Probability (IMS)* **5**, 41–53.
- [8] Ethier, S.N. & Kurtz, T.G. (1986). *Markov Processes: Characterization and Convergence*, John Wiley & Sons, New York.
- [9] Swishchuk, A. (1999). Stochastic stability and optimal control of semi-Markov risk processes in insurance mathematics, in *Semi-Markov Models and Applications*, J. Janssen & N. Limnios, eds, Kluwer Academic Publishers, Dordrecht, pp. 313–323.
- [10] Davis, M. (1993). *Markov Models and Optimisation*, Chapman & Hall, London.
- [11] Devoght, J. (1997). Dynamic reliability, *Advances in Nuclear Science and Technology* **25**, 215–278.
- [12] Pinsky, M. (1991). *Lectures on Random Evolutions*, World Scientific, Singapore.

- [13] Korolyuk, V.S. & Swishchuk, A. (1995). *Evolution of Systems in Random Media*, CRC Press.
- [14] Korolyuk, V.S. & Korolyuk, V.V. (1999). *Stochastic Models of Systems*, Kluwer Academic Publishers, Dordrecht.
- [15] Hersh, R. (1974). Random evolutions: a survey of results and problems, *Rocky Mountain Journal of Mathematics* **4**, 443–477.
- Further Reading**
- Cogburn, E. (1984). The ergodic theory of Markov chains in random environment, *Zeitschrift für Wahrscheinlichkeitstheorie und Verwandte Gebiete* **66**, 109–128.
- Cogburn, R. & Hersh, R. (1973). The limit theorems for random differential equations, *Indiana University Mathematics Journal* **22**, 1067–1089.
- Comets, F. & Pardoux, E. (eds) (2001). *Milieux Aléatoires*, Société Mathématique de France, Vol. 12.
- Gihman, I.I. & Skorohod, A.B. (1974). *Theory of Stochastic Processes*, Springer-Verlag, Berlin, Vol 2.
- Griego, R. & Hersh, R. (1971). Theory of random evolutions with applications to partial differential equations, *Transactions of the American Mathematical Society* **156**, 405–418.
- Hersh, R. & Papanicolaou, G. (1972). Non-commuting random evolutions, and an operator-valued Feynman-Kac formula, *Communications on Pure and Applied Mathematics* **30**, 337–367.
- Hersh, R. & Pinsky, M. (1972). Random evolutions are asymptotically Gaussian, *Communications on Pure and Applied Mathematics* **25**, 33–44.
- Jefferies, B. (1996). *Evolution Processes and the Feynman-Kac Formula*, Kluwer Academic Publishers, Dordrecht.
- Kolesnik, A. (1998). The equations of markovian random evolution on the line, *Journal of Applied Probability* **35**, 27–35.
- Korolyuk, V.S. (1999). Stochastic models of systems in reliability problems, in *Statistical and Probabilistic Models in Reliability*, D.C. Ionescu & N. Limnios, eds, Birkhäuser, Boston, pp. 127–141.
- Korolyuk, V.S. & Turbin, A.F. (1993). *Mathematical Foundation of the State Lumping of Large Systems*, Kluwer Academic Publishers, Dordrecht.
- Limnios, N. & Oprea, G. (2001). *Semi-Markov Processes and Reliability*, Birkhäuser, Boston.
- Papanicolaou, G. (1987). *Random Media*, Springer-Verlag, Berlin.
- Papanicolaou, G., Kohler, W. & White, B. (1991). *Random Media, Lectures in Applied Mathematics Vol. 27*, SIAM, Philadelphia.
- Zacks, S. (2004). Generalized integrated telegraph processes and their distribution of related stopping times, *Journal of Applied Probability* **41**, 497–507.

VLADIMIR S. KOROLIUK AND NIKOLAOS
LIMNIO

Damage Processes

Introduction and Background

There is extensive and burgeoning material on the topic of damage and its associated factors like **aging**, cumulative damage, **degradation**, deterioration, fatigue, health status, and quality of life. This material appears in both the biostatistical and the engineering reliability literatures. However, these notions suffer from the feature that they lack a precise definition. Rather they convey an abstract but intuitive import in the sense of a decrease in residual (or remaining) life. This decrease in residual life is conceptualized *via* the feature that an item experiencing aging and degradation will fail when the damage hits some barrier or threshold. Alternatively, it is supposed that at inception, every item is endowed with a resource that gets depleted because of damage, and that the item fails when the resource gets exhausted. Thus, for example, to engineers like Bogdanoff and Kozin [1], “Degradation is the irreversible accumulation of damage throughout life that leads to failure.” The term damage is not made precise; however it is claimed that damage reveals itself *via surrogates* or *markers*, such as cracks that grow in size, corrosion, measured wear (i.e., depletion of material), and so on. Similarly, Sobczyk [2] sees fatigue as “a phenomenon which takes place in units experiencing time-varying external actions which manifest in a deterioration of the unit’s resistance to carry its intended loading”. In the biostatistical literature, *aging* pertains to a unit’s position in a state space wherein the probabilities of failure are greater than its former position. Aging manifests itself in terms of biomedical and physical difficulties experienced by individuals, and in certain scenarios, *via* things like low-CD4 cell counts; these serve as biomedical surrogates, or what are known as *biomarkers*.

The markers mentioned above are, in most cases, observable and measurable entities. Much of the recent work on what is known as *degradation modeling* centers around assessing lifetimes *via* an analysis of the observed markers and their hitting times to a threshold (*cf.* Doksum [3], Doksum and Normand [4], Lu and Meeker [5], Ebrahimi [6], and Lehmann [7]). However, treating the observable markers as substitutes for the unobservable degradation process

that actually causes failure is tantamount to putting the cart before the horse. This is because the unobservable degradation process spawns the observable marker process, and is therefore its cause. An exception, however, is the work of Whitmore *et al.* [8] and of Lee *et al.* [9], who treat the degradation and the marker as separate but related processes. Also, see Nair [10], who makes the point that data on the observable surrogates of degradation help *sharpen* lifetime assessments. In this vein, a noteworthy contribution is by Cox [11] who systematically articulates the roles that the observable and the unobservable play in lifetime assessments. The premise upon which our bivariate stochastic process model with a random threshold is based has been inspired by the papers of Whitmore *et al.* [8] and Cox [11], and our work on hazard potential (see Singpurwalla [12; 13, p. 79]).

Preliminaries: The Hazard Potential

For an appreciation of the bivariate stochastic process model as a description of the damage process, some preliminaries on the notion of a hazard potential would be helpful. Accordingly, let T denote the lifetime of a unit and let $h(t)$ be the hazard rate of $P(T \geq t), t \geq 0$; let $H(t) = \int_0^t h(u) du$ be the cumulative hazard function at t . Then it is easy to see that

$$\begin{aligned} P(T \geq t; h(t), t \geq 0) \\ = \exp(-H(t)) = P(X \geq H(t)) \end{aligned} \quad (1)$$

where X has an exponential distribution with scale one. The random variable X is called the *hazard potential* of the item, and it represents an unknown “resource” that the item is endowed with at inception. Furthermore $H(t)$ is a measure of the amount of resource consumed by time t , and $h(t)$, the rate at which the resource is consumed at t . The unit fails when $H(t)$ exceeds X ; that is when $H(t)$ hits the random threshold X .

When the rate at which a unit’s resource gets consumed is random, $h(t)$ is described by a stochastic process, making $\{H(t); t \geq 0\}$ a stochastic process as well. However, this latter process has to be nondecreasing. The unit fails when the process $\{H(t); t \geq 0\}$ hits a barrier X , where X is also random with a unit exponential distribution. Candidate stochastic processes for $\{H(t); t \geq 0\}$ are

also alluded to in Singpurwalla [12]. Since $H(t)$ is, in principle, nondecreasing in t , $t \geq 0$, the process $\{H(t); t \geq 0\}$ is a candidate for describing a damage process. Furthermore, since the conventional view claims that an item fails when the damage hits a threshold, the cumulative hazard and the damage reflect a parallel feature. This motivates us to view the (cumulative) damage as being isomorphic with the cumulative hazard. Doing so makes our perspective different from that which is currently being discussed in literature. Like (cumulative) damage, the cumulative hazard is not observable. However, the cumulative hazard does influence the time to failure. Consequently, the cumulative hazard and the (cumulative) damage are to be seen as *latent variables*, and for that matter, so is the hazard potential X .

A Stochastic Process Model for Damage and its Markers

Because markers are closely linked with damage, any suitable model for the damage process should be accompanied by some sort of description for the markers as well. The most general way to do this would be to assume that the markers are realizations of stochastic processes, just as the (cumulative) damage is a stochastic process. The simplest way to proceed would be to suppose that there is only one marker to focus attention upon, so that a bivariate stochastic process $\{H(t), Z(t); t \geq 0\}$ would be a suitable description of the damage and its marker. As stated in the section titled “Preliminaries: The Hazard Potential”, the process $\{H(t); t \geq 0\}$ is nondecreasing in t . However, the process $\{Z(t); t \geq 0\}$ need not be restricted to being nondecreasing. Indeed, markers such as crack growth and CD4 cell counts fluctuate around some trend, and thus one is free to choose any suitable model for the process $\{Z(t); t \geq 0\}$. A Wiener process appears to be the model of choice, but this need not be so.

Thus to summarize, our proposed model for the (cumulative) damage and its associated marker is a bivariate stochastic process $\{H(t), Z(t); t \geq 0\}$ with $H(t)$ nondecreasing in t , and $Z(t)$ free to fluctuate around some constant or trend. We term such a process a *degradation process*. Since $H(t)$ spawns $Z(t)$, the two processes $\{H(t); t \geq 0\}$ and $\{Z(t); t \geq 0\}$ need to be linked; that is, they need

to be *cross-correlated*. Without such linkage, the marker process cannot serve as predictor of failure, and the statistical exercise of degradation modeling is not meaningful. One way to achieve this linkage is to describe $\{Z(t); t \geq 0\}$ by a Wiener process, and the unobservable (cumulative) damage process by a *Wiener maximum process*, namely,

$$H(t) = \sup_{0 \leq s \leq t} \{Z(s); s \geq 0\} \quad (2)$$

This strategy has been proposed in Singpurwalla [14], wherein a Bayesian approach for inference about lifetimes, using data on the marker process, is also described. The item fails when $H(t)$ hits the (random) threshold X . Whereas the model of equation (2) could be a starting point, there is a *caveat* that needs to be addressed. Specifically, since $Z(t)$ is spawned by $H(t)$, the latter is the cause of the former. This means that $H(t)$ must lead $Z(t)$, and so any linkage between the two processes in question should incorporate a time lag. The model of equation (2) does not do this because here $H(t)$ is determined retrospective to $Z(t)$ and therefore lags $Z(t)$, instead of the other way around. Thus $H(t)$ and $Z(t)$ need to be connected, with the observable $Z(t)$ lagging the unobservable $H(t)$. This is a possible topic for future research.

In the section titled “Candidate Processes for Damage and Markers”, we give an overview of some modeling strategies that have been proposed for the damage process $\{H(t); t \geq 0\}$, as well as for the marker process $\{Z(t); t \geq 0\}$ when each are treated separately; that is when no distinction is made between the damage (or degradation) process and the marker process. Supplementary material on the above can also be found in Chapters 7 and 8 of Singpurwalla [13].

Candidate Processes for Damage and Markers

The origins of the work on threshold crossing of cumulative damage as a basis for failure goes back to Epstein [15], Esary [16], and Gaver [17]. The idea of describing cumulative damage as a stochastic process can be traced, to the best of our knowledge, to Cox [18, p. 91], and to the Ph.D. thesis of Morey [19]. However, the granddaddy of all work on damage processes is the remarkable paper of Esary

et al. [20], who (without articulating what damage means) describe damage by a compound **Poisson process** with increments that are positive and have the Markov property. Failure occurs when the said process hits a random barrier whose distribution is exponential. The choice of an exponential distribution for the barrier is arbitrary, and Esary *et al.* [20] show that the hitting time has an exponential distribution. More recently Zacks [21, 22] has elaborated on this theme.

In the biostatistical arena there is a setup parallel to that of Esary *et al.* [20], which does not allude to damage, deterioration, or aging, but to the number of mice at some time t , that have typhoid organisms. The growth of such mice is described by a pure birth process, and the first passage time to a barrier is investigated. The specifics are in Cox and Miller [23, p. 160], and in Cox [11].

The Esary *et al.* [20] architecture is enhanced by Lemoine and Wenocur [24], who model wear (i.e., damage) by a suitable random process but who also allow for failure due to trauma. The latter is described by a Poisson process, the rate of which depends on the state of wear. An item fails when the wear reaches a threshold or when the item experiences fatal trauma. Thus in the model of Lemoine and Wenocur [24], the wear and the trauma processes compete with each other for an item's lifetime. The random process considered by the above authors is a diffusion process that is driven by a Wiener process. In a subsequent paper, Lemoine and Wenocur [25] describe wear by a shot-noise process. A disadvantage of the diffusion and the shot-noise process is that the wear (to us damage) is not monotonically nondecreasing. To rectify this deficiency, Wenocur [26] considers a gamma process for describing wear. His development of the gamma process proceeds along the following lines.

Partition the time interval into subintervals of length h , and let $X(n)$ denote the damage (or wear) at time nh , $n = 1, 2, \dots$. Suppose that the damage at time $(n+1)h$ is prescribed *via* the relationship

$$X(n+1) - X(n) = \alpha(X(n))\varepsilon_n + \beta(X(n))h \quad (3)$$

where α, β are constants, and $\{\varepsilon_n\}$ is a sequence of independent and identically distributed random variables having a gamma distribution with shape parameter $h > 0$. Letting $h \downarrow 0$, we have

$$dX(t) = \alpha(X(t^-)) d\gamma(t) + \beta(X(t^-)) \quad (4)$$

where $\{\gamma(t)\}$ is a *gamma process*.

In integral terms, equation (4) becomes the *stochastic integral*

$$X(t) = X(0) + \int_0^t \alpha(X(s^-)) d\gamma(s) + \int_0^t \beta(X(s^-)) ds \quad (5)$$

Since the gamma process has nonnegative increments, the wear (or damage) process is increasing. For an overview of the gamma process and their constructions, see Singpurwalla [27], or van der Weid [28]. Whereas a gamma process model may be attractive in scenarios wherein the damage causing shocks occur frequently, the models by Zacks [21, 22] for the compound Poisson process case and for the compound renewal process case, respectively, seem to be more appropriate when the shocks are infrequent.

Candidate Marker Processes

In engineering reliability, an archetypical marker process is crack growth, whereas in biostatistical studies, it appears that CD4 cell counts is a commonly mentioned biomarker. With archetypical markers come archetypical stochastic processes for $\{Z(t); t \geq 0\}$, and one such process is the Wiener process with a drift parameter η and a diffusion parameter $\sigma^2 > 0$; see Doksum [3] and Whitmore *et al.* [8]. As mentioned before, the marker is often viewed as a proxy for damage, and failure is said to occur when the marker process hits a threshold. As is well known, the hitting time to the threshold (assumed fixed and known) of a Wiener process has an *inverse Gaussian distribution* (see Singpurwalla [13, p. 68 and 136], for a discussion of this distribution).

The Wiener process has independent increments, so does a gamma process. This amounts to saying that the increments of crack size are independent of the existing crack length. This latter phenomenon is not always true. The bigger the crack, the bigger is its growth. This motivates one to consider transformations of the Wiener process. Furthermore, the crack growth phenomenon also exhibits abrupt growth. The Wiener process does not encapsulate such abruptness of growth. With the above in mind, Schabe [29] proposes the following as a model for $X(t)$, the size of a crack at time t , $t \geq 0$. Let $X(t) = (M(t))^a$, where

$$M(t) = bt + W(t) + \mu P(t) \quad (6)$$

Here $W(t)$ is a Wiener process with variance $\sigma^2 t$, and $P(t)$ is a Poisson process with intensity

4 Damage Processes

λ . The constant $b \geq 0$ describes a trend, and the constant a is such that $a > (<)1$ encapsulates a progressive (regressive) velocity with which the crack grows. Under the model of equation (6), Schabe [29] obtains the hitting time of $X(t)$ to $h > 0$, a barrier. This distribution does not have a closed-form solution. However, the mean and the variance of this distribution are available; these are $h^{h/a}(b + \mu\lambda)$ and $h^{1/a}(\sigma^2 + \lambda\mu^2)/(b + \mu\lambda)^3$, respectively.

In Ebrahimi [6], a strategy that parallels those of Lemoine and Wenocur [24] and of Schabe [29] is taken, but instead of looking at the growth of a single crack, an ensemble of k cracks, each having its own growth rate is considered. Specifically, it is supposed that the growth of the i th crack, $i = 1, \dots, k$ is governed by the stochastic differential equation

$$dX_i(t) = \lambda_i(t)X_i(t) + \sigma X_i(t) dW(t) \quad (7)$$

where $\lambda_i(t)$ is the growth rate of the cracks, $\sigma > 0$ is a constant, and $\{W(t); t \geq 0\}$ is a standard Wiener process with mean 0 and variance t . Using standard results (i.e., the Ito formula) it can be seen that

$$X_i(t) = X(0) \exp \left[\Lambda_i(t) - \frac{\sigma^2}{2}t + \sigma W(t) \right] \quad (8)$$

where $\Lambda_i(t) = \int_0^t \lambda_i(s) ds$, and $X(0)$ is the initial crack size, assumed to be known, and is the same for all the k cracks. The item fails when the size of the largest crack hits some threshold, say a . If T denotes the passage time of the largest crack to the threshold a , then it can be seen that for $0 \leq u \leq t$

$$P(T \geq t) = P \left[W(u) < \left(\min_{1 \leq i \leq k} \left(\frac{\sigma}{2}u - \frac{1}{\sigma} \Lambda_i(u) \right) + c \right) \right] \quad (9)$$

where $c = b/\sigma$ and $b = \log a - \log X(0)$.

Computation of the above follows from results on the the times taken by the Weiner process to hit a threshold.

Motivated by a model of Durham and Padgett [30], Park and Padgett [31] propose a model for cumulative damage, assuming that the damage is an observable measurable entity. This amounts to interpreting cumulative damage as a marker, like crack growth. The scheme proposed by Park and Padgett [31] is noteworthy, because it facilitates the introduction of both a **Brownian motion** process and

a gamma process as the driving processes for crack growth. Here, for some functions $c(\cdot)$ and $h(\cdot)$, it is assumed that

$$dc(X(t)) = h(X(t)) dD(t) \quad (10)$$

where $D(t)$ is the damage at t , and $X(t)$ is the cumulative damage at t . As a consequence of the above

$$\int_0^t \frac{1}{h(X(u))} dc(X(u)) = \int_0^t dD(u) = D(t) - D(0) \quad (11)$$

By choosing various forms for the function $c(\cdot)$ and $h(\cdot)$, and a stochastic process for $\{D(t); t \geq 0\}$, different models for $X(t)$ can be attained. For example, with $h(u) = 1$, $c(u) = \log u$, and a Brownian motion (or Wiener process) for $\{D(t); t \geq 0\}$, we obtain a geometric Brownian motion process for $X(t)$. With $h(u) = 1$, $c(u) = u$ and a Brownian motion process for $\{D(t); t \geq 0\}$ we obtain a Gaussian process for $X(t)$. With $h(u) = 1$, $c(u) = u$ and a gamma process for $\{D(t); t \geq 0\}$ we obtain a gamma process for $X(t)$. Whereas the Gaussian process is not always positive, the geometric Brownian motion process is always positive but not increasing. In contrast, the gamma process is both positive and increasing. Characteristics of the hitting times of the geometric Brownian motion and the gamma process to a fixed and known threshold are also obtained by Park and Padgett [31]. This completes our overview of stochastic process models for the damage process and the marker process – when viewed separately – save for the work of Desmond [32], who articulates on a two-parameter family of life distributions introduced by Birnbaum and Saunders [33]. This distribution is motivated *via* the notion that failure caused by fatigue is due to the initiation, growth, and extension of a dominant crack past some critical length.

The essence of Desmond's [32] idea is based on the notion that it is the environmental stresses called *impulses* that cause a crack to grow, so that if X_i is the size of the crack after the i th impulse, then

$$X_{i+1} = X_i + \Pi_{i+1}g(X_i), i = 0, 1, 2, \dots \quad (12)$$

here Π_i is taken to be the magnitude of the i th impulse; e.g. the stress caused by the i th impulse. The Π_i 's are assumed to be random. If $\nabla X_i = X_{i+1} - X_i$

is taken to be sufficiently small, then

$$\sum_1^n \Pi_i \approx \int_{X_0}^{X_n} \frac{dy}{g(y)} \quad (13)$$

is approximately normal; $g(y)$ is some function of y . The quantity X_0 is the initial size of the dominant crack.

With $g(y) = 1$, and assuming that the Π_i 's have a common mean μ and variance σ^2

$$I(X(t)) \stackrel{\text{def}}{=} \int_{X_0}^{X_n} \frac{dy}{g(y)} \sim N(t\mu, t\sigma^2) \quad (14)$$

If X_c denotes the critical crack size, and T the time to failure of the unit experiencing the impulses, then

$$T = \inf \{t: X(t) > X_c\} \quad (15)$$

Simple manipulations show that

$$P(T \leq t) = \Phi \left(\frac{t\mu - I(X_c)}{\sigma\sqrt{t}} \right) \quad (16)$$

Where $\Phi(\cdot)$ is the unit **normal distribution** function. The distribution function given above is a member of the Birnbaum–Saunders [33] family of distributions.

Acknowledgment

Supported by the US Army Research Office Grant W911NF-05-01-2009 and the Office of Naval Research Contract N00014-06-1-037 with the George Washington University.

References

- [1] Bogdanoff, J.L. & Kozin, F. (1985). *Probabilistic Models of Cumulative Damage*, John Wiley & Sons, New York.
- [2] Sobczyk, K. (1987). Stochastic models for fatigue damage of materials, *Advances in Applied Probability* **19**, 652–673.
- [3] Doksum, K.A. (1991). Degradation models for failure time and survival data, *CWI Quarterly*, Amsterdam **4**, 195–203.
- [4] Doksum, K.A. & Normand, S.L.T. (1995). Gaussian models for degradation processes-part I: methods for the analysis of biomarker data, *Lifetime Data Analysis* **1**(2), 131–144.
- [5] Lu, C.J. & Meeker, W.Q. (1993). Using degradation measures to estimate a time-to-failure distribution, *Technometrics* **35**(2), 161–174.
- [6] Ebrahimi, N. (2004). System reliability based on diffusion models for fatigue crack growth, *Naval Research Logistics* **52**(1), 46–57.
- [7] Lehmann, A. (2006). Degradation-threshold-shock models, in *Probability, Statistics and Modeling in Public Health*, M. Nikulin, D. Commenges & C. Huber, eds, Springer, pp. 286–298.
- [8] Whitmore, G.A., Crowder, M.J. & Lawless, J.F. (1998). Failure inference from a marker process based on a bivariate Wiener process, *Lifetime Data Analysis* **4**, 229–251.
- [9] Lee, M.L.T., Grutolla, V.D. & Schoenfeld, D. (2000). A model for markers and latent health status, *Journal of the Royal Statistical Society. Series B* **62**(4), 747–762.
- [10] Nair, V.N. (1998). Discussion of estimation of reliability in field performance studies by J.D. Kalbfleisch and J.F. Lawless, *Technometrics* **30**, 379–383.
- [11] Cox, D.R. (1999). Some remarks on failure-times, surrogate markers, degradation, wear, and the quality of life, *Lifetime Data Analysis* **5**, 307–314.
- [12] Singpurwalla, N.D. (2006). The hazard potential: introduction and overview, *Journal of the American Statistical Association* **101**(476), 1705–1717.
- [13] Singpurwalla, N.D. (2006). *Reliability and Risk: A Bayesian Perspective*, John Wiley & Sons.
- [14] Singpurwalla, N.D. (2006). On competing risk and degradation processes, in *The Second Erich L. Lehman Symposium – Optimality*, Institute of Mathematical Statistics – Monograph Series, Vol. 49, J. Rojo, ed, The Institute of Mathematical Statistics, pp. 289–304.
- [15] Epstein, B. (1958). The exponential distribution and its role in life testing, *Industrial Quality Control* **15**, 2–7.
- [16] Esary, J.D. (1957). A stochastic theory of accident survival and fatality, Ph.D. dissertation, University of California, Berkeley.
- [17] Gaver Jr, D.P. (1963). Random hazard in reliability problem, *Technometrics* **5**, 211–226.
- [18] Cox, D.R. (1962). *Renewal Theory*, Methuen, London.
- [19] Morey, R.C. (1965). Stochastic wear processes, Technical Report ORC 65-16, Operations Research Center, University of California, Berkeley.
- [20] Esary, J.D., Marshall, A.W. & Proschan, F. (1973). Shock models and wear processes, *Annals of Probability* **1**, 627–649.
- [21] Zacks, S. (2004). Distributions of failure times associated with non-homogeneous compound Poisson damage processes, *A Festschrift for Herman Rubin*, Institute of Mathematical Statistics – Lecture Notes, Vol. 45, pp. 396–407.
- [22] Zacks, S. (2006). Failure distributions associated with general compound renewal damage processes, in *Probability, Statistics and Modeling in Public Health*, M. Nikulin, D. Commenges & C. Huber, eds, Springer, pp. 466–474.

- [23] Cox, D.R. & Miller, H.D. (1965). *The Theory of Stochastic Processes*, Chapman and Hall, London.
- [24] Lemoine, A.J. & Wenocur, M.L. (1985). On failure modeling, *Naval Research Logistics Quarterly* **32**, 497–508.
- [25] Lemoine, A.J. & Wenocur, M.L. (1986). A note on shot-noise and reliability modeling, *Operations Research* **34**, 320–323.
- [26] Wenocur, M.L. (1989). A reliability model based on gamma process and its analytic theory, *Advances in Applied Probability* **21**, 899–918.
- [27] Singpurwalla, N.D. (1997). Gamma processes and their generalizations: an overview, in *Engineering Probabilistic Design and Maintenance for Flood Protection*, R. Cook, M. Mendel & H. Vrijling, eds, Kluwer Academic Publishers, pp. 67–73.
- [28] van der Weid, H. (1997). Gamma processes, in *Engineering Probabilistic Design and Maintenance for Flood Protection*, R. Cook, M. Mendel & H. Vrijling, eds, Kluwer Academic Publishers, pp. 77–83.
- [29] Schabe, H. (1990). A new stochastic model for crack propagation, *Quality and Reliability Engineering International* **6**, 341–344.
- [30] Durham, S.D. & Padgett, W.J. (1997). A cumulative damage model for system failure with application to carbon fibers and composites, *Technometrics* **39**, 34–44.
- [31] Park, C. & Padgett, W.J. (2005). Accelerated degradation models for failure based on geometric Brownian motion and gamma processes, *Lifetime Data Analysis* **11**, 511–527.
- [32] Desmond, A. (1985). Stochastic models of failure in random environments, *The Canadian Journal of Statistics* **13**(2), 171–183.
- [33] Birnbaum, Z.W. & Saunders, S.C. (1969). A new family of life distributions, *Journal of Applied Probability* **6**, 319–327.

Related Articles

Cumulative Damage Models Based on Gamma Processes; Degradation and Failure; Degradation Processes.

NOZER D. SINGPURWALLA

Precedence Tests

Review on Precedence-Type Procedures

In measuring reliability, we observe lifetimes of products or systems of interest. In order to gain a sound knowledge about product or system failure-time distributions, *life-testing and reliability experiments* are carried out before (and while) products are put on the market or systems are used. Life-testing experimentation is a practical way to determine the product reliability and product quality. During the life test, successive times to failure are noted and lifetime data are thus collected. These lifetime data are often used to make decisions on accepting a new design over an existing design. Life testing is useful in many industrial environments including the automobile, materials, chemical, telecommunications, electronics, and pharmaceutical industries. For instance, a manufacturer of electronic components may wish to compare two designs A and B in terms of life. Specifically, he/she may want to abandon design A if there is evidence at the 0.05 level that it has shorter life in comparison to design B.

Assume that a random sample of n_1 units is from F_X (existing design), another independent sample of n_2 units is from F_Y (new design), and that all these sample units are placed simultaneously on a life-testing experiment. We use $X_1, X_2, \dots, X_{n_1}$ to denote the sample from distribution F_X , and $Y_1, Y_2, \dots, Y_{n_2}$ to denote the sample from distribution F_Y . A natural null hypothesis is that the two distributions are equal (i.e., there is no difference between the existing design and the new design) and we are generally concerned with the alternative models where one distribution is stochastically larger than the other; for example, the alternative that F_Y is stochastically larger than F_X (the new design produces units with longer life as compared to the existing design) which can be expressed as

$$H_0 : F_X = F_Y \text{ against } H_1 : F_X > F_Y \quad (1)$$

There has been a considerable amount of research work carried out on this problem from a nonparametric point of view; for more details, see the books [1–5].

In order to save on test time and/or to save on the number of test items, it is common that not all

the items under test will be observed until failure in an industrial setting. In other words, some of the test items will be withdrawn or removed from the life test. This situation is known as *censoring*. For example, in the situation mentioned above with regard to making a decision on accepting a new design over an existing design, usually the cost of producing items from the new design will be much higher than the cost of producing items from the existing design; therefore, we may decide to observe only a few failures (say, the first r failures) from the new design in order to save a number of items that could potentially be used for some other tests. In this scenario, *precedence test* is useful to test the hypotheses in equation (1). The precedence test is a distribution-free two-sample life test based on the order of early failures, and was first proposed by Nelson [6]. The precedence test allows a simple and robust comparison of two distribution functions. Precedence test, based on the number of X failures that precede the r th Y failure, will be useful (a) when life tests involve expensive units since the units that had not failed could be used for some other testing purposes, and (b) to make quick and reliable decisions early on in the life-testing experiment.

We denote the order statistics from the X sample and the Y sample by $X_{1:n_1} \leq X_{2:n_1} \leq \dots \leq X_{n_1:n_1}$ and $Y_{1:n_2} \leq Y_{2:n_2} \leq \dots \leq Y_{n_2:n_2}$, respectively. Without loss of generality, we assume that $n_1 \leq n_2$. Moreover, we let M_1 be the number of X failures before $Y_{1:n_2}$ and M_i be the number of X failures between $Y_{i-1:n_2}$ and $Y_{i:n_2}$, $i = 2, 3, \dots, r$. We denote the observed value of M_i by m_i . It is of interest to mention here that these M_i 's are related to the so-called exceedance statistics whose distributional properties have been discussed, for example, by Fligner and Wolfe [7] and Randles and Wolfe [2].

Nelson [6] provided tables of critical values, which cover all combinations of sample sizes up to 20 for one-sided (two-sided) significance levels of 0.05 (0.10), 0.025 (0.05), and 0.005 (0.01). After Nelson [6] first introduced the precedence life test, many authors have studied the **power** properties of the precedence test and have also proposed some alternatives. For example, Eilbott and Nadler [8] investigated the properties of precedence tests under the assumption that both underlying distributions are exponential. They obtained closed-form expressions for the small-sample and asymptotic power under the exponential distribution. Then, Shorack [9] actually showed that these expressions are valid for a large

2 Precedence Tests

class of distributions. Subsequently, Young [10] presented some asymptotic results for the precedence test. He also applied the precedence sampling to a population selection problem concerning population **quantiles**. Katzenbeisser [11] derived the distribution as well as the first two **moments** of precedence (he used the term *exceedance* instead) test statistics under Lehmann alternatives. Moreover, Katzenbeisser [12] studied the exact power of two-sample location tests based on precedence test statistics against shifts in exponential, logistic, and rectangular distributions. Liu [13] investigated the properties of the precedence probabilities and their applications, while Lin and Sukhatme [14] considered the best precedence test and compared the power of the best precedence test with other nonparametric and **parametric tests**.

Nelson [15] examined the power of the precedence test when the underlying distributions were normal. After Lin and Sukhatme [14] discussed the best precedence test, van der Laan and Chakraborti [16] used their results to derive the best precedence test under Lehmann alternatives. Recently, a book by Balakrishnan and Ng [17] provides a comprehensive overview of theoretical and applied results relating to a wide variety of problems in which precedence test and its alternatives have been used. They have demonstrated the effectiveness of these tests in life-testing situations designed for making quick and reliable decisions in the early stages of a life-testing experiment.

Precedence Test

Following the aforementioned setting and notation, the precedence test statistic $P_{(r)}$ is defined as the number of failures from the X sample that precede

the r th failure from the Y sample; that is,

$$P_{(r)} = \sum_{i=1}^r M_i \quad (2)$$

For example, from Figure 1, with $r = 4$, the precedence test statistic takes on the value $P_{(4)} = 0 + 3 + 4 + 1 = 8$.

It is evident that large values of $P_{(r)}$ lead to the rejection of H_0 in equation (1). For a fixed level of significance α , the critical region will be $\{s, s + 1, \dots, n_1\}$, where

$$\alpha = \Pr(P_{(r)} \geq s | F_X = F_Y) \quad (3)$$

For specified values of n_1, n_2, s , and r , an expression for α in equation (3) is given by

$$\alpha = \frac{\sum_{j=s}^{n_1} \binom{s+r-1}{j} \binom{n_1+n_2-s-r+1}{n_1-j}}{\binom{n_1+n_2}{n_2}} \quad (4)$$

with the summation terminating as soon as any of the factorials involve negative arguments. The critical value s and the exact level of significance α , as close as possible to 5 and 10%, are presented in Balakrishnan and Ng [17] for different choices of the sample sizes n_1 and n_2 and r .

Maximal Precedence Test

Since the precedence test is based on the sum of the frequencies of failures from the X sample between

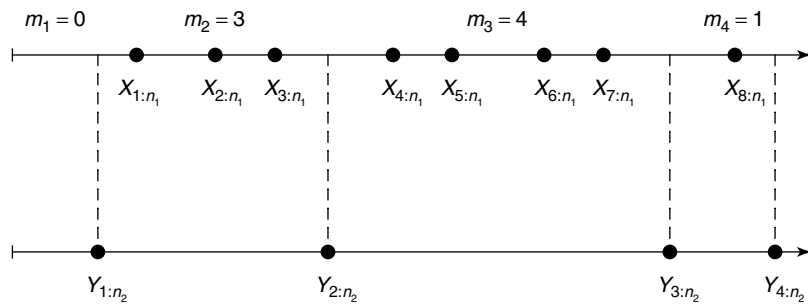


Figure 1 Schematic representation of a precedence life test

the first r failures of the Y sample, it does not take into consideration the distribution of these frequencies of failures. As a result, there will be a *masking effect* present. Balakrishnan and Frattina [18] noted that such a “masking effect” affects the performance of the precedence test and, therefore, proposed a maximal precedence test which is based on the maximum of the numbers of X failures occurring before the first and between the first and second Y -failures. Balakrishnan and Ng [19] generalized this procedure for up to r Y failures which is called *general maximal precedence test*. It is a test procedure based on the maximum number of failures occurring from the X sample before the first, between the first and the second, ..., between the $(r - 1)$ th and the r th failures, from the Y sample. It is a test procedure useful for testing the hypotheses in equation (1) and it avoids the masking effect.

With the same notation as above, the general maximal precedence test statistic is defined as

$$M_{(r)} = \max(M_1, M_2, \dots, M_r) \quad (5)$$

For example, if we refer to Figure 1, with $r = 4$, the maximal precedence test statistic is $M_{(4)} = \max(0, 3, 4, 1) = 4$. Large values of $M_{(r)}$ will lead to the rejection of H_0 and in favor of H_1 in equation (1). The null distribution of $M_{(r)}$, for the special case when $r = 2$, was derived by Balakrishnan and Frattina [18]. Subsequently, Balakrishnan and Ng [19] derived the joint distribution of $M_1, \dots, M_r$, under $H_0 : F_X = F_Y$, as

$$\begin{aligned} \Pr(M_1 = m_1, \dots, M_r = m_r \mid F_X = F_Y) \\ = \frac{\binom{n_1 + n_2 - \sum_{i=1}^r m_i - r}{n_2 - r}}{\binom{n_1 + n_2}{n_2}} \quad (6) \end{aligned}$$

and they used equation (6) to derive the null distribution of the general maximal precedence test statistic $M_{(r)}$ for $r \geq 2$. They examined the power properties of the maximal precedence test and compared them with those of the precedence test and Wilcoxon's rank-sum test (based on complete samples). Through their simulation work, they demonstrated that the maximal precedence test does

avoid the masking effect affecting the precedence test.

Wilcoxon-Type Precedence Test

We can see that even though an alternative to the precedence test has been proposed in the form of maximal precedence test, it too is based on frequencies of failures preceding the r th Y failure. Although Balakrishnan and Ng [19] demonstrated that the maximal precedence test does avoid the masking effect, it still suffers from a loss of power when compared to the classical Wilcoxon's rank-sum test (based on complete samples). For this reason, Ng and Balakrishnan [20, 21] extended the precedence test in another way by constructing Wilcoxon-type rank based precedence test that takes into account the magnitude of the failure times.

The Wilcoxon rank-sum test is a well-known non-parametric testing procedure for testing the hypothesis in equation (1) based on complete samples. For testing the hypotheses in equation (1), if complete samples of sizes n_1 and n_2 were available, then one can use the standard Wilcoxon's rank-sum statistic, proposed by Wilcoxon [22], which is simply the sum of ranks of X observations in the combined sample.

Ng and Balakrishnan [20] proposed the Wilcoxon-type rank-sum precedence tests for testing the hypotheses in equation (1) when the Y sample is Type-II censored. This test is a variation of the precedence test and a generalization of the Wilcoxon rank-sum test. Three Wilcoxon-type rank-sum precedence test statistics – the minimal, maximal, and expected rank-sum statistics – have been proposed by these authors. In order to test the hypotheses in equation (1), instead of using the maximum of the frequencies of failures from the X sample between the first r failures of the Y sample, one could use the sum of the ranks of those failures. More specifically, suppose $m_1, m_2, \dots, m_r$ denote the number of X failures that occurred before the first, between the first and the second, ..., between the $(r - 1)$ th and the r th Y failures, respectively; see Figure 1. Let W be the rank sum of the X failures that occurred before the r th Y failure. The Wilcoxon's test statistic will be smallest when all the remaining $(n_1 - \sum_{i=1}^r m_i)$ X failures occur between the r th and $(r + 1)$ th Y failures. The test statistic in this case

4 Precedence Tests

would be

$$\begin{aligned}
 W_{\min,r} &= W + \left[\left(\sum_{i=1}^r m_i + r + 1 \right) \right. \\
 &\quad \left. + \left(\sum_{i=1}^r m_i + r + 2 \right) + \cdots + (n_1 + r) \right] \\
 &= \frac{n_1(n_1 + 2r + 1)}{2} - (r + 1) \sum_{i=1}^r m_i + \sum_{i=1}^r i m_i
 \end{aligned} \tag{7}$$

This is called the *minimal rank-sum statistic*.

The Wilcoxon's test statistic will be the largest when all the remaining $(n_1 - \sum_{i=1}^r m_i)$ X failures occur after the n_2 th Y failure. Such a test statistic is called the *maximal rank-sum statistic* and is given by

$$\begin{aligned}
 W_{\max,r} &= \frac{n_1(n_1 + 2n_2 + 1)}{2} \\
 &\quad - (n_2 + 1) \sum_{i=1}^r m_i + \sum_{i=1}^r i m_i
 \end{aligned} \tag{8}$$

We could similarly propose a rank-sum statistic using the expected rank sums of failures from the first sample between the r th and the $(r + 1)$ th, ..., after the n_2 th failures of the second sample, denoted by $W_{E,r}$. It can be shown that $W_{E,r}$ is simply the average of $W_{\min,r}$ and $W_{\max,r}$, and is given by

$$\begin{aligned}
 W_{E,r} &= \frac{n_1(n_1 + n_2 + r + 1)}{2} \\
 &\quad - \left(\frac{n_2 + r}{2} + 1 \right) \sum_{i=1}^r m_i + \sum_{i=1}^r i m_i
 \end{aligned} \tag{9}$$

For example, from Figure 1, when $n_1 = n_2 = 10$ with $r = 4$, we have

$$\begin{aligned}
 W_{\min,4} &= 2 + 3 + 4 + 6 + 7 + 8 + 9 \\
 &\quad + 11 + 13 + 14 = 77, \\
 W_{\max,4} &= 2 + 3 + 4 + 6 + 7 + 8 + 9 \\
 &\quad + 11 + 19 + 20 = 89, \\
 W_{E,4} &= \frac{77 + 89}{2} = 83
 \end{aligned} \tag{10}$$

It is evident that small values of $W_{\min,r}$, $W_{\max,r}$, and $W_{E,r}$ lead to the rejection of H_0 and in favor of H_1 in equation (1). Moreover, in the special case of $r = n_2$ (i.e., when we observe all the failures from the Y sample), we have $W_{\min,n_2} = W_{\max,n_2} = W_{E,n_2}$ and in this case they are all equivalent to the classical Wilcoxon's rank-sum statistic mentioned earlier.

Ng and Balakrishnan [21] derived the null distributions of these three Wilcoxon-type rank-sum precedence test statistics based on equation (6), while Ng and Balakrishnan [20] observed that the large-sample normal approximation for the null distribution is not satisfactory in the case of small or moderate sample sizes. For this reason, they developed an Edgeworth expansion to approximate the significance probabilities. They also derived the exact power function under the Lehmann alternative and examined the power properties of the Wilcoxon-type rank-sum precedence tests under a location-shift alternative through extensive **Monte Carlo** simulations.

Weighted Precedence and Maximal Precedence Tests

Weighted precedence and maximal precedence tests for testing the hypotheses in equation (1) is another logical extension of the precedence and maximal precedence tests. The motivation for this is explained in the following two cases based on the example mentioned earlier.

Example 1 A manufacturer of electronic components wishes to compare two designs A and B in terms of their life lengths. Specifically, he/she wishes to abandon design A if there is enough evidence at 5% level of significance that it has shorter life in comparison to design B. $n_1 = 10$ components from design A and $n_2 = 10$ components from design B are placed simultaneously on a life test, and the test is to be terminated when the 5th failure from design B occurs. Figures 2 and 3 show two possible outcomes of the life-testing experiment in this example.

The precedence and maximal precedence test statistics are equal in both cases, and they are $P_{(5)} = 8$ and $M_{(5)} = 5$. We have the critical values with $n_1 = n_2 = 10$ and $r = 5$ as 9 (with level of significance 0.02864) and 6 (with level of significance 0.02709) for the precedence and maximal precedence tests, respectively. Therefore, we will not reject the null

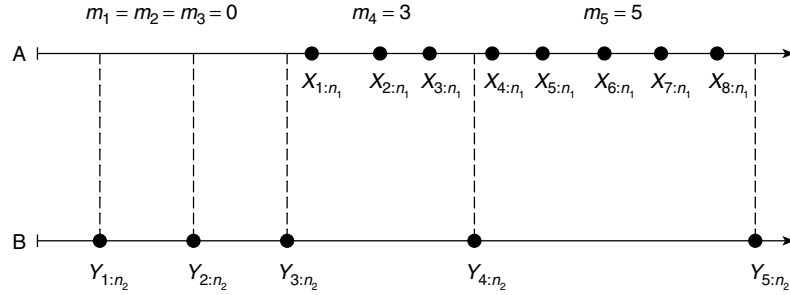


Figure 2 Case 1 of the precedence life test for Example 1

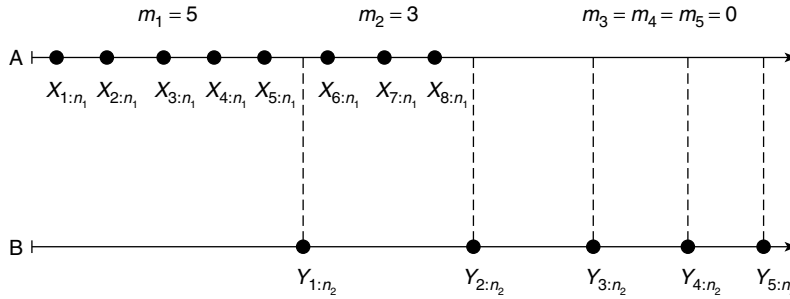


Figure 3 Case 2 of the precedence life test for Example 1

hypothesis that two distributions are equal in both cases at the same level of significance. However, we feel that Case 2 provides much more evidence that design B is better than design A. This suggests that we should develop a test procedure that distinguishes Cases 1 and 2. On this basis, Ng and Balakrishnan [23] proposed the weighted precedence and maximal precedence tests.

The weighted precedence and maximal precedence tests are defined by giving decreasing weights to m_i for increasing i . For example, the weighted precedence test statistic $P_{(r)}^*$ is defined as

$$P_{(r)}^* = \sum_{i=1}^r (n_2 - i + 1) m_i \quad (11)$$

and the weighted precedence test statistic $M_{(r)}^*$ is defined as

$$M_{(r)}^* = \max_{1 \leq i \leq r} \{(n_2 - i + 1) m_i\} \quad (12)$$

For example, from Figure 1, when $n_1 = n_2 = 10$ with $r = 4$, we have

$$\begin{aligned} P_{(4)}^* &= (9 \times 3) + (8 \times 4) + (7 \times 1) = 66, \\ M_{(4)}^* &= \max \{(9 \times 3), (8 \times 4), (7 \times 1)\} = 32 \end{aligned} \quad (13)$$

Once again, it is clear that large values of $P_{(r)}^*$ or $M_{(r)}^*$ would lead to the rejection of H_0 and in favor of H_1 in equation (1).

Ng and Balakrishnan [23] derived the null distributions of these test statistics and the exact power functions under the Lehmann alternative. They also compared the power (under location shift) of the weighted precedence and maximal precedence tests with those of the original precedence and maximal precedence tests.

We note here that in a precedence test, since the life testing continues until the occurrence of the r th Y failure, it may be viewed as a test based on a Type-II right censored Y sample. Therefore, if we want to save the testing units at the early stage of the life test,

a Type-II progressive censoring scheme may instead be employed on the Y sample. For a comprehensive treatment on progressive censoring, one may refer to Balakrishnan and Aggarwala [24]. This will then be a logical extension of the precedence and maximal precedence tests.

Suppose a Type-II progressive censoring scheme is to be adopted on the Y sample which means that a number $r < n_2$ is prefixed for the number of complete failures to be observed and the censoring scheme $(R_1, R_2, \dots, R_r)$ with $R_j \geq 0$ and $\sum_{j=1}^r R_j + r = n_2$ is to be employed at the failure times. During the life-testing experiment, at the time of the j th failure, R_j functioning items are randomly removed from the test. We denote such an observed ordered Y sample by $Y_{1:r:n_2} \leq Y_{2:r:n_2} \leq \dots \leq Y_{r:r:n_2}$. Moreover, we denote by M_1 the number of X failures before $Y_{1:r:n_2}$ and by M_i the number of X failures between $Y_{i-1:r:n_2}$ and $Y_{i:r:n_2}$, $i = 2, 3, \dots, r$.

An extension of the weighted precedence and weighted maximal precedence tests to Type-II progressive censoring has been discussed by Ng and Balakrishnan [23] and Balakrishnan and Ng [17]. Then, based on $M_1, M_2, \dots, M_r$ obtained under Type-II progressive censoring on the Y sample with censoring scheme $(R_1, R_2, \dots, R_r)$, we propose the weighted precedence test statistic $P_{(r)}^*$ (analogous to equation (11)) as

$$P_{(r)}^* = \sum_{i=1}^r \left[n_2 - \left(\sum_{j=1}^{i-1} R_j \right) - i + 1 \right] m_i \quad (14)$$

and the weighted maximal precedence test statistic $M_{(r)}^*$ (analogous to equation (12)) as

$$M_{(r)}^* = \max_{1 \leq i \leq r} \left\{ \left[n_2 - \left(\sum_{j=1}^{i-1} R_j \right) - i + 1 \right] m_i \right\} \quad (15)$$

For example, from Figure 4, with $n_2 = 10$, $r = 4$ and censoring scheme $(2, 1, 1, 2)$, the weighted precedence test statistic becomes $P_{(4)}^* = (10 \times 0) + (7 \times 3) + (5 \times 4) + (3 \times 1) = 44$, while the weighted maximal precedence test statistic becomes $M_{(4)}^* = \max\{(10 \times 0), (7 \times 3), (5 \times 4), (3 \times 1)\} = 21$. It is clear that large values of $P_{(r)}^*$ or $M_{(r)}^*$ would lead to the rejection of H_0 and in favor of H_1 in equation (1).

It is important to mention here that the weighted precedence and maximal precedence test statistics defined earlier in equations (11) and (12) become special cases when we set $R_1 = R_2 = \dots = R_{r-1} = 0$ and $R_r = n_2 - r$.

Exact Power under Lehmann Alternative

The Lehmann alternative $H_1 : [F_X]^\gamma = F_Y$ for some γ is a subclass of the alternative $H_1 : F_X > F_Y$ when $\gamma > 1$; see Lehmann [1] and Gibbons and Chakraborti [5]. The joint probability mass function of $M_1, \dots, M_r$ under the Lehmann alternative is

$$\Pr\{M_1 = m_1, \dots, M_r = m_r \mid [F_X]^\gamma = F_Y\}$$

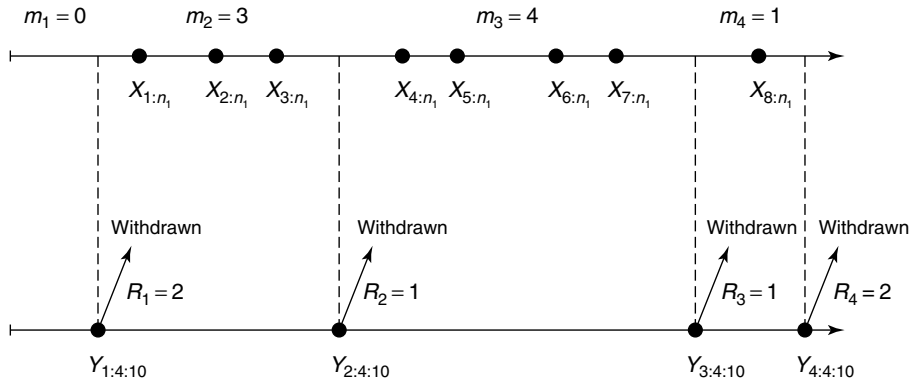


Figure 4 Schematic representation of a precedence life test with progressive censoring

$$\begin{aligned}
&= \frac{n_1!n_2!\gamma^r}{m_1!(n_2-r)!} \left\{ \prod_{j=1}^{r-1} \frac{\Gamma\left(\sum_{i=1}^j m_i + j\gamma\right)}{\Gamma\left(\sum_{i=1}^{j+1} m_i + j\gamma + 1\right)} \right\} \\
&\quad \times \sum_{k=0}^{n_2-r} \binom{n_2-r}{k} (-1)^k \\
&\quad \times \frac{\Gamma\left(\sum_{i=1}^r m_i + (r+k)\gamma\right)}{\Gamma(n_1 + (r+k)\gamma + 1)} \quad (16)
\end{aligned}$$

The power of a test is the probability of rejecting the null hypothesis when the alternative hypothesis is indeed true. So under the Lehmann alternative, the power function for the precedence test and its alternatives can be determined. Note that the null distribution of the test procedures are obtained simply by putting $\gamma = 1$ in the above formulas. We can see that $H_1 : [F_X]^\gamma = F_Y$ is a subclass of the alternative $H_1 : F_X > F_Y$ when $\gamma > 1$.

The power values of the precedence, maximal precedence, weighted precedence, weighted maximal precedence, and Wilcoxon-type rank-sum precedence tests under the Lehmann alternative have all being computed and presented in Balakrishnan and Ng [17]. From those values, it is seen that the Wilcoxon-type rank-sum precedence tests discussed here give higher power values under the Lehmann alternative for all the cases considered.

With a Type-II progressive censoring scheme $(R_1, R_2, \dots, R_r)$ on the Y sample, under the Lehmann alternative $[F_X]^\gamma = F_Y$, $\gamma > 1$, the joint probability mass function of $M_1, M_2, \dots, M_r$ is given by

$$\begin{aligned}
&\Pr(M_1 = m_1, M_2 = m_2, \dots, M_r = m_r \mid [F_X]^\gamma = F_Y) \\
&= C\gamma^r \sum_{j_i (i=1,2,\dots,r)=0}^{R_i} \binom{R_1}{j_1} \binom{R_2}{j_2} \cdots \binom{R_r}{j_r} \\
&\quad \times (-1)^{\left(\sum_{i=1}^r j_i\right)} \\
&\quad \times \left\{ \prod_{k=1}^{r-1} B\left(\sum_{i=1}^k m_i + \gamma\left(\sum_{i=1}^k j_i\right) + k\gamma, m_{k+1} + 1\right) \right\}
\end{aligned}$$

$$\begin{aligned}
&\times B\left(\sum_{i=1}^r m_i + \gamma\left(\sum_{i=1}^r j_i\right) + r\gamma, n_1 - \sum_{i=1}^r m_i + 1\right) \quad (17)
\end{aligned}$$

Note that when $R_1 = R_2 = \dots = R_{r-1} = 0$ and $R_r = n_2 - r$, the joint probability mass function in equation (17) reduces to the expression in equation (16). The exact power values for different choices of n_1, n_2, r , and γ for several progressive censoring schemes have been computed and presented in Balakrishnan and Ng [17].

Selection of the Best Population

When we have two or more competing populations, a problem of interest often is to make a decision as to which is the best if we reject the null (homogeneity) hypothesis in equation (1). In the literature of ranking and selection methodology, parametric and nonparametric procedures for selecting the best among k populations based on samples from these populations have been discussed in great detail; see, for example, Bechhofer *et al.* [25], Gibbons *et al.* [26], and Gupta and Panchapakesan [27]. In order to select the best population based on observing only a few failures from two or more samples under life testing, test procedures based on the nonparametric precedence-type statistics with selecting the best population as an objective in the alternative hypotheses will be useful, especially when budget and/or facility constraints are in place in an industrial setting.

Recently, Balakrishnan *et al.* [28] proposed a nonparametric test procedure based on the precedence statistic to test the equality of two distributions against alternatives that a specified population is better than the other. They also studied the k -sample situation and showed that the procedure works well for different underlying distributions by means of Monte Carlo simulations. Similarly, Ng *et al.* [29] used the minimal Wilcoxon rank-sum precedence statistic instead of the precedence statistic to select the best population based on a test for equality. They also compared the probabilities of correct selection of the minimal Wilcoxon rank-sum precedence statistic with those of the precedence statistic and showed that the performance of minimal Wilcoxon rank-sum precedence statistic is better for selecting the

best population. For more details about the use of precedence-type statistics in ranking and selection, one may refer to the recent book by Balakrishnan and Ng [17].

References

- [1] Lehmann, E.L. (1975). *Nonparametrics: Statistical Methods Based on Ranks*, McGraw-Hill, New York.
- [2] Randles, R.H. & Wolfe, D.A. (1979). *Introduction to the Theory of Nonparametric Statistics*, John Wiley & Sons, New York.
- [3] Hettmansperger, T.P. & McKean, J.W. (1998). *Robust Nonparametric Statistical Methods*, Edward Arnold, London.
- [4] Hollander, M. & Wolfe, D.A. (1999). *Nonparametric Statistical Methods*, 2nd Edition, John Wiley & Sons, New York.
- [5] Gibbons, J.D. & Chakraborti, S. (2003). *Nonparametric Statistical Inference*, 4th Edition, Marcel Dekker, New York.
- [6] Nelson, L.S. (1963). Tables of a precedence life test, *Technometrics* **5**, 491–499.
- [7] Fligner, M.A. & Wolfe, D.A. (1976). Some applications of sample analogues to the probability integral transformation and a coverage property, *The American Statistician* **30**, 78–85.
- [8] Eilbott, J. & Nadler, J. (1965). On precedence life testing, *Technometrics* **7**, 359–377.
- [9] Shorack, R.A. (1967). On the power of precedence life tests, *Technometrics* **9**, 154–158.
- [10] Young, D.H. (1973). A note on some asymptotic properties of the precedence test and applications to a selection problem concerning quantiles, *Sankhyā Series B* **35**, 35–44.
- [11] Katzenbeisser, W. (1985). The distribution of two-sample location exceedance test statistics under Lehmann alternatives, *Statistical Papers* **26**, 131–138.
- [12] Katzenbeisser, W. (1989). The exact power of two-sample location tests based on exceedance statistics against shift alternatives, *Mathematische Operationsforschung und Statistik* **20**, 47–54.
- [13] Liu, J. (1992). Precedence probabilities and their applications, *Communications in Statistics: Theory and Methods* **21**, 1667–1682.
- [14] Lin, C.H. & Sukhatme, S. (1992). On the choice of precedence tests, *Communications in Statistics: Theory and Methods* **21**, 2949–2968.
- [15] Nelson, L.S. (1993). Tests on early failures – the precedence life test, *Journal of Quality Technology* **25**, 140–143.
- [16] van der Laan, P. & Chakraborti, S. (2001). Precedence tests and Lehmann alternatives, *Statistical Papers* **42**, 301–312.
- [17] Balakrishnan, N. & Ng, H.K.T. (2006). *Precedence-Type Tests and Applications*, John Wiley & Sons, Hoboken.
- [18] Balakrishnan, N. & Frattina, R. (2000). Precedence test and maximal precedence test, in *Recent Advances in Reliability Theory: Methodology, Practice, and Inference*, N. Limnios & M. Nikulin, eds, Birkhäuser, Boston, pp. 355–378.
- [19] Balakrishnan, N. & Ng, H.K.T. (2001). A general maximal precedence test, in *System and Bayesian Reliability*, Y. Hayakawa, T. Irony & M. Xie, eds, World Scientific Publishing, Singapore, pp. 105–122.
- [20] Ng, H.K.T. & Balakrishnan, N. (2002). Wilcoxon-type rank-sum precedence tests: large-sample approximation and evaluation, *Applied Stochastic Models in Business and Industry* **18**, 271–286.
- [21] Ng, H.K.T. & Balakrishnan, N. (2004). Wilcoxon-type rank-sum precedence tests, *Australia and New Zealand Journal of Statistics* **46**, 631–648.
- [22] Wilcoxon, F. (1943). Individual comparisons by ranking methods, *Biometrics Bulletin* **1**, 80–83.
- [23] Ng, H.K.T. & Balakrishnan, N. (2002). Weighted precedence and maximal precedence tests and an extension to progressive censoring, *Journal of Statistical Planning and Inference* **135**, 197–221.
- [24] Balakrishnan, N. & Aggarwala, R. (2000). *Progressive Censoring: Theory, Methods and Applications*, Birkhäuser, Boston.
- [25] Bechhofer, R.E., Santner, T.J. & Goldsman, D.M. (1995). *Design and Analysis of Experiments for Statistical Selection, Screening, and Multiple Comparisons*, John Wiley & Sons, New York.
- [26] Gibbons, J.D., Olkin, I. & Sobel, M. (1999). *Selecting and Ordering Populations: A New Statistical Methodology, Classics in Applied Mathematics, Vol. 26*, Society for Industrial and Applied Mathematics, Philadelphia. Unabridged reproduction of the same title, John Wiley & Sons: New York, 1977.
- [27] Gupta, S.S. & Panchapakesan, S. (2002). *Multiple Decision Procedures: Theory and Methodology of Selecting and Ranking Populations, Classics in Applied Mathematics, Vol. 44*, Society for Industrial and Applied Mathematics, Philadelphia. Unabridged reproduction of the same title, John Wiley & Sons: New York, 1979.
- [28] Balakrishnan, N., Ng, H.K.T. & Panchapakesan, S. (2006). A nonparametric procedure based on early failures for selecting the best population using a test for equality, *Journal of Statistical Planning and Inference* **136**, 2087–2111.
- [29] Ng, H.K.T., Balakrishnan, N. & Panchapakesan, S. (2007). Selecting the Best Population Using a Test for Equality Based on Minimal Wilcoxon Rank-sum Precedence Statistic, *Methodology and Computing in Applied Probability* **9** 263–305.

NARAYANASWAMY BALAKRISHNAN AND
HON KEUNG TONY NG

Change-Point Models

Introduction

Change-point models have originally been developed in connection with applications in **quality control**, where a change from the *in-control* to the *out-of-control* state has to be detected based on the available random observations. Up to now, various change-point models have been suggested for a broad spectrum of applications like quality control, reliability, econometrics, or medicine.

The general change-point problem can be described as follows: a random process indexed by time is observed and we want to investigate whether a change in the distribution of the random elements occurs. In other words, we are interested in determining whether the observed stochastic process is homogeneous. Formally, in the discrete-time case, let $X_1, X_2, \dots$ denote a sequence of independent random variables, where the elements $X_1, \dots, X_{\theta-1}$ have an identical distribution function F_0 and $X_\theta, X_{\theta+1}, \dots$ are distributed according to F_1 and the change point θ is unknown. Several statistical tests of the null hypothesis $F_0 = F_1$ against the alternative $F_0 \neq F_1$ for some $\theta > 1$ have been suggested. In addition, estimates for the change point have been proposed and their properties have been investigated. A lot of work has been done in this research field and it is impossible to give an exhaustive overview. The Zentralblatt MATH returns almost 3000 hits, when entering the word “change point”.

Change-point problems can be classified in different ways. Approaches in the classical framework as in the Bayesian framework have been made. Also, there exist models in continuous time as well as in discrete time. Furthermore, the analysis of change points can be partitioned in sequential and *a posteriori* detection models (*ex post* analysis). The problem can also be viewed in a parametric or nonparametric context. Another characterization is whether only one or more than one change points exist. The early work in change-point analysis is described in the survey article of [1]. Other comprehensive reviews are given in [2] and in [3] for nonparametric models. For an overview of limit theorems in change-point problems, we refer to [4]. Change-point methods in **process control** are extensively studied in **Change-point methods**. In this article, we want to concentrate

on a review of models and methods for sequentially observed data and for change points in regression and hazard rate models. We contemplate methods used for sequentially observed data in discrete and continuous time in the first section. The second section presents an overview of change points in regression and hazard rate models. The selection of models follows our personal view.

Detection of a Change Point in Sequentially Observed Data

The aim of so-called disorder or detection problems is to detect the change point “as soon as possible” but avoiding too many false alarms. We distinguish discrete-time and continuous-time models.

Discrete-Time Models

In the discrete case, let $X_1, X_2, \dots$ be a sequence of independent random variables which are observed sequentially. The first $X_1, \dots, X_{\theta-1}$, $\theta > 1$ are distributed according to some known distribution F_0 , while $X_\theta, X_{\theta+1}, \dots$ have some known distribution function $F_1 \neq F_0$. The change point θ is unknown. The time of alarm (a change point has occurred) is determined by a stopping rule which takes the random observations into account. Concerning the change point, there exist Bayesian and non-Bayesian approaches. One of the first non-Bayesian methods is the **CUSUM**-procedure proposed by Page [5], which was further investigated by Lorden [6] and Moustakides [7].

A Bayesian formalization of the disorder problem goes back to Shiryaev [8]. He postulated that the change point θ has a geometric *a priori* distribution with some parameter p and considered the following risk function $R(\tau)$ for stopping at τ : $R(\tau) = P(\tau < \theta) + cE(\tau - \theta)^+$. Here the penalty costs of a false alarm are normalized to 1, and the costs for the delay of stopping after the change point are c per unit time. Now the Bayes stopping rule is to stop at the smallest n for which the posterior probability $\Pi_n = P(\theta \leq n | X_1, \dots, X_n)$ of a change up to n is greater than some threshold A for some $0 < A < 1$.

Continuous-Time Model

Shiryaev [8] was also one of the first to present a model in continuous time in a Bayesian framework:

2 Change-Point Models

the change point is assumed to be a random variable with some prior distribution. Shiryaev considered the following observation process

$$W_t = B_t + r(t - \theta)^+, \quad t \in [0, \infty) \quad (1)$$

where B denotes a standard **Brownian motion**, r is a known fixed constant, and θ is an unknown (random) change point, which is assumed to be independent of B and to have a mixed exponential prior distribution: $P(\theta = 0) = p$ and $P(\theta > t) = (1 - p)e^{-\lambda t}$, $p \in [0, 1)$, $\lambda > 0$, $t \geq 0$. The stopping time τ with respect to the filtration generated by W should signal the change in the drift as soon as possible. The speed of detection is measured by the risk function $R(\tau)$, which is the same as in the discrete-time case. A Bayes solution τ^* should minimize the risk

$$R(\tau^*) = \inf_{\tau} R(\tau) \quad (2)$$

The optimal stopping time τ^* can be determined by means of the posterior distribution $\Pi_t = P(\theta \leq t | \mathcal{F}_t^W)$, with $\mathcal{F}_t^W = \sigma\{W_s : s \leq t\}$. Then the optimal stopping time is

$$\tau^* = \inf\{t > 0 | \Pi_t \geq p^*\} \quad (3)$$

for some properly chosen $p^* \in [0, 1)$. Details about this approach can be found in [9]. An explicit expression for the optimal threshold p^* , and further ramifications can be found in [10, 11].

Another type of change-point problem has been studied in recent years, namely the Poisson disorder problem. Instead of considering a Wiener process with changing drift, a **Poisson process** with changing intensity is observed. Then the problem is to determine a stopping time that signals a change of the intensity of the observed Poisson process. Formally, a point process (T_n) , $n \in \mathbb{N}$ and its corresponding counting process $N_t = \sum_{n=1}^{\infty} I_{\{T_n \leq t\}}$ are observed, where I is the indicator function. At an unknown random time θ , the intensity of N switches from μ_0 to $\mu_1 > \mu_0$. This means that if θ is given, N is a Poisson process with intensity μ_0 up to θ and a Poisson process with intensity μ_1 after θ . [12] described the structure of the solution in the general case. They assumed that θ is a random variable independent of the Poisson process having a mixed exponential prior distribution with parameters p and λ as before. The objective is to stop as soon as possible after θ . The associated risk function

$R(\tau) = P(\tau < \theta) + cE(\tau - \theta)^+$ is given as before and the optimal stopping time in the case $\mu_1 > \mu_0$ is $\tau^* = \inf\{t \geq 0 | \Pi_t \geq B^*\}$, where

$$B^* = \begin{cases} \frac{\lambda}{\lambda + c}, & c \geq \mu_1 - \mu_0 - \lambda \\ \frac{\lambda(\mu_1 - c)}{\lambda\mu_1 + c\mu_0}, & \mu^* \leq c < \mu_1 - \mu_0 - \lambda \end{cases} \quad (4)$$

is the optimal stopping threshold and $\mu^* = \mu_0\mu_1(\mu_0 - \mu_1 - \lambda)/\mu_0\mu_1 + (\mu_1 - \mu_0)(\lambda + \mu_0)$ and (Π_t) , $t > 0$ is the posterior probability process $\Pi_t = P(\theta \leq t | \mathcal{F}_t^N)$, with respect to the filtration $\mathcal{F}_t^N = \sigma\{N_s : s \leq t\}$ and the prior distribution. This stopping time τ^* is obtained by the use of the same methodology as in the Wiener process case. Brown and Zacks [13] extended this result for $\mu^* > c$ using dynamic programming.

A different approach was made by Herberts and Jensen [14]. They solved the detection problem by deriving a semimartingale representation of a gain process whose expectation corresponds to the risk process. In a Bayesian framework, the change point has an exponential prior with parameter λ . The gain process, considered by Herberts and Jensen [14] is a generalization of the one introduced by Shiryaev and is given by

$$Z_t = c_0 \min(t, \theta) + c_1(t - \theta)^+ - k_0(1 - I_{\{\theta \leq t\}}) - k_1 I_{\{\theta \leq t\}} \quad (5)$$

if the process is stopped at time t . There is a gain of c_0 per unit of time and costs k_0 for stopping the process before the change occurs, and a gain c_1 per unit time and stopping costs k_1 afterward. The problem of detecting the change point is then formulated as an optimal stopping problem: find a stopping time ζ with respect to the given observation filtration $\mathbb{F}$ such that $EZ_{\zeta} = \sup_{\tau} \{EZ_{\tau} : \tau \in C^{\mathbb{F}}\}$ where (Z_t) , $t \in \mathbb{R}_+$ is the gain process and $C^{\mathbb{F}}$ is the set of admissible stopping times. Their method goes beyond classical approaches, particularly because it can be used for different information schemes, that is for sequentially observed data, *ex post* decision after observing the point process up to a fixed time t^* , and a combination of both observation schemes. In the sequential case, the observation filtration $\mathbb{F} = (\mathcal{F}_t)$ is the point process history $\mathcal{F}_t = \mathcal{F}_t^N$. In the *ex post* case, the history of N is always known for the whole observation period up to t^* ; therefore, we have in this case $\mathcal{F}_t = \mathcal{F}_{t^*}^N$ for all t , $0 \leq t \leq t^*$. In both cases, the

optimal stopping time can be determined explicitly. In the case of sequential observation, it is given by

$$\zeta^{\mathbb{F}^N} = \inf\{t \in \mathbb{R}_+ : \widehat{Y}_t^{\mathbb{F}^N} \geq z^*\} \quad (6)$$

where $\widehat{Y}_t^{\mathbb{F}^N} = P(\theta \leq t | \mathcal{F}_t^N)$ can be calculated explicitly and $z^* = c_0 + \lambda(k_0 - k_1)/c_0 - c_1 + \lambda(k_0 - k_1)$. In the *ex post* case the solution is a bit more complex and can be found, among other details and more references, in [14].

In recent years, not only homogeneous but also inhomogeneous Poisson process models have been studied. We refer to [15] and [16].

Change Points in Regression and Hazard Rate Models

Regression Models

In the literature two different types of change-point regression models can be found: on the one hand, so-called time-varying regression models, in which the model parameters change at some unknown point in time, and on the other hand, two-phase regression models. Both models are presented briefly in the following.

In the time-varying model, a change in the regression coefficients takes place from the early to the late observations of a sequence (X_n) , $n \in \mathbb{N}$. Let $(X_1, Y_1), \dots, (X_n, Y_n)$ be independent random vectors. Then the model in the random design is given by

$$Y_i = \begin{cases} \alpha_0 + \alpha_1 X_i + \epsilon_i, & i \leq \tau \\ \beta_0 + \beta_1 X_i + \epsilon_i, & i \geq \tau + 1 \end{cases} \quad (7)$$

where (X_i) and (ϵ_i) , $i = 1, \dots, n$ are mutually independent, independent and identically distributed (i.i.d.) sequences with $E(\epsilon_i) = 0$ and $E(\epsilon_i^2) = 1$ and $(\alpha_0, \alpha_1) \neq (\beta_0, \beta_1)$. If $1 \leq \tau \leq n - 1$, then τ is a change point. This design is called *fixed*, if the sequence (X_i) , $i = 1, \dots, n$ is nonstochastic. A two-phase regression model is a regression model with piecewise linear regression functions over two different domains of the design variable. The random design of a two-phase regression model is given by

$$\begin{aligned} Y_i &= (\alpha_0 + \alpha_1 X_i)I_{\{X_i \leq \tau\}} + (\beta_0 + \beta_1 X_i)I_{\{X_i > \tau\}} + \epsilon_i \\ &= m(X_i) + \epsilon_i \end{aligned} \quad (8)$$

The two-phase regression models can be classified further into a restricted and an unrestricted case.

In the restricted case, the regression function f is continuous but not differentiable, whereas in the unrestricted case the regression function is discontinuous. The discontinuity can be expressed in the form of a fixed jump size or a contiguous jump size, in which the jump size tends to zero as the **sample size** tends to infinity.

Hinkley [17] was one of the first authors to investigate a maximum-likelihood estimator (MLE) of the point of intersection for the special case of two line segments under normally distributed errors. A generalization of his model with multiple change points was considered by Feder [18, 19], who investigated least-squares estimates and showed that these estimates are consistent under suitable identifiability assumptions and the asymptotic distributions of these estimates are obtained by “classical” methods.

Koul and Qian [20] and Koul *et al.* [21] considered M-estimators in the unrestricted two-phase random design with a fixed jump size. The M-process corresponding to a function $\phi : \mathbb{R} \rightarrow [0, \infty)$ is defined as

$$M_n(\theta) = \sum_{i=1}^n \phi(Y_i - m(X_i, \theta)) \quad (9)$$

where $m(X; \theta)$ is the linear regression function of model (8) and the M-estimator $\widehat{\theta}$ is given as the minimizer of the M-process:

$$M_n(\widehat{\theta}) = \inf_{\theta} M_n(\theta) \quad a.s. \bullet \quad (10)$$

They showed that the estimate of the jump point converges with rate $O_p(n^{-1})$, whereas the consistency rate of the coefficient parameters is $O_p(n^{-1/2})$. The normalized M-process is asymptotically equivalent to the sum of two processes. One is a quadratic form in the standardized coefficient parameter vector, the other is a jump-point process in the change-point parameter. This result can be exploited to show weak convergence. The suitably standardized M-estimator of the change point converges weakly to the minimizer of a compound Poisson process. The estimates of the regression coefficients are asymptotically normal and independent of the jump-point M-estimator. This is remarkable because the results differ from the restricted and unrestricted contiguous nonrandom design cases.

All models considered above assumed a parametric setting. Of course, there exist various nonparametric

4 Change-Point Models

models as well. Müller [22] studied the following fixed-design nonparametric regression model

$$Y_{in} = g(t_{in}) + \epsilon_{in}, \quad t_{in} \in [0, 1], \quad 1 \leq i \leq n \quad (11)$$

where Y_{in} are noisy measurements of the regression function g taken at points t_{in} and $\epsilon_{in} \sim \mathcal{N}(0, \sigma^2)$ are i.i.d. errors. The assumption is made that there is a change point for the v th derivative $g^{(v)}$ at τ , $0 < \tau < 1$ in the following sense: there exists a function $f \in \mathcal{C}^{k+v}([0, 1])$ with $v \geq 0$ and $k \geq 2$ an even integer, such that

$$g^{(v)}(t) = f^{(v)}(t) + \Delta_v I_{[\tau, 1]}(t), \quad \Delta_v > 0, \quad 0 \leq t \leq 1 \quad (12)$$

The case $\Delta_v < 0$ can be treated analogously. Now, the jump size at the possible change point τ of the v th derivative of g is given by

$$\Delta_v = g_+^{(v)}(\tau) - g_-^{(v)}(\tau) \quad (13)$$

where $g_+^{(v)}(x) = \lim_{y \downarrow x} g^{(v)}(y)$ and $g_-^{(v)}(x) = \lim_{y \uparrow x} g^{(v)}(y)$ are the one-sided limits of the derivative $g^{(v)}(x)$. Hence, the idea is to base the inference of the change points on differences between the left- and right-sided estimates of $g^{(v)}(t)$, which can be done by suitably chosen one-sided kernel estimates. The location of the maximum of these differences is a reasonable estimator of the location of the change point. Let τ be an element of a closed interval $T \subset (0, 1)$. Then the estimator is

$$\hat{\tau} = \inf\{\rho \in T : \hat{\Delta}_v(\rho) = \sup_{x \in T} \hat{\Delta}_v(x)\} \quad (14)$$

In this setting, Müller [22] proved weak convergence of the estimator $\hat{\tau}$. Loader [23] considered a similar nonparametric regression model in which the mean function may have a discontinuity at an unknown point. His estimate is similar in principle to that studied by Müller [22]. However, since he imposed different conditions on the kernel K , his estimate has different properties. It is shown that the change-point estimate converges in probability with rate $O_p(n^{-1})$ and that it has the same asymptotic distribution as maximum-likelihood estimates in parametric models.

The same convergence rate is attained by Müller and Song [24] in a two-step **estimation** of the change point in a nonparametric fixed-design regression model with fixed jump size, whereas the convergence rate in the contiguous case is $O_p(n^{-1} \Delta_n^{-2})$, where

Δ_n is a sequence of jump sizes which tends to zero.

Another important problem in modeling data is the question of whether an unknown function, which cannot parametrically be specified, should be modeled as a globally smooth function or a smooth function with isolated change points. Müller and Stadtmüller [25] proposed statistics, which provide relevant information for this decision.

Hazard Rate Models

Hazard rate models often occur in medical follow-up studies after major surgery. The simplest one with a change point can be expressed as follows

$$\lambda(t) = \begin{cases} \lambda_1 & t \leq \tau \\ \lambda_2 & t > \tau \end{cases}, \quad t \geq 0, \quad \tau \geq 0 \quad (15)$$

with constants $\lambda_1, \lambda_2 > 0$ and change point τ . A first attempt to estimate these three parameters was made by Anderson and Senthilselvan [26]. They investigated, as a special case of this simple model, an extended **Cox model** with $\lambda_1 = e^{\alpha^T Z}$, $\lambda_2 = e^{\gamma^T Z}$, where α and γ are parameter vectors and Z is a vector of covariates. The parameters are estimated by the conditional log likelihood given the value of τ and then the baseline hazard $\lambda(t)$ is estimated by a penalized maximum-likelihood method conditioning on the parameter estimates. Liang *et al.* [27] proposed a slightly different Cox model

$$\lambda(t) = \lambda_0(t) \exp\{(\beta + \theta I_{[t \leq \tau]})Z + \gamma^T X\} \quad (16)$$

where Z is a one-dimensional **covariate** which should be included in possibly different magnitudes over time, X is another confounding covariate vector and the change point at an unknown time is given by τ . They tested the hypothesis of $H_0 : \theta = 0$ using a test statistic

$$M = \sup_{\tau \in [a, b]} S(\tau) \quad (17)$$

where S is a function of the first two derivatives of the partial log-likelihood function with respect to β and γ .

A further Cox model is presented by Luo and Boyett [28]:

$$\lambda(t) = \lambda_0(t) \exp\{\beta I_{[X \leq \theta]} + \alpha^T Z\} \quad (18)$$

where a constant is added to a covariate beyond a threshold, which is characterized by a random variable X . They proved consistency of their partial MLE.

A Cox model for i.i.d. censored survival times with a change point according to the unknown threshold of a covariate was introduced by Pons [29]:

$$\lambda(t) = \lambda_0(t) \exp\{\alpha^T \mathbf{Z}_1(t) + \beta^T \mathbf{Z}_2(t) I_{\{Z_3 \leq \zeta\}} + \gamma^T \mathbf{Z}_2(t) I_{\{Z_3 > \zeta\}}\} \quad (19)$$

In this model, it is shown that the partial MLE of the change point ζ is n consistent, i.e. the convergence rate is $O_p(n^{-1})$. Such a convergence rate was also attained for the change point in the unrestricted two-phase random linear regression design with a fixed jump size (see [20]). Furthermore, Pons [29] proved that the estimates of the regression parameter vectors α, β, γ are $n^{1/2}$ consistent and that $n(\hat{\zeta}_n - \zeta)$ converges weakly to a random variable $\hat{v}_Q$, which is a maximizer of a certain jump process. The estimates of the regression parameters are asymptotically normal.

Gandy *et al.* [30] investigated a similar model. But instead of a jump there is a smooth change at the change point ξ :

$$\lambda_i(t, \theta) = \lambda_0(t) R_i(t) \exp\{\beta_1^T \mathbf{Z}_{1i}(t) + \beta_2 Z_{2i} + \beta_3 (Z_{2i} - \xi)^+\} \quad (20)$$

where $\theta = (\xi, \beta^T)^T$ with $\beta = (\beta_1^T, \beta_2, \beta_3)^T \in \mathcal{B} \subset \mathbb{R}^{p+2}$ is the vector of regression parameters and R_i the at-risk indicator. The change point is denoted by ξ lying in a closed interval of known parameters $[\xi_1, \xi_2]$. For brevity, $\tilde{\mathbf{Z}}_i(t; \xi) = (\mathbf{Z}_{1i}^T(t), Z_{2i}, (Z_{2i} - \xi)^+)^T$ is considered resulting in $\lambda_i(t, \theta) = \lambda_0(t) R_i(t) \exp\{\beta^T \tilde{\mathbf{Z}}_i(t; \xi)\}$. In this new, extended Cox model θ is estimated by the value $\hat{\theta}_n$ that maximizes the logarithm of the partial likelihood:

$$\begin{aligned} \log L(\theta) = & \sum_{i=1}^n \int_0^\tau \beta^T \tilde{\mathbf{Z}}_i(t; \xi) dN_i(t) \\ & - \int_0^\tau \log \left(\sum_{i=1}^n R_i(t) \exp(\beta^T \tilde{\mathbf{Z}}_i(t; \xi)) \right) \\ & \times d \left(\sum_{i=1}^n N_i(t) \right) \end{aligned} \quad (21)$$

The maximization can be carried out in the following way:

For fixed ξ , let $\hat{\beta}_n(\xi) = \arg\max_{\beta \in \mathcal{B}} \log L(\xi, \beta)$ and $\log L(\hat{\xi}) = \log L(\hat{\xi}, \hat{\beta}_n(\hat{\xi}))$. Then $\hat{\xi}$ can be estimated by $\hat{\xi}_n$ satisfying

$$\hat{\xi}_n = \inf \left\{ \xi \in [\xi_1, \xi_2] : \log L(\xi) = \sup_{\xi \in [\xi_1, \xi_2]} \log L(\xi) \right\} \quad (22)$$

The partial MLE of θ is $\hat{\theta}_n = (\hat{\xi}_n, \hat{\beta}_n)$, where $\hat{\beta}_n = \hat{\beta}_n(\hat{\xi}_n)$.

It is proved that the partial MLE of all parameters is $n^{1/2}$ consistent. In contrast to Pons' model $n^{1/2}(\hat{\xi} - \xi_0)$ is asymptotically normal. The proof uses results of the theory of M -estimators as presented in [31]. The model can be extended further by substituting a vector $\mathbf{Z}_2$ for Z_2 , which may also be time dependent, and a vector ξ for ξ allowing more than just one change point. The asymptotic results are still the same.

Other approaches relying on the Cox model have recently been suggested by Dupuy [32]. He considered a model with a change point in both hazard and regression parameters. Estimates of the change point, hazard, and regression parameters are proposed and shown to be consistent.

Acknowledgments

Financial support of the *Deutsche Forschungsgemeinschaft* through Research Unit 460 is gratefully acknowledged.

References

- [1] Zacks, S. (1982). Classical and Bayesian approaches to the change-point problem: fixed sample and sequential procedures, *Statistical Analysis Donnees* 7, 48–81.
- [2] Bhattacharya, P. (1994). Some aspects of change-point analysis, in *Change-Point Problems* E. Carlstein, H. Müller & D. Siegmund, eds, Institute of Mathematical Statistics, Hayward, pp. 28–56.
- [3] Brodsky, B. & Darkhovsky, B. (1993). *Nonparametric Methods in Change-Point Problems*, Kluwer Academic Publishers, Dordrecht.
- [4] Csörgő, M. & Horváth, L. (1997). *Limit Theorems in Change-Point Analysis*, John Wiley & Sons, Chichester.
- [5] Page, E. (1954). Continuous inspection schemes, *Biometrika* 41, 100–115.
- [6] Lorden, G. (1971). Procedures for reacting to a change in distribution, *Annals of Mathematical Statistics* 42, 1897–1908.

- [7] Moustakides, G. (1986). Optimal stopping times for detecting changes in distribution, *Annals of Statistics* **14**, 1379–1387.
- [8] Shiryaev, A. (1963). On optimum methods in quickest detection problems, *Theory of Probability and its Applications* **8**, 22–46.
- [9] Shiryaev, A. (1978). *Optimal Stopping Rules*, Springer, New York.
- [10] Beibel, M. (1994). Bayes problems in change-point models for the Wiener process, in *Change-Point Problems*, E. Carlstein, H. Müller & D. Siegmund, eds, Institute of Mathematical Statistics, Hayward, pp. 1–6.
- [11] Beibel, M. (1996). A note on Ritov's Bayes approach to the minimax property of the CUSUM-procedure, *Annals of Statistics* **24**, 1804–1812.
- [12] Peskir, G. & Shiryaev, A. (2002). Solving the Poisson disorder problem, in *Advances in Finance and Stochastics*, K. Sandmann & P. Schönbacher, eds, Springer, Berlin, pp. 295–312.
- [13] Brown, M. & Zacks, S. (2006). A note on optimal stopping for possible change in the intensity of an ordinary Poisson process, *Statistics and Probability Letters* **76**, 1417–1425.
- [14] Herberts, T. & Jensen, U. (2004). Optimal detection of a change point in a Poisson process for different observation schemes, *Scandinavian Journal of Statistics* **31**, 347–366.
- [15] Dabrye, A., Farinetto, C. & Kutoyants, Y. (2003). On Bayesian estimators in misspecified change-point problems for Poisson process, *Statistics and Probability Letters* **61**, 17–30.
- [16] Ruggeri, F. & Sivaganesan, S. (2005). On modeling change points in non-homogeneous Poisson processes, *Statistical Inference for Stochastic Processes* **8**, 311–329.
- [17] Hinkley, D. (1971). Inference in two-phase regression, *Journal of the American Statistical Association* **66**, 736–743.
- [18] Feder, P. (1975). The log likelihood ratio in segmented regression, *Annals of Statistics* **3**, 84–97.
- [19] Feder, P. (1975). On asymptotic distribution theory in segmented regression problems, *Annals of Statistics* **3**, 49–83.
- [20] Koul, H. & Qian, L. (2002). Asymptotics of maximum likelihood estimator in a two-phase linear regression model, *Journal of Statistical Planning and Inference* **108**, 99–119.
- [21] Koul, H., Qian, L. & Surgailis, D. (2003). Asymptotics of M-estimators in two-phase linear regression models, *Stochastic Processes and their Applications* **103**, 123–154.
- [22] Müller, H.-G. (1992). Change-points in nonparametric regression analysis, *Annals of Statistics* **20**, 737–761.
- [23] Loader, C. (1996). Change point estimation using nonparametric regression, *Annals of Statistics* **24**, 1667–1678.
- [24] Müller, H.-G. & Song, K.-S. (1997). Two-stage change-point estimators in smooth regression, *Statistics and Probability Letters* **34**, 323–335.
- [25] Müller, H.-G. & Stadtmüller, U. (1999). Discontinuous versus smooth regression, *Annals of Statistics* **27**, 299–337.
- [26] Anderson, J. & Senthilselvan, A. (1982). A two-step regression model for hazard functions, *Applied Statistics* **31**, 44–51.
- [27] Liang, K.-Y., Self, S. & Liu, X. (1990). The Cox proportional hazard model with change-point: an epidemiologic application, *Biometrics* **46**, 783–793.
- [28] Luo, X. & Boyett, J. (1997). Estimations of a threshold parameter in Cox regression, *Communications in Statistics: Theory and Methods* **26**, 2329–2346.
- [29] Pons, O. (2003). Estimation in a Cox regression model with a change-point according to a threshold in a covariate, *Annals of Statistics* **31**, 442–463.
- [30] Gandy, A., Jensen, U. & Lütkebohmert, C. (2005). A Cox model with a change-point applied to an actuarial problem, *Brazilian Journal of Probability and Statistics* **19**, 93–109.
- [31] Van der Vaart, A. (1998). *Asymptotic Statistics*, Cambridge University Press, New York.
- [32] Dupuy, J.-F. (2006). Estimation in a change-point hazard regression model, *Statistics and Probability Letters* **76**, 182–190.

Related Article

Change-Point Methods.

UWE JENSEN AND CONSTANZE
LÜTKEBOHMERT

System Signatures

Introduction

Structural reliability is a subfield of reliability theory which is concerned with the precise way in which systems are designed, that is, the way that the system's components work individually, and in tandem with others, to make the system itself function or fail. In a seminal paper in the field, Birnbaum *et al.* [1] defined a coherent system as one which is monotonic (i.e., fixing a failed component cannot cause a system to fail) and in which every component is relevant (i.e., for every i , there's at least one configuration of component states in which the i^{th} component's failure will make the system fail). They argued persuasively that these were the only systems meriting further study, as any incoherent system could be replaced by a coherent system of the same or smaller size having performance characteristics as good as or better than the original system.

For a system with n components, Birnbaum *et al.* [1] defined the state vector $\mathbf{x} = (x_1, x_2, \dots, x_n) \in \{0, 1\}^n$ of the system's components as the vector of 0s and 1s such that $x_i = 1$ if the i^{th} component is working and $x_i = 0$ if it is not working. The structure function $\varphi : \{0, 1\}^n \rightarrow \{0, 1\}$ of a given system is then defined as a mapping that associates those state vectors $\mathbf{x}$ for which the system works with the value 1 and those state vectors $\mathbf{x}$ for which the system fails with the value 0.

These two ideas – coherence and structure – form the foundation of modern structural reliability and, at least in theory, allow one to study the performance of a given system and to compare any pair of systems of interest. They lead, for example, to the calculation of the reliability $h(\mathbf{p})$ of a system in independent components for which $P(X_i = 1) = p_i$ for $i = 1, 2, \dots, n$; specifically, $h(\mathbf{p}) = E\varphi(\mathbf{X})$. The reliability function $h(\mathbf{p})$ is multilinear, that is, linear in each p_i , and when $p_i \equiv p$ for all i , a case of special interest, it can be written as the n^{th} degree polynomial referred to as the *reliability polynomial* for the system. For further detail on these ideas and their consequences, see Barlow and Proschan [2].

While the ideas discussed above are without doubt two of the more important conceptual breakthroughs in the mathematical theory of reliability, as tools for the computation of system reliability and for

the comparative study of different system designs, they have a number of limitations. As informative as structure functions are in determining how systems work, the class has the same cardinality as the class of **coherent systems** themselves and so their number grows exponentially in n . Further, they are complex algebraic expressions that usually have many equivalent forms. It is thus awkward, and often unfeasible, to use them as a way of indexing and/or comparing systems. In a 1985 paper, Samaniego introduced an alternative index which, although less general than structure functions, has the virtues of being both quite useful and easily interpreted. For systems of order n , this index is of fixed dimension; in fact, it resides in a bounded simplex in n -dimensional Euclidean space. This index, called the *system's signature*, is defined as follows.

Definition 1 Let τ represent a coherent system of order n . Assume that the system's n components are independent and identically distributed (i.i.d) according to the (continuous) lifetime distribution F . The *signature* of the system τ , denoted by $\mathbf{s}_\tau$, or simply by $\mathbf{s}$ when the corresponding system is clear from the context, is an n -dimensional probability vector whose i^{th} element s_i is equal to the probability that the i^{th} component failure causes the system to fail. In brief, $s_i = P(T = \mathbf{X}_{i:n})$ for $i = 1, 2, \dots, n$, where T is the lifetime distribution of the system and $\mathbf{X}_{1:n} < \mathbf{X}_{2:n} < \dots < \mathbf{X}_{n:n}$ are the ordered failure times of the system's n components.

Some comments on the i.i.d. assumption on component lifetimes would seem to be in order. The primary use of signatures to date has been the comparison of system designs. It is quite apparent that a comparison between two systems with quite different component characteristics may well be misleading or inconclusive. For example, a series system in four highly reliable components is likely to outperform a four-component parallel system with relatively poor components. But it's clear that **parallel systems** are preferable to **series systems** of the same order. Indeed, the two systems' structure functions are uniformly ordered in $\mathbf{x}$. On the other hand, once the i.i.d. assumption is made, any remaining differences in system performance must be attributable to the system's design. In that sense, the assumption levels the playing field so that one has a basis for comparing the designs themselves. From an analytical

2 System Signatures

point of view, signatures, as defined above, make available for use the tools of combinatorial mathematics for computational purposes and the well-known distribution theory for the order statistics of an i.i.d. sample for studying the performance of a particular system. It should, of course, be acknowledged that specific applications of signature-related results in non-i.i.d. settings must be done, if at all, with caution. This caveat notwithstanding, it is also worth adding that results based on an i.i.d. assumption, while being inexact in such applications, are not necessarily irrelevant. If, for example, the component lifetimes could be considered independent and, while not identically distributed, nonetheless roughly comparable, it seems reasonable to conclude from the results cited in the sequel that selecting the system with a better signature should lead to better performance. Naturally, an exact analysis is desirable, but it will often be intractable. The signature-related results to be discussed here should be of some use in some ‘neighborhood’ of the i.i.d. setting that has been studied via system signatures.

The following simple example illustrates the type of reasoning which proves useful in the computation of system signatures. Consider the system pictured in Figure 1 below. The failure times $\mathbf{X}_1$, $\mathbf{X}_2$, and $\mathbf{X}_3$ of the three components of this system can be ordered in $3! = 6$ ways, and these six possible permutations are equally likely due to the i.i.d. assumption. The ‘order-statistic equivalent’ for the system failure time T is shown in Table 1 below for each permutation of the component failure times.

Clearly, the system in Figure 1 has signature vector $\mathbf{s} = (1/3, 2/3, 0)$. Note that in this example, and in general, the signature of a system depends only on the permutation distribution of the n failure times and does not depend in any way on the continuous lifetime distribution F . The five distinct coherent systems of order 3 are easily seen to have signatures $(1, 0, 0)$, $(0, 1, 0)$, $(0, 0, 1)$, $(1/3, 2/3, 0)$, and $(0, 2/3, 1/3)$.

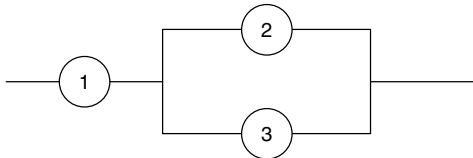


Figure 1 A three-component system with structure function $\varphi^*(\mathbf{x}) = x_1(x_2 + x_3 - x_2x_3)$

Table 1 The ordered component failure time which causes failure of system φ^* above

Ordered component failure times	Order statistic equal to system failure time T
$\mathbf{X}_1 < \mathbf{X}_2 < \mathbf{X}_3$	$\mathbf{X}_{1:3}$
$\mathbf{X}_1 < \mathbf{X}_3 < \mathbf{X}_2$	$\mathbf{X}_{1:3}$
$\mathbf{X}_2 < \mathbf{X}_1 < \mathbf{X}_3$	$\mathbf{X}_{2:3}$
$\mathbf{X}_2 < \mathbf{X}_3 < \mathbf{X}_1$	$\mathbf{X}_{2:3}$
$\mathbf{X}_3 < \mathbf{X}_1 < \mathbf{X}_2$	$\mathbf{X}_{2:3}$
$\mathbf{X}_3 < \mathbf{X}_2 < \mathbf{X}_1$	$\mathbf{X}_{2:3}$

Signature-Based Representations

Samaniego [3] showed that the distribution of the system lifetime T , given i.i.d. components lifetimes $\sim F$, can be written in terms of $\mathbf{s}$ and F alone:

Theorem 1 Let $\mathbf{X}_1, \dots, \mathbf{X}_n$ be the i.i.d. component lifetimes of a coherent system of order n , and let T be the system lifetime. Then

$$\bar{F}_T(t) \equiv P(T > t) = \sum_{j=0}^{n-1} \left(\sum_{i=j+1}^n s_i \right) \binom{n}{j} \times (F(t))^j (\bar{F}(t))^{n-j} \quad (1)$$

Assuming i.i.d components and letting t_0 be a fixed point in time, set $p = \bar{F}(t_0)$ and $q = F(t_0)$. The reliability of the system at $t = t_0$, or equivalently, the reliability polynomial $h(p)$, may then be written as

$$h(p) = \sum_{j=1}^n \left(\sum_{i=n-j+1}^n s_i \right) \binom{n}{j} p^j q^{n-j} \quad (2)$$

It is evident from the representation in (1) that the survival function of T can also be written as

$$P(T > t) = \sum_{i=1}^n s_i P(\mathbf{X}_{i:n} > t) \quad (3)$$

Integrating both sides of equation (3) over $t \in (0, \infty)$, we obtain the useful identity

$$ET = \sum_{i=1}^n s_i E\mathbf{X}_{i:n} \quad (4)$$

Samaniego [3] also established the following representations for the density and failure rate of T when

F is absolutely continuous:

$$f_T(t) = \sum_{i=1}^n i s_i \binom{n}{i} (F(t))^{i-1} (\bar{F}(t))^{n-i} f(t) \quad (5)$$

and

$$r_T(t) = \frac{\sum_{i=1}^n i s_i \binom{n}{i} (F(t))^{i-1} (\bar{F}(t))^{n-i+1}}{\sum_{i=1}^n s_i \sum_{j=0}^{i-1} \binom{n}{j} (F(t))^j (\bar{F}(t))^{n-j}} r(t) \quad (6)$$

where f and r are the density and failure rate of the underlying distribution F .

The following is a brief overview of some of the main applications to date of signatures in reliability. The first of these results shows that certain stochastic relationships enjoyed by pairs of system signatures are preserved by the lifetimes of the corresponding systems. For definitions of stochastic, hazard rate and likelihood ratio ordering between random variables (or their distributions), see Shaked and Shanthikumar [4]. In the following result, due to Kochar *et al.* [5], signatures are seen as the distributions of discrete variables taking values in $\{1, 2, \dots, n\}$.

System Comparisons

Theorem 2 Let $\mathbf{s}_1$ and $\mathbf{s}_2$ be the signatures of the two systems of order n , both based on components with i.i.d. lifetimes with common distribution F . Let T_1 and T_2 be the respective system lifetimes. Then (a) $\mathbf{s}_1 \leq_{\text{st}} \mathbf{s}_2 \Rightarrow T_1 \leq_{\text{st}} T_2$, (b) $\mathbf{s}_1 \leq_{\text{hr}} \mathbf{s}_2 \Rightarrow T_1 \leq_{\text{hr}} T_2$, and (c) $\mathbf{s}_1 \leq_{\text{lr}} \mathbf{s}_2 \Rightarrow T_1 \leq_{\text{lr}} T_2$.

Block *et al.* [6] noted that the signature conditions in Theorem 2 are sufficient but not necessary, and generalized the result, giving explicit necessary and sufficient conditions (NASCs) for $T_1 \leq T_2$ in each of the three stochastic senses considered above.

While stochastic relationships between signatures are clearly useful tools in the comparison of competing systems, they provide only a partial ordering among systems. Kochar *et al.* [5] give examples of systems that are noncomparable according to the standard orderings. There is, however, a metric under which any two systems are comparable. Arcones *et al.* [7] introduced the notion of ‘stochastic

precedence’, an alternative way to quantify the fact that one random variable is smaller than another. They defined the ‘sp’ relationship as follows.

Definition 2 Let Y_1 and Y_2 be independent random variables with respective distributions F_1 and F_2 . Then Y_1 is said to *stochastically precede* Y_2 (written $Y_1 \leq_{\text{sp}} Y_2$) if and only if $P(Y_1 \leq Y_2) \geq 1/2$. The variables are *equivalent* if both $Y_1 \leq_{\text{sp}} Y_2$ and $Y_2 \leq_{\text{sp}} Y_1$ hold. Continuous variables Y_1 and Y_2 are sp-equivalent if and only if $P(Y_1 \leq Y_2) = 1/2$.

Hollander and Samaniego [8] consider the comparison of T_1 and T_2 via stochastic precedence. As we shall see, $P(T_1 \leq T_2)$ can be computed in closed form under an i.i.d. assumption, and one may then classify the second system as better than, equivalent to or worse than the first system according to whether $P(T_1 \leq T_2)$ is greater than, equal to or less than $1/2$. It is therefore apparent that any two systems of arbitrary order are comparable via stochastic precedence. Below, an explicit formula for calculating the probability $P(T_1 \leq T_2)$ when all components have i.i.d. lifetimes with continuous distribution F . Hollander and Samaniego [8] establish the following results.

Theorem 3 Consider coherent systems τ_1 and τ_2 of orders n and m respectively, and assume that their component lifetimes are i.i.d. according to continuous distributions F_1 and F_2 , respectively. Let $\mathbf{s}_1$ and $\mathbf{s}_2$ be their signatures and T_1 and T_2 be their lifetimes. Then

$$P(T_1 \leq T_2) = \sum_{i=1}^n \sum_{j=1}^m s_{1i} s_{2j} P(\mathbf{X}_{i:n} \leq \mathbf{Y}_{j:m}) \quad (7)$$

The following computational formula for $P(T_1 \leq T_2)$, applicable when $F_1 = F_2$, follows from the explicit formula obtained by Kvam and Samaniego [9] for the probability of a given ordering among independently drawn order statistics.

Theorem 4 Under the conditions of Theorem 3, and given that $F_1 = F_2$,

$$P(T_1 \leq T_2) = \sum_{i=1}^n \sum_{j=1}^m s_{1i} s_{2j} \sum_{k=i}^n \left[\frac{\binom{n}{k} \binom{m}{j}}{\binom{n+m}{k+j}} \right] \times \left[\frac{j}{(k+j)} \right] \quad (8)$$

Example It is easy to verify that the two four-component systems with respective minimal cut sets $\{\{1\}, \{2, 3, 4\}\}$ and $\{\{1, 2\}, \{1, 3\}, \{1, 4\}, \{2, 3\}\}$ are noncomparable in the ‘st’ sense (or, of course, also in the ‘hr’ or ‘lr’ senses). The first of these systems has signature $\mathbf{s}_1 = (1/4, 1/4, 1/2, 0)$ and the second has signature $\mathbf{s}_2 = (0, 2/3, 1/3, 0)$. Assuming that all components are i.i.d. $\sim F$, one obtains from equation (8) that $P(T_1 \leq T_2) = 0.519$. From this one may conclude that the second system will last longer than the first, slightly more than half the time. Thus, in the ‘sp’ ordering, the second system is to be preferred to the first.

The expression in (8) is clearly distribution free, immediately yielding NASCs for stochastic precedence between two systems of arbitrary order with i.i.d. component lifetimes. Indeed, if the RHS of equation (8) is denoted by $W(\mathbf{s}_1, \mathbf{s}_2)$, then

$$P(T_1 \leq T_2) > \frac{1}{2} \left(\text{or} = \frac{1}{2} \text{ or } < \frac{1}{2} \right) \text{ if and only if} \\ W(\mathbf{s}_1, \mathbf{s}_2) > \frac{1}{2} \left(\text{or} = \frac{1}{2} \text{ or } < \frac{1}{2} \right) \quad (9)$$

Related Literature

There are a variety of further applications of system signatures which merit discussion. The following is a brief summary. Using the representation of the system failure rate in equation (6), Samaniego [3] provides a characterization, in terms of system signatures, of systems whose lifetime distributions have an increasing failure rate whenever their components have i.i.d. lifetimes with an increasing failure rate. Boland [10] derives the signatures of indirect majority systems and executes a comparison of such systems with simple majority systems of the same size. Shaked and Suarez-Llorens [11] use conditions on system signatures to obtain comparisons of systems *via* the information and the dispersive orderings. Belzunce and Shaked [12] obtain results for pairs of signatures with specific relationships *via* their theory of failure profiles. Navarro and Shaked [13] utilize signatures in showing hazard-rate ordering among independently drawn minima $\{\mathbf{X}_{1:1}, \mathbf{X}_{1:2}, \dots, \mathbf{X}_{1:n}, \dots\}$ and in studying the limiting behavior of the failure rates of selected systems. Dugas and Samaniego [14] study optimization problems in reliability economics

and identify optimal designs (i.e., optimal stochastic mixtures of coherent systems) relative to a class of criterion functions of the form $m(\bar{a}, \bar{c}, \bar{s})$, where $\bar{a}$ is a measure of the system’s performance, $\bar{c}$ is a measure of the system’s cost and $\bar{s}$ is the system’s signature.

Satyanarayana and Prabhakar [15] introduced an important computational approach called *domination theory* aimed at simplifying the use of the inclusion-exclusion formula in the computation of the reliability of complex communication networks. While this tool has been a substantial boon computationally, the method does not appear to have any clear interpretive content. For a coherent system in i.i.d. components, it provides an efficient way to obtain the reliability polynomial

$$h(\mathbf{p}) = \sum_{j=1}^n d_j p^j \quad (10)$$

where $\mathbf{d} = (d_1, d_2, \dots, d_n)$ is the vector of ‘signed dominations’. Boland *et al.* [16] explicitly identify the relationship between the domination and signature vectors in the form $\mathbf{s} = \mathbf{Q}\mathbf{d}$, where $\mathbf{Q}$ is a specific $n \times n$ matrix (thereby reconciling equations (2) and (10)). They illustrate the correspondence and joint utility of $\mathbf{s}$ and $\mathbf{d}$ by demonstrating that the lifetimes of two complex networks (both with 9 vertices and 27 edges) are hazard-rate ordered. The work facilitates the concurrent use of dominations for reliability computation and signatures for comparative analysis of systems or networks.

Details on much of the work described in this article can be found in the cited references, and in particular, in the monograph by Samaniego [17].

References

- [1] Birnbaum, Z.W., Esary, J.D. & Saunders, S.C. (1961). Multicomponent systems and structures, and their reliability, *Technometrics* **3**, 55–77.
- [2] Barlow, R.E. & Proshan, F. (1981). *Statistical Theory of Reliability*, McArdle Press, Silver Springs.
- [3] Samaniego, F.J. (1985). On closure of the IFR class under formation of coherent systems, *IEEE Transactions on Reliability Theory* **TR-34**, 69–72.
- [4] Shaked, M. & Shanthikumar, J.G. (1994). *Stochastic Orders and Their Applications*, Academic Press, San Diego.
- [5] Kochar, K., Mukerjee, H. & Samaniego, F.J. (1999). The ‘signature’ of a coherent system and its application to comparisons among systems, *Naval Research Logistics* **46**, 507–523.

-
- [6] Block, H., Dugas, M. & Samaniego, F.J. (2007). *Signature-related results on system failure rates and lifetimes*, in *Advances in Statistical Modeling and Inference: Essays in Honor of Kjell Doksum*, V.N. Nair, ed, Singapore: World Scientific, 115–130.
 - [7] Arcones, M., Kvam, P. & Samaniego, F.J. (2002). On nonparametric estimation of distributions subject to a stochastic precedence constraint, *Journal of the American Statistical Association* **97**, 170–182.
 - [8] Hollander, M. & Samaniego, F.J. (2007). A nonparametric, signature-based method for comparing the reliability of arbitrary systems via stochastic precedence, Technical Report # 472, Department of Statistics, University of California, Davis.
 - [9] Kvam, P. & Samaniego, F.J. (1993). On the Inadmissibility of Empirical Averages as Estimators in Ranked Set Sampling, *Journal of Statistical Planning and Inference* **36**, 39–55.
 - [10] Boland, P.J. (2001). Signatures of indirect majority systems, *Journal of Applied Probability* **38**, 597–603.
 - [11] Shaked, M. & Suarez-Llorens, A. (2003). On the comparison of reliability experiments based on the convolution order, *Journal of the American Statistical Association* **98**, 693–702.
 - [12] Belzunce, F. & Shaked, M. (2004). Failure profiles of coherent systems, *Naval Research Logistics* **51**, 477–490.
 - [13] Navarro, J. and Shaked M. (2006). Hazard rate ordering of order statistics and systems, *Journal of Applied Probability* **43**, 391–408.
 - [14] Dugas, M. & Samaniego, F.J. (2007). On optimal system design in reliability-economics frameworks, *Naval Research Logistics*, in press.
 - [15] Satyanarayana, A. & Prabhakar, A. (1978). A new topological formula and rapid algorithm for reliability analysis of complex networks, *IEEE Transactions on Reliability* **TR-30**, 82–100.
 - [16] Boland, P.J., Samaniego, F.J. & Vestrup, E. (2003). *Linking Dominations and Signatures in Network Reliability Theory*, in *Mathematical and Statistical Methods in Reliability*, B. Lindqvist & K. Doksum, eds, Singapore: World Scientific, 89–103.
 - [17] Samaniego, F.J. (2007). *System Signatures and Their Applications to Engineering Reliability*, New York: Springer, in press.

FRANCISCO J. SAMANIEGO

Bayesian Reliability Demonstration

Introduction

Many systems are required to be highly reliable, and prior to their practical use such reliability may need to be demonstrated through testing. This requires determination of the amount and type of testing that is needed in order to demonstrate a certain level of reliability. In many situations, in particular where high reliability is required, one would not accept faults occurring during such testing, as these are evidence against the system's ability to perform its task well and are likely to lead to rejection of the system. As such reliability demonstration testing clearly deals with uncertainty and information to reduce this uncertainty, statistical methods can provide valuable insights and can guide such testing. We focus on **Bayesian methods**, which allow previous information to be taken into account in a subjective manner *via* the prior distribution chosen, and, as such, can lead to reduced test effort. However, when emphasizing the *demonstration* of reliability, it might not be considered appropriate to rely on prior information, be it based on historical data from similar systems or expert judgement. We will discuss this feature in detail, as it strongly influences the choice of prior distribution. But first, we provide a brief overview of some key contributions to the literature on Bayesian reliability demonstration (BRD), and we focus on the criterion used for "reliability".

In line with the increasing development and popularity of Bayesian statistical inference in the late 1960s and 1970s, many researchers in reliability developed BRD test philosophies and procedures during this period. An excellent overview of contributions to BRD during this period is given in Chapter 10 of Martz and Waller's well-known book *Bayesian Reliability Analysis* [1]. A key topic is the choice of BRD criterion. Traditionally [1], criteria considered were closely linked to those in related statistical areas such as **quality control**, and mostly used producer's and consumer's risks expressed in terms of acceptance or rejection of a system after BRD testing, conditional on assumptions on characteristics of the system's failure distribution. Here, we use "failure distribution" as a generic term for the

probability distribution of the random variables of interest in BRD, which might, for example, include number of failures over defined periods of testing (as is typical in so-called "attribute testing") or times between failures. The characteristics typically used are a parameter representing the proportion of failures, as a population characteristic in the standard statistical sense for Bernoulli random quantities (attribute testing) or **Mean-Time-To-Failure** (MTTF) or **Mean-Time-Between-Failures** (MTBF). It is important to notice that, although such characteristics have intuitive appeal, they are not directly observable and only meaningful in combination with additional statistical modeling assumptions (mostly exchangeability of the random quantities of interest). The Bayesian approach to reliability demonstration distinguishes itself from other statistical approaches to such problems *via* the use of **prior distributions** on these parameters, which can reflect prior knowledge – a crucial question is, of course, whether or not such prior knowledge should be included in BRD and, if perhaps not, what role the Bayesian prior then might have and how it can be chosen (we return to this important issue in the section titled "The Prior Distribution"). Once one accepts the more or less traditional framework for **sample size** determination *via* producer's and consumer's risks, a variety of probabilities can be considered for the detailed specification of the BRD criterion, and all of these lead to many different solutions, which are however closely related from philosophical perspectives. Such probabilities include posterior risks, average risks and rejection probabilities, and study of corresponding BRD approaches which led to a significant number of papers in the literature until the late 1970s [1]. A typical contribution to this literature is the study of behavior of some quantities used in BRD tests, by Higgins and Tsokos [2]. They show that the use of two different priors for MTBF, which cannot be easily distinguished between at the prior stage, can lead to substantial differences in the posterior risks in BRD. From this, we can conclude that choice of the prior distribution was known to be crucial and problematic in BRD, yet it has received relatively little attention in the BRD literature.

Since the early 1980s, BRD seemed to disappear as a main topic from the reliability literature, with further development mostly in industry, so less frequently reported in the academic literature, and often embedded in wider programmes for achieving

quality and reliability. Examples of such progress on methods for reliability demonstration, including Bayesian approaches, in the automotive industry are presented by Lu and Rudy [3, 4]. Recently, Yadav *et al.* [5] raised interesting issues on reliability demonstration, emphasizing the need for a comprehensive approach which integrates reliability demonstration into the product development process. They emphasize that three dimensions of product design must be considered: physical structure, functional requirements, and time in service. They propose some statistical methodology for specific distributions and test plans, but although they mention the importance of including prior information, they do not formalize this nor do they use a Bayesian approach. This work, however, provides interesting arguments for renewed interest in BRD.

In recent years, the current authors have presented, in part jointly with Rahrhrouh, a different perspective on BRD, focussing solely on attribute testing [6–10]. This work has mainly been motivated by the observation that methods presented in the literature seemed not to take explicitly into account characteristics of the process in which the system is to be used after BRD testing, nor practical constraints on such testing. For example, they distinguished between the cases of catastrophic and noncatastrophic system failure in the process after testing, and they also studied optimal decisions for both BRD and choice of **redundancy** level in a system, from overall cost perspective. They also considered the role and choice of the prior distribution in BRD in detail, which is crucial yet far from trivial and had not yet received much attention. We briefly summarize main results in the sections “Zero-Failure Reliability Demonstration” and “The Prior Distribution”, illustrated by some examples in the section titled “Examples”, and discuss some further aspects in the section titled “Summary and Related Topics”.

Zero-Failure Reliability Demonstration

Martz and Waller [11] considered the before mentioned risk-based BRD from the perspective of only tests that reveal zero failures. This has two advantages. First, in practical demonstration of, particularly, high reliability, occurrence of one or more failures goes against the very nature of reliability demonstration, and is likely to lead to improvement

actions of the system, where possible, followed by renewed reliability demonstration. This cycle brings with it very important challenges, also called “test-fix-test”, which are beyond the scope of this paper. The second advantage of restricting attention to inference only valid when testing reveals zero failures, is the relative mathematical simplicity of the resulting Bayesian posterior and corresponding predictive distributions, enabling these to be embedded in wider decision-theoretic frameworks involving aspects of the process in which the system is required after testing, and inclusion of costs and time considerations and constraints. This and the following two sections provide some details of such a predictive approach to zero-failure BRD, with attention restricted to attribute testing [6–10].

We consider systems which have to deal satisfactorily with tasks of k different types, where we assume complete independence between tasks of different type (both in the success probability of the system to deal with such different tasks, and in the statistical sense). In a standard Bayesian setting [1, 12, 13], we use a Binomial model for the number of failures in n_i tests of tasks of type i , for $i = 1, \dots, k$, with parameter θ_i representing the unknown probability of a failure when the system has to deal with a task of type i . We assume a Beta (α_i, β_i) prior distribution for θ_i , and we further assume that zero failures will be observed in the n_i tests, denoted as data $(n_i, 0)$, so these inferences are only valid for zero-failure BRD. The corresponding posterior predictive probability of zero failures in m_i tasks of type i , in the process after such testing, is

$$P(0|m_i, (n_i, 0)) = \prod_{j=1}^{m_i} \frac{j + \beta_i + n_i - 1}{j + \alpha_i + \beta_i + n_i - 1} \quad (1)$$

For $n_i > 0$, (1) holds for $\alpha_i > 0$ and $\beta_i \geq 0$. Because no failures are observed we cannot use $\alpha_i = 0$, as the corresponding posterior distribution would be improper [12]. The hyperparameters α_i and β_i can be interpreted in terms of results of an imaginary earlier test of $\alpha_i + \beta_i$ tasks of type i , in which the system failed to deal with α_i such tasks, but performed β_i such tasks successfully. For BRD, this implies that inferences are quite insensitive to choice of β_i , where an increase in β_i means that the minimal required n_i is reduced by the same number. However, such inferences are highly sensitive to the choice of α_i , as effectively the reliability demonstration requires n_i

tests without failures to counter the prior information of α_i “imaginary test failures”. We discuss the choice of the prior distribution in more detail in the section titled “The Prior Distribution”, where we will advocate the use of Beta prior distributions with $\alpha_i = 1$ and $\beta_i = 0$. It should be noted that this chosen value for β_i leads to an improper prior distribution, so formally not really a probability distribution as it does not integrate to one over the parameter space. No statistical inferences can be based on improper distributions, but in this setting we assume explicitly that, for posterior inference for demonstrating reliability, the data contain a positive number of observed successes, which ensures that the corresponding posterior distribution is proper.

We briefly state some results from [8], on required test numbers n_i , such that zero failures in testing ensures a probability p of no failures in the process, either for deterministic number of tasks or for a stochastic process considered over a fixed period of time. This setting is, for example, suitable for catastrophic failures, that is, a failure leads to break-down of the whole system, in the sense that the clear target is to demonstrate reliability *via* an inferred (very) high probability of zero failures in the process after testing, when the system is actually in service. In such cases (e.g. alarm systems for chemical plants), it is often difficult to bring cost considerations into the argument, either due to inherent lack of knowledge of consequences of rare events or due to ethical issues, so requirements directly on the predictive probability of zero failures are attractive. Of course, such cost considerations are likely to influence choice of an appropriate value of p .

For the deterministic case, where the system must perform a known number m_i of tasks of type i in the process, the n_i must be such that

$$P(0|m_i, (n_i, 0), i = 1, \dots, k) = \prod_{i=1}^k \prod_{j=1}^{m_i} \frac{j + \beta_i + n_i - 1}{j + \alpha_i + \beta_i + n_i - 1} \geq p \quad (2)$$

If we aim at satisfying this requirement with a minimal total number of tests, there is no closed-form solution for n_i . For the special case with $\alpha_i = \alpha$, $s_i = \beta_i + n_i = s$ and $m_i = m$, for all i , an approximation for the minimum required value of s to meet this

reliability requirement, with p close to 1, is

$$\frac{p^{1/(k\alpha)}}{1 - p^{1/(k\alpha)}} \times m \quad (3)$$

For the case with $\alpha = 1$ and $\beta_i = 0$, for all $i = 1, \dots, k$, which we advocate for use in reliability demonstration (see the section titled “The Prior Distribution”), (equation 3) gives the exact value of the optimal (real-valued) n_i ’s [8] (the integer-valued solution is found by checking this probability value for the neighboring integer-valued k -vectors). The fact that, to minimize the total number of zero-failure tests required in this situation, one would indeed wish to use equal values of s_i , is proven in [8].

For p close to 1, the expression (3) is close to

$$\frac{p}{1 - p} \times k\alpha m \quad (4)$$

which can be proved *via* the same argument as used in [8] to show that equation (3) is approximately linear in k . The expression (4) makes clear that the required number of zero-failure tests is very sensitive to the choice of α . Clearly, the required value of n_i is less sensitive to small changes in β_i , as the sum of n_i and β_i must be a constant. The influence of the α_i and β_i on the required number of zero-failure tests is similar in more general situations with unequal α_i , s_i and m_i .

Of course, in many real-world processes the number of tasks m_i are not deterministic, which is for example typically the case in alarm systems. An important and interesting result is the following: let M_i be the random number of tasks of type i that the system has to perform in the process after testing. Then, for all possible probability distributions for M_i with the same expected value $E(M_i)$, the most BRD testing is actually required for the deterministic case with $m_i = E(M_i)$ (this is a consequence of Jensen’s inequality). Hence, as long as one can specify an expected value for the number of tasks the system will be required to perform (or of course an upper bound to stay on the safe side), then the amount of testing required by assuming this number to be deterministic leads to conservative test numbers. It also turned out from many examples studied, although formal proof was not achieved, that if M_i actually has a Poisson distribution, which is often suitable to model accidental events, then the optimal numbers of BRD tests of tasks of type i are

very close to the corresponding optimal test numbers in the deterministic case, which in practice means that over-testing would be prevented in such cases if you based your calculations upon the deterministic case instead of the Poisson case [8]. Example 1 in the section titled “Examples” illustrates this approach.

A different predictive BRD approach was presented in [9], which is suited for noncatastrophic system failure during the process after testing. This implies that the system will be able to continue operating in case such a failure occurs, and also that cost considerations become more important, so both costs of testing and of failure are taken into account. That approach aims at minimization of overall expected costs for testing and the process after testing, and takes constraints on testing (both budget and time) into account. Optimal BRD plans are derived by solving relatively straightforward constrained optimization problems, with several analytical results making clear the influence of constraints on testing. The reader is referred to [9] for more details, we wish to emphasize that it is again the predictive nature of this approach, with reliability criteria explicitly expressed in terms of the system’s performance in the process after testing, and using only observable random quantities (mostly the number of failures per type of task), which makes this an attractive solution to BRD.

In [10] a similar approach, yet again more suitable for catastrophic failures, is considered in which reliability of a system, to be demonstrated, is controlled both by tests and decisions on in-built system redundancy. This also considers minimization of total expected costs, taking into account the costs of additional redundancy which of course enables fewer zero-failure tests in order to demonstrate a required level of reliability, the latter typically included as a constraint for the optimization problem. We refer to [10] for further details, Example 2 in the section titled “Examples” shows the potential of this approach.

The important considerations of the role and choice of the prior distribution are similar for all the different predictive BRD approaches mentioned here, and are commented on in the next section.

The Prior Distribution

A main concern in Bayesian statistics is the choice of the prior distribution. In many applications, if sufficient data are observed then the inferences based

on the posterior distribution are pretty robust with regard to the choice of the prior distribution. In the zero-failure BRD setting, however, the posterior distribution remains highly sensitive to some changes in the prior distribution. Hence, it is not straightforward to use so-called “noninformative” priors [12] for such zero-failure testing [7, 8].

Design of experiments within the Bayesian framework is also strongly influenced by the prior distribution used [12]. It is usually advocated that the main role of the prior distribution is to take subjective information into account. From this perspective, the main aim of design of experiments, in Bayesian statistics, is for the person whose beliefs are modeled, to learn in some optimal manner from the data resulting from the experiment [14]. In case of testing for high reliability, one would often have quite strong prior beliefs about the quality of the system, in the sense that one would already expect the system to be highly reliable, such that one would be somehow surprised if there were still one or more failures observed during testing. If one takes such beliefs into account *via* a prior distribution, then this will lead to relatively few zero-failure tests being required in order to meet a reliability criterion. Using one’s prior beliefs as such, the intention of the tests would be to learn, in a way that is dependent on the subjective input. In the extreme case, if one is already so confident about the system’s reliability that one judges testing unnecessary, one can choose a corresponding prior distribution such that indeed no further testing is required. Of course, this may not convince others. A similar difference between personal learning and convincing others, affects general Bayesian design of experiments and the role of **randomization** [15].

We propose that the prior distribution used for Bayesian reliability *demonstration*, plays a different, nonsubjective, role, as reliability demonstration is not normally intended to convince only the person whose beliefs are modeled, of the reliability of a system, but to satisfy externally imposed rules or to convince others of the system’s reliability. From this perspective, the prior distribution can be regarded as reflecting neutral, or even pessimistic, prior information, which the test data must counter in order to *demonstrate* reliability.

From this perspective of reliability demonstration, we advocate the use of the Beta prior for θ_i , with hyperparameters $\alpha_i = 1$ and $\beta_i = 0$. The choice

of $\beta_i = 0$ is natural, given its interpretation and its effect on the required number of zero-failure tests [8]. An appropriate value of α_i is more open to discussion, and this choice has great impact on the number of zero-failure tests required. With the interpretation as “having seen α_i tests failing” before testing, our test criterion demands sufficiently many zero-failure tests, such that, on the basis of the combined information represented by this prior distribution and the test results, the reliability requirement is satisfied. The choice $\alpha_i = 1$ seems reasonable from this perspective, as it asks for sufficient zero-failure tests to outweigh one suggested earlier failure. We consider $\alpha_i = 1$ a fairly conservative choice, in the sense that many zero-failure tests are required to demonstrate high reliability. Hartigan [16] suggested the same prior distribution for similar conservative inferences. A further argument in favor of this choice of α_i and β_i results from analysis of this same problem from a different foundational perspective [6].

Examples

Example 1 This example illustrates the BRD approach from [8], with failures assumed to be catastrophic. It considers the optimal test numbers for a system that has to perform $k = 4$ types of tasks, using prior distributions with $\alpha_i = 1$ and $\beta_i = 0$ for all $i = 1, \dots, 4$. We aim at minimization of the total number of tests, assuming they reveal no failures, such that the resulting predictive probability (2) of zero failures in the process is at least p . For the deterministic case, let $\underline{m} = (1, 2, 4, 9)$. The corresponding optimal test numbers are given in Table 1.

This illustrates that very high reliability can only be *demonstrated* by very many zero-failure tests. For this setting, if tasks are assumed to arrive according to **Poisson processes**, with their expected numbers equal to the m_i used above, then the numbers of

zero-failure tests required are indeed nearly identical to those in Table 1. If we increase the m_i , then the required test numbers increase by about the same factor, for example optimal testing for $\underline{m} = (10, 20, 40, 90)$, and $p = 0.90$, is achieved by taking $\underline{n} = (699, 985, 1387, 2067)$.

Example 2 In this example we briefly illustrate the optimal zero-failure test numbers and numbers of components, for systems with redundancy [10]. We consider a **2-out-of-y system**, so for the system to function a minimum of 2 out of y components must function, and we assume components to function independently. We only consider the deterministic case, with the numbers of tasks of $k = 3$ independent types required in the process after testing equal to one, three and six, respectively. The optimal BRD zero-failure test numbers and optimal number of components y are presented in Table 2 for several cases. The cost per component installed is Q , C is the cost incurred if the system does not deal successfully with all 10 tasks in the process after testing, and $\underline{c}$ represents the costs of one test per type of task. Again we use Beta prior distributions with $\alpha_i = 1$ and $\beta_i = 0$ for all types of tasks, and the optimization criterion chosen is minimal expected overall costs, so for testing and the system’s functioning during the process considered after testing, with the constraint that the posterior predictive probability of zero failures in those 10 tasks after testing should be at least $p = 0.95$. The cases considered are:

- (1) $\underline{c} = (1, 1, 1)$, $C = 10\,000$
- (2) $\underline{c} = (20, 50, 50)$, $C = 10\,000$
- (3) $\underline{c} = (20, 50, 50)$, $C = 1\,000\,000$.

Obviously, higher built-in redundancy (larger y) requires fewer zero-failure tests. Cases (1) and (2) in Table 2 illustrate the fact that increasing testing cost c_i per test of type i , reduces the optimal number of zero-failure tests n_i , and may increase

Table 1 Optimal test numbers

p	n_1	n_2	n_3	n_4
0.90	70	98	139	207
0.95	144	204	287	429
0.99	737	1042	1474	2209
0.995	1479	2091	2956	4433
0.999	9457	10 553	13 943	21 505

Table 2 Optimal numbers of zero-failure tests ($\underline{n}$) and components (y)

Case(s)	$\underline{n}, y$		
	$Q = 300$	$Q = 1000$	$Q = 3000$
(1)	(47, 68, 86), 3	(47, 68, 86), 3	(201, 347, 489), 3
(2)	(8, 9, 10), 5	(12, 12, 15), 4	(26, 28, 35), 3
(3)	(14, 14, 16), 8	(20, 21, 23), 6	(41, 43, 51), 4

the optimal y . For case (2) with $Q = 1000$, Table 2 shows that the optimal solution is to install $y = 4$ components with $\underline{n} = (12, 12, 15)$ zero-failure tests. If we increase the cost per component to $Q = 3000$ in the same setting, the optimal solution for case (2) is $y = 3$ components with $\underline{n} = (26, 28, 35)$ tests. If we consider a 2-out-of-2 system instead (not presented in Table 2), the optimal zero-failure test numbers, in this case, are (296, 322, 454). This illustrates that the option of building in redundancy can greatly reduce the test requirements and the corresponding expected costs. Cases (2) and (3) illustrate that increasing process failure cost C requires an increased number of components and/or an increased number of zero-failure tests to minimize the total costs and to demonstrate the required reliability level.

Summary and Related Topics

BRD is a topic of great interest, both for applications and from the perspective of foundations of statistics, the latter due to wide variety of possible reliability criteria and the nonstandard role prior distributions play when test data are required to *demonstrate* reliability. The predictive approach to BRD, recently presented by the current authors and discussed above, is promising in its flexibility with regard to reliability criteria used and the possible inclusion of constraints in the optimization problems considered to take practical considerations into account [17]. As this approach has so far only been developed for attribute testing, there are important challenges for a similar way for BRD involving (continuous) time random quantities. This is particularly challenging in zero-failure situations, where observations would only consist of observed periods without failures, so all observations would be right-censored (see [18] for a nonparametric predictive approach to a related problem in probabilistic safety assessment). There are many topics in statistics, quality control and reliability, which consider problems that are closely related to BRD, for example, **software reliability** verification and testing [19–21] and **acceptance sampling**, as well as topics such as fault removal and re-testing, and accelerated life testing [22], and linking these to BRD for large-scale systems and developing engineering projects is a major challenge, so BRD promises to offer exciting opportunities for research and applications over the next decades.

References

- [1] Martz, H.F. & Waller, R.A. (1982). *Bayesian Reliability Analysis*, John Wiley & Sons, New York.
- [2] Higgins, J.J. & Tsokos, C.P. (1976). On the behavior of some quantities used in Bayesian reliability demonstration tests, *IEEE Transactions on Reliability* **25**, 261–264.
- [3] Lu, M.W. & Rudy, R.J. (2001). Reliability demonstration test for a finite population, *Quality and Reliability Engineering International* **17**, 33–38.
- [4] Lu, M.W. & Rudy, R.J. (2001). Laboratory reliability demonstration test considerations, *IEEE Transactions on Reliability* **50**, 12–16.
- [5] Yadav, O.P., Singh, N. & Goel, P.S. (2006). Reliability demonstration test planning: a three dimensional consideration, *Reliability Engineering and System Safety* **91**, 882–893.
- [6] Coolen, F.P.A. & Coolen-Schrijner, P. (2005). Nonparametric predictive reliability demonstration for failure-free periods, *IMA Journal of Management Mathematics* **16**, 1–11.
- [7] Coolen, F.P.A. & Coolen-Schrijner, P. (2006). On zero-failure testing for Bayesian high reliability demonstration, *Proceedings of the Institute of Mechanical Engineers, Part O: Journal of Risk and Reliability* **220**, 35–44.
- [8] Coolen, F.P.A., Coolen-Schrijner, P. & Rahrour, M. (2005). Bayesian reliability demonstration for failure-free periods, *Reliability Engineering and System Safety* **88**, 81–91.
- [9] Coolen, F.P.A., Coolen-Schrijner, P. & Rahrour, M. (2006). Bayesian reliability demonstration with multiple independent tasks, *IMA Journal of Management Mathematics* **17**, 131–142.
- [10] Rahrour, M., Coolen, F.P.A. & Coolen-Schrijner, P. (2006). Bayesian reliability demonstration for systems with redundancy, *Proceedings of the Institute of Mechanical Engineers, Part O: Journal of Risk and Reliability* **220**, 137–145.
- [11] Martz, H.F. & Waller, R.A. (1979). A Bayesian zero-failure (BAZE) reliability demonstration testing procedure, *Journal of Quality Technology* **11**, 128–138.
- [12] Bernardo, J.M. & Smith, A.F.M. (1994). *Bayesian Theory*, John Wiley & Sons, Chichester.
- [13] Gelman, A., Carlin, J.B., Stern, H.S. & Rubin, D.B. (2003). *Bayesian Data Analysis*, 2nd Edition, Chapman & Hall, London.
- [14] Chaloner, K. & Verdinelli, I. (1995). Bayesian experimental design: a review, *Statistical Science* **10**, 273–304.
- [15] Berry, S.M. & Kadane, J.B. (1997). Optimal Bayesian randomization, *Journal of the Royal Statistical Society. Series B* **59**, 813–819.
- [16] Hartigan, J.A. (1983). *Bayes Theory*, Springer, New York.

- [17] O'Connor, P.D.T. (2001). *Test Engineering: A Concise Guide to Cost-Effective Design, Development, and Manufacture*, John Wiley & Sons, Chichester.
- [18] Coolen, F.P.A. (2006). On probabilistic safety assessment in case of zero failures, *Proceedings of the Institute of Mechanical Engineers, Part O: Journal of Risk and Reliability* **220**, 105–114.
- [19] Kuo, L. (1999). Software reliability, in *Encyclopedia of Statistical Sciences*, S. Kotz, ed, John Wiley & Sons, Vol. 3, pp. 671–680.
- [20] Singpurwalla, N.D. & Wilson, S.P. (1999). *Statistical Methods in Software Engineering*, Springer, New York.
- [21] Wooff, D.A., Goldstein, M. & Coolen, F.P.A. (2002). Bayesian graphical models for software testing, *IEEE Transactions on Software Engineering* **28**, 510–525.
- [22] Meeker, W.Q. & Escobar, L.A. (1998). *Statistical Methods for Reliability Data*, John Wiley & Sons, New York.

Related Articles

Bayesian Design of Reliability Experiments; Bayesian Reliability Analysis.

FRANK P.A. COOLEN AND
PAULINE COOLEN-SCHRIJNER

Software Reliability: Bayesian Analysis

Introduction

Software reliability is usually defined as the probability that a piece of software runs without failing under certain operational conditions for a given time. In the software testing phase, it is normally assumed that after a failure, the software is debugged, so that on average, **reliability growth** is observed. The most common software reliability growth models, starting from the approach of Jelinski and Moranda [1] (hereafter JM), try to explain the behavior of the interfailure times $T_1, T_2, \dots$ (type I models) or the point processes that they generate, $N_1(t_1), N_2(t_2), \dots$ (type II models). For a full review, see **Software Reliability Modeling and Analysis**.

Bayesian Inference: General Aspects

An advantage of the Bayesian approach is that it directly incorporates prior information *via* the elicitation (see **Prior Distribution Elicitation**) of expert opinions (see **Expert Opinion in Reliability**) or the use of failure data from previous projects (which are often available in software reliability problems) in the form of a prior distribution. Thus, given a model $\mathcal{M}$ for interfailure times T_i (similarly for other data types) and a prior distribution $f(\theta|\mathcal{M})$ for the parameters θ , then when data $\mathbf{t}_n = (t_1, \dots, t_n)$ are observed, the posterior distribution for θ is calculated from Bayes theorem as

$$f(\theta|\mathcal{M}, \mathbf{t}) \propto f(\theta|\mathcal{M})l(\theta|\mathcal{M}, \mathbf{t}_n) \quad (1)$$

where $l(\theta|\mathcal{M}, \mathbf{t}_n)$ is the likelihood function. Prediction is straightforward, for example the survivor function of the time to next failure T_{m+1} is

$$P(T_{m+1} > t|\mathcal{M}, \mathbf{t}_n) = \int P(T_{m+1} > t|\mathcal{M}, \mathbf{t}_n, \theta) \times f(\theta|\mathcal{M}, \mathbf{t}_n) d\theta \quad (2)$$

An important aspect of software reliability modeling is the software testing strategy. Various testing plans have been considered, for example, one stage testing for a single fixed time period, say T .

The choice of T is a decision problem and the Bayesian approach is to define a utility function $U(T|\mathcal{M}, \theta)$ which evaluates the gains and losses of each possible choice. The optimal Bayes decision is then to choose the option which maximizes the expected utility function

$$E[U(T|\mathcal{M}, \theta)] = \int U(T|\mathcal{M}, \theta)f(\theta|\mathcal{M}) d\theta \quad (3)$$

An initial application in software testing is [2] and a good general review is [3], Chapter 6.

Although many software reliability models have been proposed, no single model has been demonstrated to be globally superior and thus, model selection is important. The Bayesian approach uses the prequential likelihood and Bayes factors. As successive data, $t_1, t_2, \dots, t_n$ are observed, then the prequential likelihood of $\mathbf{t}_i = (t_1, \dots, t_i)$ for $i = 1, \dots, n$, given model $\mathcal{M}$ is

$$f(\mathbf{t}_i|\mathcal{M}) = \int f(\mathbf{t}_i|\mathcal{M}, \theta)f(\theta|\mathcal{M}) d\theta \quad (4)$$

see, for example [4]. The performance of two different models $\mathcal{M}_1$ and $\mathcal{M}_2$ can be compared by graphing the prequential likelihood ratio

$$B_2^1(\mathbf{t}_i) = \frac{f(\mathbf{t}_i|\mathcal{M}_1)}{f(\mathbf{t}_i|\mathcal{M}_2)} \quad (5)$$

Such graphics are used to illustrate the ranges of data for which one model outperforms another. The prequential likelihood ratio given the whole data set is called the *Bayes factor*, see [5], and is used to compare the relative credibility of the two models in a similar way as the classical likelihood ratio.

Bayesian Inference: Specific Models

Various type I “bug counting” models have been developed under the assumption that the program initially contains an unknown fixed number of faults, say N , and that as each failure occurs, then these faults are successively corrected. See, for example, [1, 6–8]. The first and simplest is the JM model, [1], which further assumes that the failure rate of T_i is directly proportional to $N - i + 1$ implying that given an additional parameter ϕ ,

$$T_i|N, \phi, \mathbf{t}_{i-1} \sim \exp((N - i + 1)\phi) \quad (6)$$

that is, an exponential distribution with rate $(N - i + 1)\phi$ where $\mathbf{t}_{i-1} = (t_1, \dots, t_{i-1})$. Given the observed data $\mathbf{t}_n$, the likelihood is

$$l(N, \phi | \mathbf{t}_n) = \frac{N!}{(N - n)!} \phi^n \times \exp \left\{ -\phi \left((N + 1) \sum_{i=1}^n t_i - \sum_{i=1}^n i t_i \right) \right\} \quad (7)$$

and in principle, inference could be carried out using maximum likelihood. However it has been demonstrated that **maximum likelihood estimation** often fails to provide sensible answers, see, for example, [9] and Bayesian methods should be preferred. In [9], the following **prior distributions** are suggested

$$N \sim Po(\lambda) \quad \phi \sim G(a, b) \quad (8)$$

that is, a Poisson distribution with mean λ for N and a gamma density for ϕ . Combining likelihood and priors in equations (7) and (8) it is easy to see that, *a posteriori*, conditional on ϕ , then $N - n$ has a Poisson distribution and given N , ϕ has a gamma distribution and therefore a Gibbs sampling scheme (see **Hierarchical Markov Chain Monte Carlo (MCMC) for Bayesian System Reliability**) as in [10] can be set up to sample the posterior parameter distributions.

A number of extensions to the basic approach have been developed. Hierarchical priors (see **Hierarchical Markov Chain Monte Carlo (MCMC) for Bayesian System Reliability**), for the parameters λ , a , b in equation (8) are introduced in [10, 11]. **Covariate** information in the form of software metrics is used in [12] and formal sensitivity analysis is undertaken in [13]. Finally, different software testing strategies were explored in [14, 15].

The alternative bug counting models cited all lead to similar likelihood functions and given a Poisson prior for N , can all be analyzed using Gibbs sampling techniques. In particular [16] provides a general Bayesian method for parametric bug counting models and a nonparametric approach is introduced in [17].

Other software reliability models are based on the idea of reliability growth, without taking fault numbers into account. Moranda [18] assumes that the failure rate of T_i is Dk^{i-1} so that

$$T_i | D, k, \mathbf{t}_{i-1} \sim \exp(Dk^{i-1}) \quad (9)$$

where $k < 1$, which implies that the failure rate always decreases (see [3, p. 166] for a Bayesian approach to this model). However, the assumption of a deterministically decreasing rate is very strong and a number of hierarchical models have been constructed by assuming exponential interfailure times $T_i | \lambda_i, \mathbf{t}_{i-1} \sim \exp(\lambda_i)$ and a further hierarchical prior for the failure rate λ_i to model the reliability growth. Littlewood and Verrall (hereafter LV), [19], proposed assuming

$$\lambda_i | \alpha, \psi(i) \sim \gamma(\alpha, \psi(i)) \quad (10)$$

where $\psi(i) = \beta_0 + \beta_1 i$ is an increasing function of i . Empirical Bayes parameter estimation (see **Empirical Bayes Estimation of Reliability**) for this model is introduced in [20] and a full Bayesian inference is studied for an extended model based on a polynomial $\psi(i)$ in [10]. In [21], a Markovian structure is used and the sequence of failure rates is assumed to be a martingale by Basu and Ebrahimi [22], where

$$\lambda_1 | \alpha, \beta \sim \gamma(\alpha, \beta),$$

$$\lambda_i | \lambda_{i-1}, \alpha \sim \gamma\left(\alpha, \frac{\alpha}{\lambda_{i-1}}\right) \text{ for } i > 1 \quad (11)$$

The majority of type II models, starting with Goel and Okumoto [23] and Musa and Okumoto [24], are based on nonhomogeneous Poisson processes so that the number of failures $N(t)$ in a time period of length t is assumed to have a Poisson distribution with some **intensity function** $\lambda(t)$. Bayesian inference for the Musa Okumoto model with the use of subjectively elicited expert priors is introduced in [25] and a general Bayesian approach for analyzing nonhomogeneous **Poisson process** (type I and type II) models is given in [16] and applied in [26]. Nonparametric estimation of the intensity function has been considered by [27, 28]. Finally, the use of covariate information such as software metrics in order to estimate the intensity function is examined in [29, 30].

An alternative approach to modeling interfailure times is to consider these as a **time series**. The following Bayesian autoregressive model (hereafter KF) is proposed in [31]

$$\log T_i = \alpha \theta_i + \epsilon_i \quad \theta_i = \alpha \theta_{i-1} + v_i \quad (12)$$

where ϵ_i and v_i are independent normal errors. Given a normal prior distribution for θ_0 , it is straightforward to calculate the predictive distributions of

future failures by Bayesian Kalman filtering, see, for example [32]. An extension to the case where α is unknown is given in [33]. An alternative, non-linear Kalman filter (NGKF) is introduced in [34] where

$$T_i | \theta_i \sim \gamma(v, \theta_i) \quad \phi_i = \frac{C\theta_i}{\theta_{i-1}} \sim \beta(\sigma, v) \quad (13)$$

that is a beta distribution, with prior distribution $\theta_0 \sim G(\sigma + v, u_0)$. When C is known, inference is straightforward and when C is unknown, various schemes for calculating the posterior distribution are available. A related idea of recent interest is the use of neural networks to model the time series. Bayesian regulation neural nets are examined in [35].

Example

Table 1, taken from [36] gives usage times between 34 successive failures, due to incorrect user documentation on a workstation.

The following models were fitted to this data:

- JM*: equation (6), with priors $N|\lambda \sim Po(\lambda)$, $\lambda \sim \gamma(0.1, 0.001)$, $\phi \sim \gamma(0.1, 0.1)$.
- LV*: equation (10) with $\alpha \sim \exp(1)$, $\beta_i \sim \exp(0.001)$ for $i = 1, 2$.
- Moranda*: equation (9) with $D \sim \gamma(0.1, 0.1)$ and $k \sim U(0, 1)$, a uniform distribution.
- KF*: equation (12) with $\epsilon_i \sim N(0, 4)$, $v_i \sim N(0, 0.004)$, and $\theta_0 \sim N(1, 1)$.
- NGKF*: equation (13) with $\theta_0 \sim \gamma(2, 1)$ and $C \sim U(0, 1)$.
- Basu*: equation (11) with $\beta \sim \gamma(0.1, 0.1)$ and $\alpha \sim U(0.25, 4)$.

Using the JM model as the base $\mathcal{M} = 1$, Figure 1 graphs the logs of the prequential likelihood ratios, (see equation (5)) in favor of the various alternatives.

The LV model performs best for most of the data but is finally overtaken by KF. As suggested in [36],

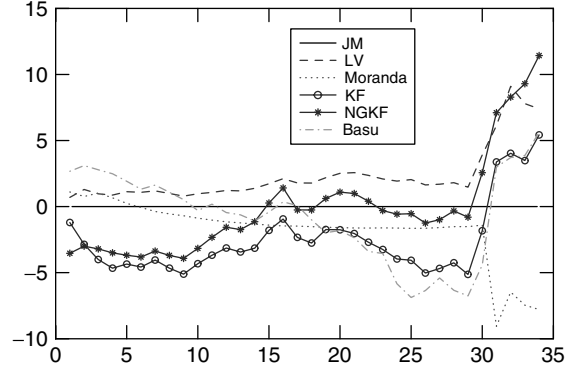


Figure 1 Log prequential likelihood ratios of alternative models versus JM

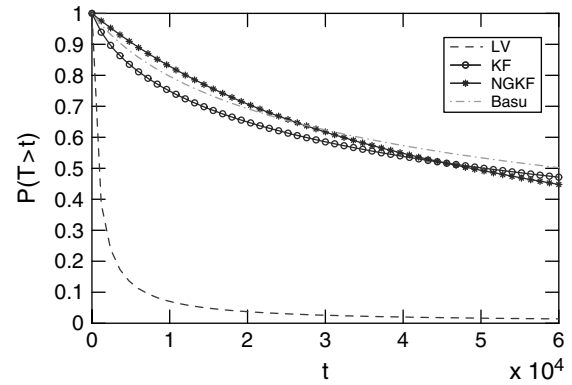


Figure 2 Predictive survivor functions for the LV, KF, NGKF, and Basu models

there appears to be a change point in the data at around the 30th failure.

Figure 2 illustrates the predicted survivor functions, see equation (2) for the time to next (35th) failure for the four best performing models.

The time series and Basu Ebrahimi models, which fit the last few interfailure times well are similar whereas the Littlewood-Verrall model provides much less confident predictions.

Table 1 Interfailure times data

1339	163	2582	583	202	337	112	373	383	44	72	77
179	87	23	40	427	165	21	75	131	251	775	175
105	561	1397	1601	552	4876	30767	710	20867	34789	—	—

Other Approaches to Software Reliability

Although this article has studied software reliability models based on interfailure times, there are other approaches. One method tries to relate software metrics, for example, lines of code to software quality. Recent work in this area has used Bayesian belief nets or graphical models, see [37]. Graphical models have also been applied to partition testing, that is, the choice of test cases in order to maximize the probability of fault detection (see e.g. [38]).

References

- [1] Jelinski, Z. & Moranda, P. (1972). Software reliability research, in *Statistical Computer Performance Evaluation*, W. Frieberger, ed, Academic Press, New York, pp. 465–484.
- [2] Dalal, S.R. & Mallows, C.L. (1986). When should one stop testing software? *Journal of the American Statistical Association* **83**, 872–879.
- [3] Singpurwalla, N.D. & Wilson, S.P. (1999). *Statistical Methods in Software Engineering*, Springer, New York.
- [4] Littlewood B. (1989). *Forecasting Software Reliability*, *Lecture Notes in Computer Science*, Vol. 341, Springer, Berlin.
- [5] Kass, R.E. & Raftery, A.E. (1995). Bayes factors, *Journal of the American Statistical Association* **90**, 773–795.
- [6] Schick, G.J. & Wolverton, R.W. (1973). Assessment of software reliability, *Proceedings in Operations Research*, Physica Verlag, Vienna, pp. 395–422.
- [7] Littlewood, B. (1981). Stochastic reliability growth: a model for fault removal in computer programs and hardware designs, *IEEE Transactions on Reliability* **30**, 313–320.
- [8] Raftery, A.E. (1987). Inference and prediction for a general order statistic model with unknown population size, *Journal of the American Statistical Association* **82**, 1163–1168.
- [9] Meinhold, R.J. & Singpurwalla, N.D. (1983). Bayesian analysis of a commonly used model for describing software failures, *The Statistician* **32**, 168–173.
- [10] Kuo, L. & Yang, T.Y. (1995). Bayesian computation of software reliability. *Journal of Computational and Graphical Statistics* **4**, 65–82.
- [11] Jewell, W.S. (1985). A Bayesian extension to a basic model of software reliability. *IEEE Transactions on Software Engineering* **11**, 1465–1471.
- [12] Wiper, M.P. & Rodriguez Bernal, M.T. (2001). Bayesian inference for a software reliability model using metrics information, *Proceedings of the European Conference on Safety and Reliability (ESREL)*, Politecnico de Torino, Torino, pp. 1999–2006, September 2001.
- [13] Wilson, S.P. & Wiper, M.P. (2000). Prior robustness in some common types of software reliability model, in *Robust Bayesian Analysis*, D. Rios Insua & F. Ruggeri, eds, Springer, New York, pp. 385–400.
- [14] Singpurwalla, N.D. (1989). Preposterior analysis in software testing, in *Statistical Data Analysis and Inference*, Y. Dodge, ed, Elsevier, Amsterdam, pp. 581–595.
- [15] Singpurwalla, N.D. (1991). Determining an optimal time for testing and debugging software, *IEEE Transactions on Software Engineering* **17**, 313–319.
- [16] Kuo, L. & Yang, T.Y. (1996). Bayesian computation for nonhomogeneous Poisson processes in software reliability, *Journal of the American Statistical Association* **91**, 763–773.
- [17] Samaniego, F.J. & Wilson, S.P. (2006). Non parametric analysis of bug counting models in software reliability, Research Report 06/03, Department of Statistics, Trinity College Dublin.
- [18] Moranda, P.B. (1975). Prediction of software reliability and its applications, in *Proceedings of the Annual Reliability and Maintainability Symposium*, Washington, DC, pp. 327–332.
- [19] Littlewood, B. & Verrall, J.L. (1973). A Bayesian reliability growth model for computer software, *Applied Statistics* **22**, 332–346.
- [20] Mazzuchi, T.A. & Soyer, R. (1988). A Bayes-empirical Bayes model for software reliability, *IEEE Transactions on Reliability* **37**, 248–254.
- [21] El Aroui, M.A. & Soler, J.L. (1996). A Bayes nonparametric framework for software-reliability analysis, *IEEE Transactions on Reliability* **45**, 652–660.
- [22] Basu, S. & Ebrahimi, N. (2003). Bayesian software reliability models based on martingale processes, *Technometrics* **45**, 150–158.
- [23] Goel, A.L. & Okumoto, K. (1979). Time dependent error detection rate for software reliability and other performance measures, *IEEE Transactions on Reliability* **28**, 206–211.
- [24] Musa, J.D. & Okumoto, K. (1984). A logarithmic Poisson execution time model for software reliability measurement, in *Proceedings of the Seventh International Conference on Software Engineering*, Orlando, pp. 230–237.
- [25] Campodónico, S. & Singpurwalla, N.D. (1994). A Bayesian analysis of the logarithmic Poisson execution time model based on expert opinion and failure data, *IEEE Transactions on Software Engineering* **20**, 677–683.
- [26] Kuo, L., Lee, J., Choi, K. & Yang, T. (1997). Bayes inference for S-shaped software reliability models, *IEEE Transactions on Reliability* **46**, 76–80.
- [27] Kuo, L. & Ghosh, S.K. (1997). Bayesian nonparametric inference for nonhomogeneous Poisson processes, Technical Report 97-18, Statistics Department, University of Connecticut.
- [28] Gutierrez Peña, E. & Nieto Barajas, L.E. (2003). Bayesian nonparametric inference for mixed Poisson

- processes, in *Bayesian Statistics*, J.M. Bernardo, M.J. Bayarri, J.O. Berger, A.P. Dawid, D. Heckerman, A.F.M. Smith & M. West, eds, Oxford University Press, Vol. 7, pp. 163–179.
- [29] Jeske, D., Qureshi, M. & Muldoon, E. (2000). A Bayesian methodology for estimating the failure rate of software, *International Journal of Reliability Quality and Safety Engineering* **7**, 153–168.
- [30] Ray, B., Liu, Z. & Ravishankar, N. (2006). Dynamic reliability models for software using time dependent covariates, *Technometrics* **48**, 1–10.
- [31] Singpurwalla, N.D. & Soyer, R. (1985). Assessing software reliability growth using a random coefficient autoregressive process and its ramifications, *IEEE Transactions on Software Engineering* **11**, 1456–1464.
- [32] Meinhold, R.J. & Singpurwalla, N.D. (1983). Understanding the Kalman filter, *The American Statistician* **37**, 123–127.
- [33] Singpurwalla, N.D. & Soyer, R. (1992). Non-homogeneous autoregressive processes for tracking (software) reliability growth, and their Bayesian analysis, *Journal of the Royal Statistical Society. Series B* **54**, 145–156.
- [34] Chen, Y. & Singpurwalla, N.D. (1994). A non-Gaussian Kalman filter model for tracking software reliability, *Statistica Sinica* **4**, 535–548.
- [35] Tian, L. & Noore, A. (2005). Evolutionary neural network modeling for software cumulative failure prediction, *Reliability Engineering and System Safety* **87**, 45–51.
- [36] Brocklehurst, S., Littlewood, B., Mellor, P., Tanner, A., Kanoun, K., Laprie, J.-C. & Metge, S. (1991). Analyses of software failure data, PDCS Technical Report 77, City University.
- [37] Neil, M., Krause, P. & Fenton, N. (2003). Software quality prediction using Bayesian networks, in *Software Engineering with Computational Intelligence, The Kluwer International Series in Engineering and Computer Science, Vol. 731*, T.M. Khoshgoftaar, ed, Kluwer Academic Publishers, Boston, pp. 136–172.
- [38] Coolen, F., Goldstein, M. & Munro, M. (2001). Generalized partition testing via Bayes linear methods, *Information and Software Technology* **43**, 783–793.

Related Articles

Bayesian Reliability Analysis; Bayesian Reliability Demonstration; Hierarchical Markov Chain Monte Carlo (MCMC) for Bayesian System Reliability; Prior Distribution Elicitation; Software Failure Data Analysis; Software Reliability Modeling and Analysis.

MICHAEL P. WIPER

Statistical Process Control, Bayesian

Introduction

Historically, **quality management** was driven largely by specification limits. In **acceptance sampling** for quality assurance, for example, items would be classified as defective or nondefective, based either on a qualitative assessment or their conformity to the specification limits. The seminal work of Shewhart [1] and more recent thinking by Taguchi [2] changed this focus, relegating specification limits to the background and concentrating on the “common cause” variability of the process. As long as the process readings lay within the range of common cause variability, the process would be left untouched, and when it appeared that the process was outside the range of common cause variability, it would be inferred that a “special cause” had intervened, and a search for and correction of the special cause would follow. From a statistical viewpoint, Shewhart’s charts can be thought of as repeated hypothesis tests, implemented graphically. For example, if the common-cause variability in a process measurement is that the measurement stream behaves like independent $N(\mu_0, \sigma_0^2)$ variables, then the Shewhart Xbar chart can be thought of as a repeated test of the null hypothesis that the current mean μ equals μ_0 , and the R or S chart as a repeated test of the null hypothesis that the current standard deviation σ equals σ_0 . Good book-length discussions of the area can be found in, for example, [3] and [4].

Bayesian Formulations of Charts

Quantifying this common cause variability in a frequentist sense, however, requires that large phase I samples be gathered. This substantial overhead, always undesirable, is impossible in short-run and start-up settings. It is also a failing of the frequentist approach typified by Shewhart charting that it is not able to make use of the prior information that is frequently available even in start-up and short-run problems. For example, when we start up a lathe to turn 6 mm bolts, previous experience using that lathe to turn 5 mm bolts is clearly relevant to this new

but similar problem, but hard to accommodate in a frequentist framework.

An early development of a Bayesian approach was that of Shirayayev [5] and subsequently Roberts [6]. These approaches may be thought of as Bayesian change point formulations [7, 8]. While the process is “in control” and only common causes operate, the process measurements X_n are modeled as following a common independent distribution $f_0(\cdot)$. At some unknown time point τ , a special cause may intervene, changing the distribution to $f_1(\cdot)$. This changed distribution may be of the same form as f_0 : for example both might be normals with the same variance but different mean, corresponding to a step change in mean, but the theory allows for general shifts. The shift in distribution remains until detected and diagnosed.

The likelihood of a sequence of data is

$$f(x_1, \dots, x_n | \tau) = \prod_{i=1}^{\min(n, \tau-1)} f_0(x_i) \prod_{j=\tau}^n f_1(x_j) \quad (1)$$

(the second product being absent if $n < \tau$).

In Shirayayev’s work, f_0 and f_1 were considered fully known, and τ modeled (*a priori*) by a geometric distribution with probability p of a shift at any observation. Shirayayev showed that $P_n = P(\tau \leq n | x_1, x_2, \dots, x_n)$, the posterior probability that the shift occurred before the n th observation is:

$$\begin{aligned} P_n &= \frac{\frac{p}{(1-p)^{n+1}} \sum_{k=1}^n (1-p)^k \prod_{j=k}^n \frac{f_1(x_j)}{f_0(x_j)}}{\frac{p}{(1-p)^{n+1}} \sum_{k=1}^n (1-p)^k \prod_{j=k}^n \frac{f_1(x_j)}{f_0(x_j)} + 1} \\ &= \frac{R_n^S}{R_n^S + 1} \end{aligned} \quad (2)$$

This posterior probability can be used to decide whether to stop the process on the grounds that the process has undergone a step change. Shirayayev in [5] provided the optimal stopping rule in the case where the distribution laws $f_0(\cdot)$ and $f_1(\cdot)$ are known: if a false alarm (stopping before there has been a change) costs 1 unit, and signaling a change at time $N > \tau$ costs $c(N - \tau)$ units, then the optimal (Bayes) solution is to stop at the smallest N for which the posterior probability $P_n \geq p^*$ for some predetermined value $0 < p^* < 1$ (close to

2 Statistical Process Control, Bayesian

1). According to Bather [9] if we wish to bound the expected false alarms by K , then the optimal $p^* = (K + 1)^{-1}$.

Letting p approach zero (following the realistic assumption that these step changes are infrequent) leads to the Shiriyayev–Roberts [6] procedure:

$$R_n^{SR} = \sum_{k=1}^n \prod_{i=k}^n \frac{f_1(x_i)}{f_0(x_i)} \quad (3)$$

and using the same loss function as above we obtain the following stopping rule:

$$N^{SR}(\gamma) = \inf \{n \geq 1 : R_n^{SR} \geq \gamma\} \quad (4)$$

The SR statistics (3) can also be given recursively: $R_0^{SR} = 0$ and for $n = 1, 2, \dots$

$$R_n^{SR} = (1 + R_{n-1}^{SR}) \frac{f_1(x_n)}{f_0(x_n)} \quad (5)$$

This equation is reminiscent of Page's [10] cumulative sum (**cusum**) statistic recursive formula:

$$W_n = \max_{1 \leq k \leq n} \prod_{i=k}^n \frac{f_1(x_i)}{f_0(x_i)} = \max\{W_{n-1}, 1\} \frac{f_1(x_n)}{f_0(x_n)} \quad (6)$$

a powerful frequentist tool in detecting (even small) persistent shifts in the distribution law, and whose performance is quite like that of the SR method under minimax criteria.

The Chernoff and Zacks [11] model is another change point with several notable differences from the Shiriyayev–Roberts approach. It avoids the fully specified SR setup, allowing unknown distributional parameters. So it is able to provide, not only inference on when the change occurred, but also the posterior distribution of the parameter of interest. More specifically the distributional model is a **normal distribution** with constant variance assumed without loss of generality (since it is assumed known) to equal 1. The mean at observation n is θ_n . If there has been no change, then $\theta_n = \theta_{n-1}$, but following a change $\theta_n = \theta_{n-1} + Z_n$, where Z_n is a normally distributed increment in the mean.

This model allows for multiple changes in the mean, and leads to computationally challenging expressions for its posterior distributions. A substantial simplification is Zacks and Kenett [12] AMOC

(at most one change) model. Reminiscent of the Shiriyayev formulation, this holds that

$$\theta_i = \begin{cases} \theta_0 & i = 1, \dots, \tau - 1 \\ \theta_1 = \theta_0 + Z, & i = \tau, \dots, n \end{cases} \quad (7)$$

but differs in that, rather than being constants known *a priori*, $\theta_0 \sim N(\mu_0, \sigma^2)$, $Z \sim N(\delta, \rho^2)$ are random variables.

This setting shows some philosophical differences from the Shewhart thinking, where the in-control mean μ_0 represents the mean when things are going right in the system. In all the change point approaches, μ_0 is simply the mean when the system is first observed. That the system is going right at that time is not an automatic given, but the result of care and attention on the part of the process owner.

A final generalization to be mentioned in this thread is that of Tsiamyrtzis and Hawkins [13], who considered a random walk model with occasional shocks for the mean of normally distributed data, that is, we observe:

$$X_n | \theta_n \sim N(\theta_n, \rho^2) \quad (8)$$

and the means obey the following:

$$\begin{aligned} &\theta_n | \theta_{n-1} \\ &\sim \begin{cases} N(\theta_{n-1}, \sigma^2) & \text{with probability } p \\ N(\theta_{n-1} + \delta, \sigma^2) & \text{with probability } 1 - p \end{cases} \end{aligned} \quad (9)$$

with $\theta_0 \sim N(\zeta, \sigma_0^2)$. Then the posterior distribution of θ_n given the data: $\mathbf{X}_n = (x_1, x_2, \dots, x_n)$ will be given as a mixture of 2^n normal distributions of the form:

$$p(\theta_n | \mathbf{X}_n) = \sum_{i=0}^{2^n-1} \alpha_i^{(n)} N(\theta_i^{(n)}, \hat{\sigma}_n^2) \quad (10)$$

where the weights, means and variance are provided recursively.

Here the equivalence between being out of control to the parameter having changed falls away; rather the process is out of control when the current mean has drifted too far from its initial acceptable level. This feature of the approach meshes with six-sigma thinking that the process can be left alone as long as the mean is no more than 1.5 standard deviations from its nominal setting.

The Kalman Filter

This last development brings us to the Kalman filter (KF) [14]. **SPC** ideas tend to be motivated by the idea of sudden discrete changes from some fixed in-control model to another fixed out-of-control model. The KF has a steadily changing parameter with no sudden jumps:

$$\theta_n = G_n \theta_{n-1} + w_n \quad (11)$$

where G_n is assumed known, and w_n induces the random drift in θ . Generalizing the standard SPC model

$$X_n \sim N(\theta_0, \sigma_0^2) \quad (12)$$

in the KF we have

$$X_n \sim N(F_n \theta_n, \sigma_0^2) \quad (13)$$

The parameter θ_n may be vector-valued. The success and wide adoption of the KF are helped by its simple recursive updating formulas.

Count Data

Shewhart methods for nonconformities are embodied in the C and U charts. The milestone in development of Bayesian alternatives is perhaps Hoadley's [15] quality measurement plan (QMP).

At time period n write X_n for the number of defects (nonconformities) in the sample, e_n for the expectation function discussed below, and θ_n for the true population index. As usual with count data, we use a Poisson model

$$X_n | \theta_n \sim \text{Poisson}(e_n \theta_n) \quad (14)$$

so the expected value of X_n involves an unknown parameter θ_n multiplying a known expectation coefficient e_n . These are calibrated so that while the process is operating correctly, θ_n should equal 1, that is, e_n should be the expected value of X_n .

The process parameter θ_n is assumed to follow a gamma prior distribution:

$$\begin{aligned} \theta_n | (\alpha, \beta) &\stackrel{iid}{\sim} \gamma(\alpha, \beta) \quad \text{with} \\ \pi(\theta_n | \alpha, \beta) &= \frac{\beta^\alpha}{\Gamma(\alpha)} \theta_n^{\alpha-1} e^{-\beta \theta_n} \end{aligned} \quad (15)$$

leading to the posterior distribution

$$\theta_n | (x_n, \alpha, \beta) \sim \gamma(\alpha + x_n, \beta + e_n) \quad (16)$$

The QMP graph is a time-ordered sequence of box and whisker plots with the box marking the 5 and 95th percentile and the whiskers the 1st and 99th percentile of the posterior distribution. Under squared error loss the Bayes (point) estimate of the parameter of interest (θ_n) is the mean of the posterior distribution, that is,

$$\hat{\theta}_n = \frac{\alpha + x_n}{\beta + e_n} \quad (17)$$

which is used as the center line in the box.

The process is deemed to be out of control if the lower whisker ends above the value 1, while an alert state, analogous to Shewhart warning limits, occurs if the box lies above the value 1.

Prior Parameters

In the parametric Bayesian approach, once the parametric form of the prior distribution has been decided (by either similar historic data or expert's opinion) we are left with the question of what the values of the prior distribution parameters (hyperparameters) will be. Two standard approaches for getting this needed information are **empirical Bayes** (EB) and purely Bayes (PB) approaches.

The EB approach is based on the fact that the marginal distribution of the data will depend on these hyperparameters and thus estimates can be obtained using the data. Therefore the EB approach requires the data in advance and so it can not be employed in the very beginning stage of the process, where the data will become available sequentially. It naturally follows that the more the data the better the estimates will be.

On the other hand in the PB approach the data are not allowed to be used twice in the modeling (first in estimating the hyperparameters and secondly in the Bayes theorem to provide the posterior distribution). Thus for the hyperparameters in cases where we are not able to provide some point estimates (coming from expert's opinion and/or other historic data—not the one we observe) a (second level) prior distribution can be used (hyper-prior). Under this scenario we can also incorporate hierarchical settings, where the prior distribution itself might depend on other

random quantities which we can simply model with a hyper-prior. So hierarchical schemes fit naturally in this context, something that is largely beyond the capabilities of straightforward frequentist models.

Finally, these approaches are not mutually exclusive; a PB approach can be used for an initial prior distribution, and then modified using EB prior to use in the actual charting.

An Evaluation of Bayesian SPC

The tension throughout statistical practice between frequentist and Bayesian approaches is present in SPC as well. It is fair to say though that Bayesian methods have gained less traction in SPC than in other areas of statistics.

Some of the reasons for this are the same inertia and resistance to change that cause some users who would benefit from a switch to cusums or exponentially weighted moving average (EWMAs) to continue their reliance on $\bar{X}$ and R charts. Some of the objections are resistance to Bayesian ideas as a matter of principle, and sometimes to the misconception that all informative priors come from some personalized and irreproducible introspection.

Along with these bad reasons is one with an element of validity: the much greater computational load involved in many of the Bayesian methods. For example, the Chernoff–Zacks Bayesian approach to the Shewhart question of a shift out of control leads, at the n th observation, to a mixture distribution with 2^n components that, unless tamed with some approximations, becomes unmanageable once the process is much beyond a dozen readings.

Beyond the straightforward natural conjugate priors, many Bayesian posteriors do not have closed forms, but must be addressed by estimation or approximation. Recent work in Markov Chain Monte Carlo (MCMC) has vastly simplified this process, but it remains beyond the scope of a nontechnical user who may be well able to use frequentist Shewhart, cusum or EWMA charts.

Turning to the other side of the comparison, what benefits accrue from using Bayesian SPC? The primary one is the ability to incorporate prior information in a natural seamless way, and to get the monitoring of the process up and running immediately, without the cumbersome preliminary phase I data-gathering exercise required by traditional frequentist

approaches. Always valuable, an early start is indispensable in short-run processes which do not run long enough to ever provide the several hundred phase I observations that are now recognized as necessary to calibrate known-parameter frequentist charts.

A further less tangible benefit is that the approach provides a full posterior distribution for any parameter of interest. This leads to the possibility of using a much richer array of inferential tools [16] than such binary observations as whether the current $\bar{X}$ lies inside or outside a Shewhart control limit. This posterior distribution allows us to check at each stage of the process, not only for quality **degradation** (as in frequentist's SPC), but for potential quality improvement as well. Furthermore within the Bayesian framework it is straightforward to obtain the predictive distribution [17] which can be used in forecasting and/or model assessment.

Last but not least is the fact that within the Bayesian framework, problems with multiple parameters are handled no differently than those with a single parameter, and there is no essential distinction between parameters of interest and nuisance parameters, except that in the final stages the latter are integrated out.

We argue that Bayesian approach should be used at all times where there is some prior information on process parameters. The benefit of the Bayesian formulation is that gives a unified framework which can take into account the prior information and seamlessly incorporates additional process observations.

As the prior information is most valuable when there is not much direct information, it is reasonable therefore to expect that the Bayesian approach will make substantial inroads in settings where the direct process measurements are less available: in short-run job-shop settings, where relevant informative prior information is common and large frequentist phase I exercises are impossible. In settings where large phase I data gathering is possible, the Bayesian approach may be attractive even to frequentists in providing a bridge allowing process monitoring to begin much sooner than is possible with the classical methods.

References

- [1] Shewhart, W.A. (1931). *Economic Control of Quality Manufactured Product*, D Van Nostrand, New York.

-
- [2] Taguchi, G. (1987). *System of Experimental Design: Engineering Methods to Optimize Quality and Minimize Costs*, Unipub/Kraus International Publications, White Plains.
 - [3] Kenett, R. & Zacks, S. (1998). *Modern Industrial Statistics - Design and Control of Quality and Reliability*, Duxbury Press, San Francisco.
 - [4] Montgomery, C.D. (2004). *Introduction to Statistical Quality Control*, John Wiley & Sons, New York.
 - [5] Shiriyayev, A.N. (1963). On optimum methods in quickest detection problems, *Theory of Probability and its Applications* **8**, 22–46.
 - [6] Roberts, S.W. (1966). A comparison of some control chart procedures, *Technometrics* **8**, 411–430.
 - [7] Carlstein, E., Müller, H. & Siegmund, D. (1994). *Change-Point Problems, Lecture Notes - Monograph Series*, Vol. 23, Institute of Mathematical Statistics.
 - [8] Zacks S. (1983). Survey of classical and Bayesian approaches to the change point problem: fixed sample and sequential procedures of testing and estimation, in *Recent Advances in Statistics*, M.A. Rizvi, J.S. Rustagi & D. Siegmund, eds, Academic Press, New York, pp. 245–269.
 - [9] Bather, J.A. (1967). On a quickest detection problem, *Annals of Mathematical Statistics* **38**, 711–724.
 - [10] Page, E.S. (1954). Continuous Inspection Schemes, *Biometrics* **41**, 1–9.
 - [11] Chernoff, H. & Zacks, S. (1964). Estimating the current mean of a normal distribution which is subjected to changes in time, *Annals of Mathematical Statistics* **40**, 999–1018.
 - [12] Zacks, S. & Kenett, R.S. (1994). Process tracking of time series with change points, in *Recent Advances in Statistics and Probability* (Proc. 4th IMSIBAC), J.P. Vilaplana & M.L. Puri, ed. VSP International Science Publishers, Utrecht, pp. 155–171.
 - [13] Tsiamyrtzis, P. & Hawkins, D.M. (2005). A Bayesian scheme to detect changes in the mean of a short run process, *Technometrics* **47**, 446–456.
 - [14] Kalman, R.E. (1960). A new approach to linear filtering and prediction problems, *Journal of Basic Engineering* **82**, 35–45.
 - [15] Hoadley, B. (1981). Quality Measurement Plan (QMP), *Bell System Technical Journal* **60**, 215–274.
 - [16] DeGroot, M.H. (1970). *Optimal Statistical Decisions*, McGraw-Hill, New York.
 - [17] Geisser, S. (1993). *Predictive Inference: An Introduction*, Chapman & Hall, London.

Related Articles

Bayesian Control Charts; Control Charts, Overview; Control Charts, Selection of; Process Capability Indices, Multivariate.

PANAGIOTIS TSIAMYRTZIS AND
DOUGLAS M. HAWKINS

Bayesian Reliability Analysis

Introduction

Bayesian methods embody the eighteenth century philosophical approach of Rev. Thomas Bayes (1702–1761), a Presbyterian minister and mathematician. The hallmark of Bayesian methods is the use of Bayes’ theorem to incorporate pertinent supplementary knowledge and information not contained in the sample data. Loosely, we define *Bayesian inference* as those statistical methods in which, prior to the consideration and use of sample data, the uncertainty regarding the unknown values of the parameters in an assumed sampling distribution for the data is expressed by means of a prior distribution. Gelman *et al.* [1] provide a nice introduction to Bayesian methods.

In the Bayesian approach, the model consists of two parts: the likelihood function and the prior distribution. The *likelihood function* is constructed from the *sampling distribution* of the data, which describes the probability of observing the data before the experiment is performed. Once we perform the experiment and observe the data, we can consider the sampling distribution as a function of the unknown parameters. This function (or any function proportional to it) is called the *likelihood function*.

We treat the unknown parameters in the likelihood function as random, and a **probability density function** describes our uncertainty about them. This probability density function is the joint *prior distribution* for the parameters. In a Bayesian reliability analysis, the likelihood function and the prior distribution are combined using Bayes’ theorem, forming the basis for parameter **estimation**, statistical inference, and test design.

One major philosophical difference from the classical frequency-based methods is the notion of probability that is considered. Classical methods use the well-known *relative frequency* notion of probability. In contrast, subjective probability forms the cornerstone of Bayesian methods, which consider probability to be a subjective assessment of the state of knowledge (also called *degree of belief*) about reliability parameters of interest.

A reliability analyst often has access to relevant supplementary information about the value of the unknown reliability parameters from such sources as physical/chemical theory, computational analysis, process development testing, past experience with similar devices, generic reliability data, and **expert opinion**. The use of such added resources is an extremely useful and powerful component in the Bayesian approach.

The prior distribution describes the analyst’s state of knowledge about the parameter’s value prior to obtaining the sample data x . We note that in reliability analysis, the data x often contain *censored* observations (see [2]). After obtaining the sample data, the uncertainty associated with the parameter is fully expressed by the *posterior distribution* that is calculated *via* Bayes’ theorem. For the case of continuous random variables, *Bayes’ theorem* combines the prior distribution and the likelihood; namely,

$$\pi(\theta|x) = \frac{f(x|\theta)\pi(\theta)}{\int f(x|\theta)\pi(\theta) d\theta} \quad (1)$$

where $\pi(\theta|x)$ denotes the posterior distribution of the parameter θ , $\pi(\theta)$ represents the prior distribution of θ , and $f(x|\theta)$ denotes the sampling distribution of x (or *likelihood function* when considered as a function of θ given x). Hamada *et al.* [3] presents a thorough treatment of Bayesian reliability analysis.

Appeal of Bayesian Reliability Methods

The use of deductive reasoning embodied in equation (1) is straightforward. Consequently, Bayesian reliability methods are easy to follow and the corresponding estimates are easy to interpret. For example, a Bayesian *credible interval* (or Bayesian **confidence interval**) may be directly interpreted as a probability statement about θ in contrast to a corresponding (frequency-based) confidence interval that has no such direct probabilistic interpretation.

There are several other appealing aspects of Bayesian reliability methods. In most cases, for large sample sizes of test data, the difference between Bayesian and classical inferences will be insignificant. However, when test data are scarce, the differences are often significant, and Bayesian credible intervals based on informative priors are often narrower than classical confidence intervals. If the prior

distribution is essentially noninformative, then the Bayesian and classical estimates are often virtually the same.

An important situation arises for a binomial or Poisson sampling model when no failures have occurred. In this case, the classical *maximum-likelihood* estimate of the binomial failure probability or Poisson failure rate is zero, which is clearly too optimistic. Although there exist a variety of classical methods for rectifying this shortcoming, Bayesian point estimates are naturally nonzero.

Because Bayesian interval estimates of parameters are true probability statements about unknown parameters, they may be easily propagated through complex system models, such as **Bayesian networks, fault trees**, event trees, and other logic models. Except in the simplest cases, it is difficult to propagate classical confidence intervals through such models. Features and nuances of real-world reliability problems, such as complex censoring, random hierarchical effects, and so on, can easily be accommodated and modeled by Bayesian methods (see the section titled “Example”). Such features are often difficult to consider when using frequentist methods.

Bayesian inference provides a natural and convenient method to update our state of knowledge about a parameter as future additional test data are obtained. Bayes’ theorem accomplishes this update. Also, the Bayesian *predictive distribution*

$$f(y|x) = \int f(y|\theta)\pi(\theta|x) d\theta \quad (2)$$

is the posterior distribution of a “future” sample observation y given the current sample data x . Such predictive distributions are extremely useful in practice.

Finally, using Bayesian methods, it is easy to numerically calculate the posterior distribution for any quantities of interest involving θ , such as mean time to failure (MTTF), **quantiles** of the lifetime distribution, reliability, and so on. This is perhaps the most useful and appealing feature of Bayesian reliability methods (see the section titled “Example”).

Bayesian computations, such as the integration in the denominator of Bayes’ theorem, can be complex; a decade or so ago they were often essentially impossible to perform. However, such computations are now straightforward using modern computer software based on **Markov chain Monte Carlo** (MCMC) algorithms. MCMC makes the Bayesian

approach to reliability feasible and easy to do, particularly for hierarchical problems (see the section titled “Example”).

There are several MCMC computer software packages available. Perhaps the most widely used package is WinBUGS [4], a Windows-based implementation of BUGS (Bayesian inference Using Gibbs Sampling). The package contains flexible software for the Bayesian analysis of complex statistical models using MCMC methods. It is available for free download from the Internet at <http://www.mrc-bsu.cam.ac.uk/bugs/>.

Example

We now illustrate an MCMC-based Bayesian reliability analysis. Gerstle and Kunz [5] provide the failure times (in hours) for pressure vessels that have been wrapped in Kevlar 49 fibers and subsequently tested at four different fiber stresses: 23.4, 25.5, 27.6, and 29.7 megapascals (MPa). The data are shown in Table 1. The fibers were taken from eight different spools of material (labeled 1–8). We are interested here in (a) determining the effect of stress on lifetime and (b) estimating the reliable life for 99% of the population (i.e. estimating the 0.01 quantile of the lifetime distribution).

In the Bayesian analysis below, we consider the eight spools to be a random sample from a population of such spools in which each spool potentially contributes a unique random “spool effect” to the lifetime. This random effect is called a *hierarchical* effect because we consider a second-order prior distribution (or *hyperprior*) on the unknown parameters of the distribution of this effect. Note also that 11 of the 23.4 MPa tests were terminated at 41 000 h (i.e. *time censored*) without having resulted in a failure. Under the assumption of a Weibull failure time distribution with a constant shape parameter for all four stress levels, these data were analyzed using classical (frequentist-based) methods by Crowder *et al.* [6]. Both the original and the Crowder analyses concluded that there is a significant spool effect.

We likewise assume a Weibull failure time model with a constant (but unknown) shape parameter. Let t_{ij} , $i = 1, \dots, 8$, $j = 1, \dots, n_i$, denote the failure time for the j th observation taken from the i th spool, and let t denote all the data (including the 11 time-truncated observations). Thus, conditional on

Table 1 Failure times of Kevlar 49 wrapped pressure vessels at four stress levels

Stress (MPa)	Spool	Failure times (h)
29.7	1	444.4, 755.2, 952.2, 1108.2
29.7	2	2.2, 8.5, 9.1, 10.2, 22.1, 55.4, 111.4, 158.7
29.7	3	12.5, 14.6, 18.7, 101.0
29.7	4	254.1, 1148.5, 1569.3, 1750.6, 1802.1
29.7	5	8.3, 13.3, 87.5, 243.9
29.7	6	6.7, 15.0, 144.0
29.7	7	4.0, 4.0, 4.6, 6.1, 7.9, 14.0, 45.9, 61.2
29.7	8	98.2, 590.4, 638.2
27.6	1	453.4, 664.5, 930.4, 1755.5
27.6	2	71.2, 199.1, 403.7, 432.2, 514.1, 544.9, 694.1
27.6	3	19.1, 24.3, 69.8, 136.0
27.6	4	876.7, 1275.6, 1536.8, 6177.5
27.6	6	514.2, 541.6, 1254.9
27.6	8	554.2, 2046.2
25.5	1	11 487.3, 14 032.0, 31 008.0
25.5	2	1134.3, 1824.3, 1920.1, 2383.0, 3708.9, 5556.0
25.5	3	1087.7, 2442.5
25.5	4	13 501.3, 29 808.0
25.5	5	11 727.1
25.5	6	225.2, 6271.1, 7996.0
25.5	7	503.6
25.5	8	2974.6, 4908.9, 7332.0, 7918.7, 9240.3, 9973.0
23.4	1	41 000, ^(a) 41 000, ^(a) 41 000, ^(a) 41 000 ^(a)
23.4	2	14 400.0
23.4	3	8616.0
23.4	4	41 000, ^(a) 41 000, ^(a) 41 000, ^(a) 41 000 ^(a)
23.4	5	9120.0, 20 231.0, 35 880.0
23.4	6	7320.0, 16 104.0, 20 233.0
23.4	7	4000.0, 5376.0
23.4	8	41 000, ^(a) 41 000, ^(a) 41 000 ^(a)

^(a)Time censored

Weibull scale parameter λ_{ij} and shape parameter β , we assume that

$$t_{ij} \sim \text{Weibull}(\lambda_{ij}, \beta) = \lambda_{ij} \beta t_{ij}^{\beta-1} \exp(-\lambda_{ij} t_{ij}^{\beta}),$$

$$t_{ij}, \lambda_{ij}, \beta > 0 \quad (3)$$

Let us consider a so-called *power-law* model for λ_{ij} , which, upon taking logarithms, becomes

$$\log(\lambda_{ij}) = \gamma_0 + \gamma_1 \log(s_{ij}) + \omega_i \quad (4)$$

where ω_i is the random contribution of the i th spool, s_{ij} denotes the value of stress s under which t_{ij} was obtained, and γ_0 and γ_1 are two regression parameters used to model the assumed linear dependency of $\log(\lambda_{ij})$ on $\log(s_{ij})$. Denoting the vector of spool effects as ω , we further assume that, conditional on λ_{ij} and β , the t_{ij} are statistically independent.

Consider the following **prior distributions**. Conditional on σ , we assume that $\omega_i \sim \text{Normal}(0, \sigma)$ and place a diffuse Gamma(0.001, 0.001) prior distribution on the associated *precision* parameter $\tau = 1/\sigma^2$ (in conformance with WinBUGS convention to consider Gaussian precision rather than variance). *A priori* we believe that β is close to 1.0; thus, we consider $\beta \sim \text{Gamma}(3, 3)$. We consider diffuse priors on the remaining two parameters γ_0 and γ_1 ; namely, $\gamma_0 \sim \text{Normal}(0, 10^3)$ and $\gamma_1 \sim \text{Normal}(0, 10^3)$.

We use MCMC to obtain a set of N joint posterior draws on $\beta, \gamma_0, \gamma_1, \omega$, and σ given t , which we denote as $\{(\beta_l, \gamma_{0l}, \gamma_{1l}, \omega_l, \sigma_l), l = 1, \dots, N\}$. The WinBUGS computer software discussed earlier was used to obtain $N = 10\,000$ of these posterior draws. The marginal posterior distributions calculated from these 10 000 draws for each of the 12 parameters are summarized in Table 2.

In Table 2 we see that the majority of the support for the marginal posterior distribution of β is slightly above 1.0. Likewise, by examining the marginal posteriors of γ_0 and γ_1 , we observe a rather strong effect of stress on the Weibull scale parameter. Indeed, we also see a strong spool effect by noting that the majority of the support for σ exceeds 1.0.

Table 2 Summaries of the marginal posterior distributions on $\beta, \gamma_0, \gamma_1, \omega$, and σ given t in the Kevlar example

Parameter	Mean	Standard deviation	5th percentile	Median	95th percentile
β	1.17	0.08	1.04	1.18	1.28
γ_0	-97.12	5.82	-104.90	-98.00	-83.74
γ_1	26.91	1.65	23.20	27.15	29.14
ω_1	-1.56	0.74	-2.71	-1.57	-0.38
ω_2	0.87	0.74	-0.20	0.83	2.02
ω_3	1.68	0.77	0.57	1.64	2.94
ω_4	-2.12	0.75	-3.30	-2.13	-0.94
ω_5	0.11	0.77	-1.06	0.09	1.32
ω_6	0.35	0.75	-0.77	0.32	1.56
ω_7	2.25	0.79	1.10	2.20	3.52
ω_8	-0.91	0.75	-2.05	-0.93	0.27
σ	1.76	0.61	1.06	1.63	2.89

4 Bayesian Reliability Analysis

The effect that spools 1, 3, 4, and 7 have on the scale parameter is especially pronounced.

Before continuing, one might question whether or not the preceding modeling assumptions are supported by the data. In other words, what is the *goodness of fit* of the model to the data? Using five bins, the value of the Bayesian χ^2 goodness-of-fit test statistic developed by Johnson [7] is computed to be 4.1. Comparing this to a χ^2 distribution with 4 **degrees of freedom**, we find that the p value for testing whether the above Weibull regression model “fits” the data is 0.39. Thus, there is ample evidence to suggest that the assumed model agrees with the data.

The effect of stress on lifetime can be seen by computing the posterior distribution of the Weibull MTTF as a function of s for a randomly selected spool of Kevlar material. Now the MTTF of the Weibull distribution in equation (3) is $\lambda^{-1/\beta} \Gamma(1 + 1/\beta)$ which, upon taking logarithms, substituting equation (4), and simplifying, becomes $-(\gamma_0/\beta) - (\gamma_1/\beta) \log(s) - (1/\beta)\omega + \log[\Gamma(1 + 1/\beta)]$. For the 10000 joint posterior MCMC draws on β , γ_0 , γ_1 , and σ given t , and a given value of s , the corresponding posterior draws on $\text{Log}(\text{MTTF})$ are given by

$$\left\{ \frac{-\gamma_{0\ell}}{\beta_\ell} - \left(\frac{\gamma_{1\ell}}{\beta_\ell} \right) \log s - \left(\frac{1}{\beta_\ell} \right) \log \omega_\ell + \log \left[\Gamma \left(1 + \frac{1}{\beta_\ell} \right) \right], \ell = 1, \dots, 10000 \right\} \quad (5)$$

where $\{\omega_\ell \sim \text{Normal}(0, \sigma_\ell), \ell = 1, \dots, 10000\}$ denotes a set of additional Gaussian posterior draws on ω . These additional draws represent the posterior variability in ω associated with a randomly selected spool. In Figure 1, we have plotted the median, 0.05, and 0.95 sample quantiles of the 10000 posterior draws in equation (5) as a function of stress s . Note the decreasing trend in the Weibull MTTF as s increases.

Now consider the reliable life for 99% of the population, that is, the 0.01 quantile of the lifetime distribution. The reliable life for 99% of the population distributed according to equation (3) is $\lambda^{-1/\beta} [-\log(0.99)]^{1/\beta}$, which, again upon taking logarithms, substituting equation (4), and simplifying, becomes $-(1/\beta)[\gamma_0 + \gamma_1 \log(s) + \omega + \log(0.99)]$. Thus, for the 10000 joint posterior MCMC draws on β , γ_0 , γ_1 , and σ given t , and a given value of s , the

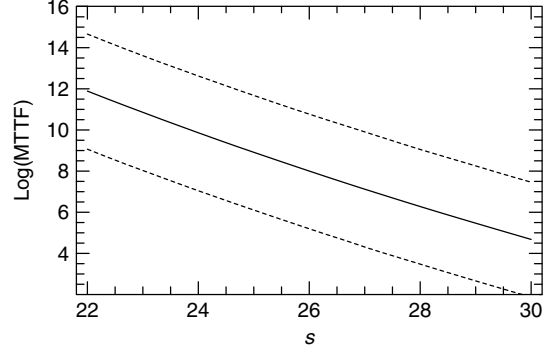


Figure 1 The median, 0.05, and 0.95 quantiles of the posterior distribution of Weibull $\text{Log}(\text{MTTF})$ as a function of stress s

corresponding posterior draws on $\text{Log}(\text{Reliable life})$ are given by

$$\left\{ \left(\frac{-1}{\beta_\ell} \right) [\gamma_{0\ell} + \gamma_{1\ell} \log(s) + \omega_\ell + \log(0.99)], \ell = 1, \dots, 10000 \right\} \quad (6)$$

In Figure 2, we have plotted the median, 0.05, and 0.95 sample quantiles of the 10000 posterior draws in equation (6) as a function of stress s . For example, at a stress level of 23 MPa, we are 95% certain that 99% of the population should survive at least $e^{4.4} = 81$ h.

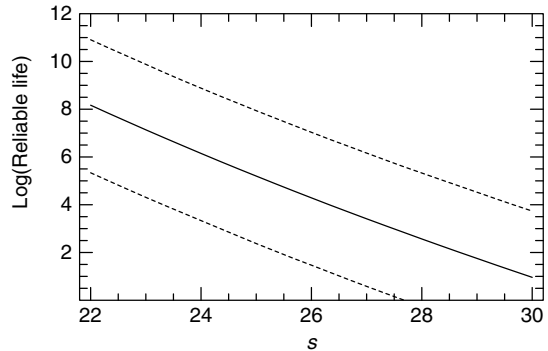


Figure 2 The median, 0.05, and 0.95 quantiles of the posterior distribution of Weibull $\text{Log}(\text{Reliable life})$ for $R = 0.99$ as a function of stress s

Summary

Bayesian methods are useful in reliability because of their ease of use (usually *via* MCMC sampling), ability to utilize other sources of relevant information (as prior information), ability to model real-world nuances, and clear probabilistic interpretation.

While hierarchical Bayesian methods place a second-order (hyperprior) distribution on unknown prior distribution parameters, *empirical Bayes* methods (see **Empirical Bayes Estimation of Reliability**) use the observed sampling data to directly estimate these parameters. These parameter estimates are usually produced using either *maximum likelihood* or the *method of moments*. It is generally agreed that empirical Bayes methods provide a good approximation to a full hierarchical Bayesian approach. However, the hierarchical Bayesian approach avoids several issues surrounding empirical Bayes; namely, the double counting of data encountered when estimating the prior parameters and the need to reflect the uncertainties in the estimated prior parameters in the posterior and predictive distributions.

References

- [1] Gelman, A., Carlin, J.B., Stern, H.S. & Rubin, D.B. (2004). *Bayesian Data Analysis*, 2nd Edition, Chapman & Hall, London.
- [2] Meeker, W.Q. & Escobar, L.A. (1998). *Statistical Methods for Reliability Data*, John Wiley & Sons, New York.
- [3] Hamada, M.S., Johnson, V.E., Martz, H.F., Reese, C.S. & Wilson, A.E. (2008). *Bayesian Analysis of Reliability Data*, Springer-Verlag, New York.
- [4] Gilks, W.R., Thomas, A. & Spiegelhalter, D.J. (1994). A language and program for complex Bayesian modeling, *The Statistician* **43**, 169–178.
- [5] Gerstle, F.P. & Kunz, S.C. (1983). Prediction of long-term failure in Kevlar 49 composites, in *Long-Term Behavior of Composites*, ASTM STP 813, T.K. O'Brien, ed, American Society for Testing and Materials, Philadelphia, pp. 263–292.
- [6] Crowder, M.J., Kimber, A.C., Smith, R.L. & Sweeting, T.J. (1991). *Statistical Analysis of Reliability Data*, Chapman & Hall, London.
- [7] Johnson, V.E. (2004). A Bayesian χ^2 test for goodness-of-fit, *Annals of Statistics* **32**, 2361–2384.

Related Articles

Bayesian Control Charts; Bayesian Reliability Demonstration; Empirical Bayes Estimation of Reliability; Hierarchical Markov Chain Monte Carlo (MCMC) for Bayesian System Reliability; Masked Failure Data: Bayesian Modeling; Software Reliability: Bayesian Analysis; Statistical Process Control, Bayesian.

HARRY F. MARTZ

Bayesian Design of Reliability Experiments

Introduction

Reliability is an important criterion in product and system design processes, as described for example in [1]. Engineers use models of various types to predict the reliability of proposed product and system designs. These include probabilistic system reliability models (e.g., [2]) and physical/empirical models such as the Paris law and Miner's rule for fatigue crack growth and finite element models to assess component strength and other physical characteristics. These models require quantitative inputs, which we will call *parameters*. Because of cost and time constraints, it is impossible to use empirical experiments (as distinguished from model-based **computer experiments**) to estimate all of the parameters relating to the reliability of a product or system. Some parameter values are available from vendor specifications, past experience, or from engineering judgment. These parameters are usually known to be within differing amounts of uncertainty, which may or may not be well understood. Sensitivity analysis and the use of engineering "factors of safety" and general conservatism are often used to deal with such uncertainties. When the available information is insufficient for particular components or parts of a system, additional empirical reliability experiments may be required. This need would arise, for example, when new materials or technologies are introduced or when existing components are to be used in new environments. Often, such experiments are censored (i.e., stopping after a fixed time period or a certain fractional failing), and they usually are "accelerated", as described, for example, in [3] and [4, Chapters 17–21].

For example, as described in [5], when testing microelectronic components at increased temperature (**accelerated life test** or ALT using temperature), engineers, if they have understanding of the failure mechanism, will sometimes assume that the activation energy is known. This is equivalent to specifying the slope of the linear relationship between log lifetime and reciprocal absolute temperature (the Arrhenius model). Of course, the activation energy is not known exactly and assuming that it is known gives a false sense of statistical precision. The extremes

(activation energy known exactly vs totally unknown) are illustrated in Figures 1 and 2, taken from an example in Chapter 19 of [4], indicating the large differences in the **estimation confidence intervals** of the failing probability function at use condition (100°C).

From a practical point of view, neither extreme is appropriate. There is, in most applications, useful but imperfect information about activation energy. An appropriate compromise analysis would use a prior distribution to describe the available information and to use Bayesian methods for inference. When planning such a reliability experiment, it is essential to account for the fact that the prior information will be used in the analysis of the experimental results and the prior should also be used in the planning. This can be seen because the optimum plan will be radically different under the two different assumptions (see, for example, [5]). Similar situations arise frequently in applications.

One basic idea of experimental design (Bayesian or not) is to use available knowledge of the process or phenomena to be studied to explore the kind of results that a particular design might produce. This is done by using analytical tools to evaluate important properties of the proposed design such as precision. For example, in non-**Bayesian reliability** testing, a useful metric for assessing a particular experimental design is the large-sample approximate variance of the maximum-likelihood estimator of a specified quantile of the life distribution. Within the Bayesian framework, a similar metric is to compute the expected value of the posterior variance of a quantity of interest, where the expectation is with respect to the data yet to be seen, according to the prior distribution. This is known as *preposterior analysis*.

Some Related Literature

Chapter 6 of [3] and Chapters 10, 20, and 22 of [4] give theory and methods for planning non-Bayesian reliability tests. Classical work in the area of Bayesian design of experiments and preposterior analysis is covered, for example, in [6]. Chaloner and Verdinelli [7] give a broad review of Bayesian design methods. Chaloner and Larntz in [8] and [9] describe Bayesian designs for logistic regression and accelerated testing, respectively. Clyde *et al.* [10] describe the Bayesian

2 Bayesian Design of Reliability Experiments

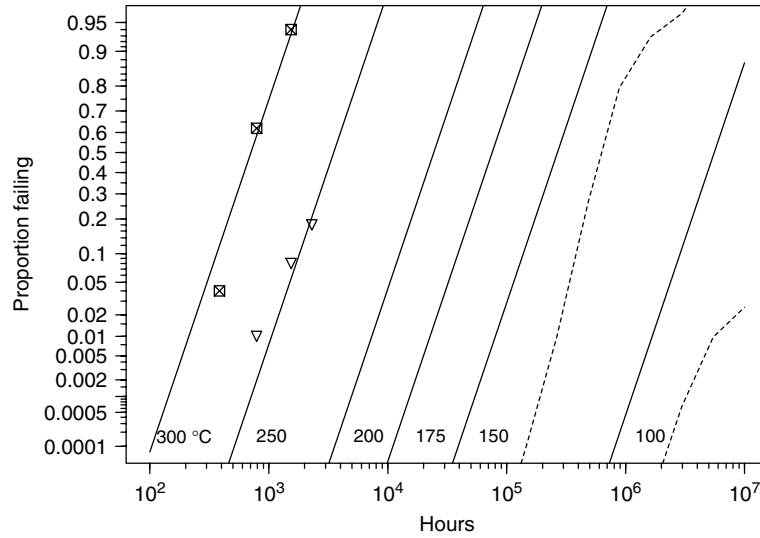


Figure 1 Accelerated test results, with the estimated proportional failing function *versus* operating hours at various temperatures, for an integrated circuit with activation energy assumed to be unknown. The dashed lines are 95% confidence intervals for the function at 100 °C [© John Wiley & Sons Ltd, New York, 1998.]

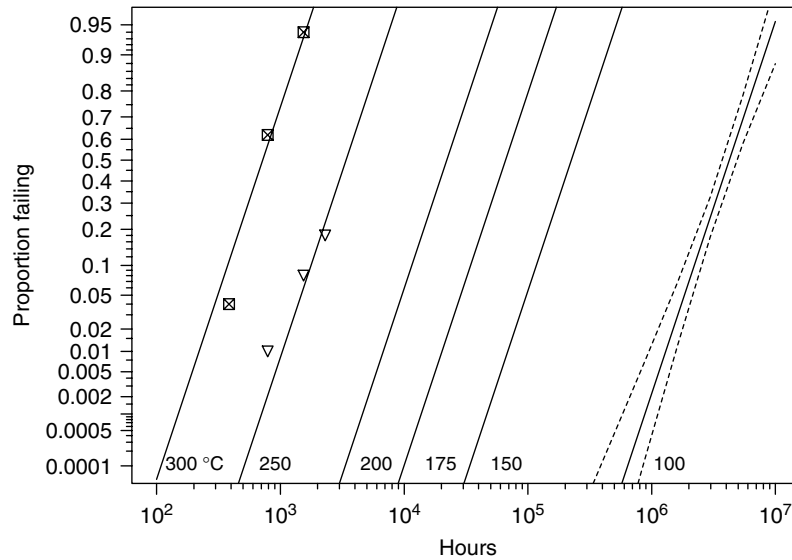


Figure 2 Accelerated test results, with the estimated proportional failing function *versus* operating hours at various temperatures, for an integrated circuit with activation energy assumed to be known. The dashed lines are 95% confidence intervals for the function at 100 °C [© John Wiley & Sons Ltd, New York, 1998.]

design methods for heart defibrillators, under a logistic regression model. They use a utility function based on preposterior expectation and large-sample approximation for both fixed-point and sequential designs.

Polson [11] describes a general decision theory for the ALT Bayesian design problem and proposes a preposterior expected information-based utility function. Verdinelli *et al.* [12] describe ALT Bayesian

design methods for predictions using utility functions based on Shannon information. Erkanli and Soyer [13] present optimum ALT Bayesian designs for the exponential lifetime distribution with no censoring, by adopting the curve-fitting optimization approaches developed in [14] and [15]. Zhang and Meeker [5] apply Bayesian methods to design an ALT with a log-location-scale distribution and censoring, using preposterior precision on a quantity of interest and large-sample approximation to provide optimum and optimized compromise test plans. Hamada *et al.* [16] show how to find near-optimal Bayesian designs for regression models by using **genetic algorithms**.

The Model

The model and presentation in this article is general, implying wide applicability of the approach outlined here. Suppose that an experiment is to be conducted to estimate the reliability of a product or system under a model that is described by a parameter vector θ of length p . More specifically, suppose that interest centers on a scalar function $g(\theta)$. This, for example, could be the reliability of the system, perhaps at a point in time, if reliability is time dependent. Alternatively, it could be a quantile of the life distribution of the product or the time to first failure of the system.

The prior information on θ is quantified by a joint distribution $\omega(\theta)$ that has a $p \times p$ variance–**covariance** matrix denoted by $\mathbf{S}$. In practice, prior information is usually specified or elicited from an engineer using parameters that are familiar and approximately independent. $\omega(\theta)$ can then be obtained by transformation from these identifiable parameters. Planning results are typically not highly sensitive to the particular shape of the prior distribution.

The Reliability Experiment

There are different types of reliability experiments that can be used to obtain information about θ . For example, [3] describes accelerated life tests and destructive **degradation** tests. Chapter 21 of [4] describes repeated measured degradation tests. The specifics of the model, type of response (e.g., time to failure or degradation), experimental conditions (e.g., temperature levels), **sample size**, and allocation of

the units to the different levels of the experimental variables determine the amount of information (or more precisely, the distribution of the amount of information) that one will get from a proposed experiment. We will let $\mathbf{D}$ denote a particular specified experimental plan.

Quantifying Precision

The expected information from an experimental plan $\mathbf{D}$ can be quantified by a $p \times p$ Fisher information matrix

$$\mathbf{I}_\theta(\mathbf{D}) = \mathbb{E} \left(-\frac{\partial^2 \ell(\theta, \mathbf{D})}{\partial \theta \partial \theta^T} \right) \quad (1)$$

where $\ell(\theta, \mathbf{D})$ is the log likelihood for the model/data combination and T denotes vector transpose. In general, $\mathbf{I}_\theta(\mathbf{D})$ depends on the model and the model parameters. Computing $\mathbf{I}_\theta(\mathbf{D})$ for a given test plan is more complicated when data are censored. Escobar and Meeker [17] describe methods and algorithms for computing $\mathbf{I}_\theta(\mathbf{D})$ for different kinds of censored data.

With a reasonably large sample size, the multivariate **normal distribution** provides a good approximation for the posterior distribution. Then the posterior variance can be expressed as a simple combination of information from the prior distribution and the data. Because $\mathbf{S}$ is the variance–covariance matrix of the prior distribution for θ , $\mathbf{S}^{-1}$ is the prior precision matrix for θ . Then the posterior variance of θ can be approximated by

$$\text{Var}(\theta|\mathbf{t}, \mathbf{D}) \approx [\mathbf{S}^{-1} + \hat{\mathbf{I}}_\theta(\mathbf{D})]^{-1} \quad (2)$$

where $\hat{\mathbf{I}}_\theta(\mathbf{D})$ is $\mathbf{I}_\theta(\mathbf{D})$ evaluated at $\hat{\theta}$, the maximum-likelihood estimator of θ . The posterior variance in equation (2) depends on the test plan $\mathbf{D}$ and the future data $\mathbf{t}$ only through the Fisher information matrix from the experiment.

The posterior variance of $g(\theta)$, the function of interest, can be approximated by

$$\text{Var}(g|\mathbf{t}, \mathbf{D}) \approx (\hat{\mathbf{g}}')^T \text{Var}(\theta|\mathbf{t}, \mathbf{D}) (\hat{\mathbf{g}}') \quad (3)$$

where $\hat{\mathbf{g}}'$ is the gradient vector $\mathbf{g}'(\theta) = \partial g / \partial \theta$, again evaluated at $\hat{\theta}$. Note that the approximate posterior variance in equation (3) depends on the data only through $\hat{\theta}$.

For Bayesian test planning, we need to use the preposterior expectation of $\text{Var}(g|\mathbf{t}, \mathbf{D})$ with respect to

4 Bayesian Design of Reliability Experiments

$\mathbf{t}$ under design $\mathbf{D}$. Taking the expectation with respect to the data is intractable. Instead, following [10], one can use a large-sample approximation,

$$E_{\mathbf{t}|\mathbf{D}}[\text{Var}(g|\mathbf{t}, \mathbf{D})] \approx \int \text{Var}(g|\mathbf{t}, \mathbf{D})_{\hat{\theta}=\theta} d(\omega(\theta)) \quad (4)$$

As explained in [5] the approximation in equation (4) is valid even with moderate-sized samples, when the variability in the prior for θ dominates the sampling variability in its estimates from the data. This integral has the same dimension as θ and for typical reliability experiments its evaluation is more manageable. In particular, if the number of parameters is small (e.g., less than five or six) it can be evaluated by using standard numerical quadrature methods. If the dimension of the integral is large (e.g., in a complicated system with multiple failure modes corresponding to components with different parameters), one could use a **Monte Carlo**-based algorithm to evaluate $E_{\mathbf{t}|\mathbf{D}}[\text{Var}(g|\mathbf{t}, \mathbf{D})]$.

Although it is possible to use other utility functions in the design problem, minimizing the expected posterior variance in equation (4) is meaningful and corresponds to what has been used in previous non-Bayesian design of reliability experiments work.

Design Optimization

There is much work in the literature on optimum design of experiments (e.g., [18] and [19]). The basic ideas used in traditional optimum design area are to optimize some measure of precision (e.g., minimize the variance of an estimator or maximize the determinant of an information matrix), subject to constraints (e.g., number of test units, time, and experimental region). Literature and examples of such optimum plans for non-Bayesian reliability applications are described in Chapter 6 of [3] and Chapter 20 of [4]. Within the Bayesian framework, there is a similar optimization to maximize some Bayes utility such as preposterior precision. Applications in reliability testing include [5, 9, 11–13].

In the area of optimum design, there is a well-known result that for a model that is linear in the parameters, there is an upper bound on the number of design points and that upper bound is equal to the number of parameters in the model. For models that are nonlinear in the parameters, including Bayesian optimum design in nonlinear

problems, the number of design points generally has no theoretical upper bound. This is discussed, for example, in [19, Chapter 19] and [7, Section 5.2]. To confirm that an optimum design has been found one can use a sequential numerical search with different numbers of design points and apply the general equivalence theorem (GET) from [20] to verify the global optimality of a particular plan. Illustrations of this approach are given in [5, 8].

Optimum plans provide useful insight into how to plan an experiment. Such plans, however, have practical deficiencies. For example, they generally provide little **power** to detect departures from the assumed model and they tend to be more sensitive to misspecification of the model and other planning information. Meeker and Escobar [4] illustrate the use of compromise test plans for more **robust design** in non-Bayesian applications. Zhang and Meeker [5] describe and evaluate compromise Bayesian test plans for an accelerated life test experiment.

Prior Distributions

An important (and somewhat controversial) part of Bayesian analysis is the specification of the prior information. When little or no prior information is available, it is common practice to use diffuse or flat **prior distributions**, resulting in an analysis for which the assumed prior is supposed to have no important effect on the posterior distribution and resulting inferences. Various technical difficulties can, however, arise when a diffuse prior distribution is combined with a dataset containing little information.

In preposterior analysis, the prior distribution is used twice – in the computation of the posterior distribution and in the computation of the expected posterior variance of the quantity of interest. Technical problems arise immediately in the preposterior expectation if one attempts to use one single diffuse prior distribution. An alternative approach is to use two different prior distributions. This can be justified because the risk associated with incorrect prior specification may be different for those conducting the experiment and those who will depend on the results of the experiment. This was, in effect, done in [9] with diffuse priors for computing the posterior distribution and informative prior distributions for computing expectations of the quantities of interest.

Acknowledgments

Figures were reproduced, with permission, from Figures 19.13 and 19.15 of [4].

References

- [1] Meeker, W.Q. & Escobar, L.A. (2004). Reliability: the other dimension of quality, *Quality Technology and Quality Management* **1**, 1–25.
- [2] Høyland, A. & Rausand, M. (1994). *System Reliability Theory: Models and Statistics Methods*, John Wiley & Sons, New York.
- [3] Nelson, W. (1990). *Accelerated Testing: Statistical Models, Test Plans, and Data Analyses*, John Wiley & Sons, New York.
- [4] Meeker, W.Q. & Escobar, L.A. (1998). *Statistical Methods for Reliability Data*, John Wiley & Sons, New York.
- [5] Zhang, Y. & Meeker, W.Q. (2006). Bayesian methods for planning accelerated life tests, *Technometrics* **48**, 49–60.
- [6] Raiffa, H. & Schlaifer, R. (1961). *Applied Statistical Decision Theory*, MIT Press, Cambridge.
- [7] Chaloner, K. & Verdinelli, I. (1995). Bayesian experimental design: a review, *Statistical Science* **10**, 273–304.
- [8] Chaloner, K. & Larntz, K. (1989). Optimal Bayesian design applied to logistic regression experiments, *Journal of Statistical Planning and Inference* **21**, 191–208.
- [9] Chaloner, K. & Larntz, K. (1992). Bayesian design for accelerated life testing, *Journal of Statistical Planning and Inference* **33**, 245–259.
- [10] Clyde, M., Müller, P. & Parmigiani, G. (1995). Optimal Design for Heart Defibrillators, in *Case Studies in Bayesian Statistics 2*, C. Gatsonis, J. Hodges, R.E. Kass & N. Singpurwalla, eds, Springer, New York, pp. 278–292.
- [11] Polson, N.G. (1993). A Bayesian perspective on the design of accelerated life tests, in *Advances in Reliability*, A.P. Basu, ed, Elsevier, New York, pp. 321–330.
- [12] Verdinelli, I., Polson, N.G. & Singpurwalla, N.D. (1993). Shannon information and Bayesian design for prediction in accelerated life testing, in *Reliability and Decision Making*, R.E. Barlow, C.A. Clariotti & F. Spizzichino, eds, Chapman & Hall, London, pp. 247–256.
- [13] Erkanli, A. & Soyer, R. (2000). Simulation-based designs for accelerated life tests, *Journal of Statistical Planning and Inference* **90**, 335–348.
- [14] Müller, P. & Parmigiani, G. (1995). Optimal design via curve fitting of Monte Carlo experiments, *Journal of the American Statistical Association* **90**, 1322–1330.
- [15] Müller, P. (1999). Simulation Based Optimal Design, in *Bayesian Statistics 6*, J.M. Bernardo, J.O. Berger, A.P. Dawid & A.F.M. Smith, eds. Oxford University Press, pp. 459–474.
- [16] Hamada, M., Martz, H.F., Reese, C.S. & Wilson, A.G. (2001). Finding near-optimal Bayesian experimental designs via genetic algorithms, *The American Statistician* **55**, 175–181.
- [17] Escobar, L.A. & Meeker, W.Q. (1998). The asymptotic covariance matrix for maximum likelihood estimators with models based on location-scale distributions involving censoring, truncation, and explanatory variables, *Statistica Sinica* **8**, 221–237.
- [18] Fedorov, V.V. (1972). *Theory of Optimal Experiments*, Academic Press, New York.
- [19] Atkinson, A.C. & Donev, A.N. (1992). *Optimum Experimental Designs*, Clarendon, Oxford.
- [20] Whittle, P. (1973). Some general points in the theory of optimal experimental design, *Journal of the Royal Statistical Society. Series B* **35**, 123–130.

Related Article

Bayesian Reliability Analysis.

WILLIAM Q. MEEKER AND YAO ZHANG

Empirical Bayes Estimation of Reliability

Introduction

Assessment of the reliability of various types of equipment relies on statistical inference about characteristics of reliability such as reliability function, mean lifetime of the devices, or failure rate. General techniques of statistical inference (estimation and hypotheses testing) are reviewed in **Estimation; Least-Squares Estimation; Maximum Likelihood; Nonparametric Tests; Hypothesis Testing; and Significance Level**, respectively.

Consider the following model. Let $T = (T_1, T_2, \dots, T_m)$, $T_i \in \mathcal{T}$, be independent identically distributed (i.i.d.) observations with the joint **probability density function** (pdf) $f(t|\theta)$, $t \in \mathcal{T}^m$, depending on an unknown parameter $\theta \in \Theta$. If we are interested in estimating θ or a known function $u(\theta)$ of θ , we can, for example, use the maximum-likelihood estimation (MLE) of θ described in **Maximum Likelihood**. This technique is, however, of very little help if we want to accommodate some *prior* information about θ which may be available from previous experiments, expert opinions, or other sources of knowledge. To take advantage of this information, we assume that parameter θ is also a random variable or vector, and the particular value of θ associated with our sample comes from a pool of possible values of θ , which have a *prior* pdf $g(\theta)$. Using Bayes formula one can obtain the updated *posterior* pdf of θ given the sample T as

$$g(\theta|T) = \frac{[f(T|\theta)g(\theta)]}{\left[\int_{\Theta} f(T|\theta)g(\theta) d\theta \right]} \quad (1)$$

The estimator

$$\hat{u}(T) = \int_{\Theta} u(\theta)g(\theta|T) d\theta \quad (2)$$

is called the *Bayes estimator* of $u(\theta)$ and it minimizes the quadratic risk

$$\begin{aligned} R(g) &= E(\hat{u}(T) - u(\theta))^2 \\ &= \int_{\mathcal{T}^m} \int_{\Theta} (\hat{u}(T) - u(\theta))^2 \end{aligned}$$

$$\times f(T|\theta)g(\theta) dT d\theta \quad (3)$$

Here E is the expectation with respect to T and θ (the concept of expectation is introduced in **Expectation**). The prior density $g(\theta)$ is usually chosen from some familiar distribution family, or by some other considerations (to minimize the information introduced by a prior distribution or to make integration in equation (2) easy, etc.). Extensive introduction to Bayesian methods can be found in [1].

In practice, however, none of the choices for the prior may offer sufficient motivation. Moreover, we may have some prior measurements that are indirectly related to the recent data of interest. In particular, suppose that n “past” data sets are available,

$$\begin{aligned} T_i &= (T_{i1}, T_{i2}, \dots, T_{im_i}), \quad T_{ij} \in \mathcal{T}, \\ i &= 1, \dots, n \end{aligned} \quad (4)$$

where, conditional on θ_i , vectors T_i have pdfs $f_i(t|\theta_i)$, and we also have a “present” data set $T = (t_1, \dots, t_m)$, which carries information about the parameter of interest τ . Now, if we assume that $\theta_1, \theta_2, \dots, \theta_n$ and τ are not just fixed constants unrelated to each other but realizations of a random variable with the unknown pdf $g(\theta)$, we can use past data $T_i, i = 1, \dots, n$, for assessment of the value of τ or some function $u(\tau)$ of τ . Decision theoretic procedures that utilize past data as a means for bypassing the necessity of identifying a completely unknown prior distribution are called *empirical Bayes* (EB) decision procedures [2, p. 605].

Empirical Bayes Estimation

Estimator and its Risk

The EB estimation was introduced by Robbins [3] and became increasingly popular in various areas of statistics (see, e.g., [4] and [5] for a general review of EB methods). To simplify our introduction, we consider the case when all conditional densities $f_i(t|\theta)$ are identical and equal to $f(t|\theta)$ (the case of different conditional densities is treated in, e.g., [6]). Let one observe random vectors $(T_1, \theta_1), \dots, (T_n, \theta_n)$, where each θ_i is distributed according to some unknown prior distribution G with the pdf $g(\theta)$ and, given θ_i , T_i has a known conditional density function $f(t|\theta_i)$, $i = 1, \dots, n$. In each pair the first

2 Empirical Bayes Estimation of Reliability

component is observable, but the second is not. For the observation T with the conditional pdf $f(t|\tau)$, the goal is to estimate τ or a function $u(\tau)$. Here (T, τ) can be a new pair of observations $T \equiv T_{n+1}$, $\tau = \theta_{n+1}$, or one of the old ones $T \equiv T_k$, $\tau = \theta_k$, $k = 1, \dots, n$.

If the prior pdf $g(\theta)$ were known, the best estimator for $u(\tau)$ would be the Bayes estimator

$$\hat{u}(T) = \Phi(T)/p(T) \quad (5)$$

where

$$\begin{aligned} \Phi(t) &= \int_{\Theta} u(\theta) f(t|\theta) g(\theta) d\theta, \\ p(t) &= \int_{\Theta} f(t|\theta) g(\theta) d\theta \end{aligned} \quad (6)$$

Here, $p(t)$ is the marginal pdf of T at the point t . As the prior density $g(\theta)$ is unknown, we construct an EB estimator $\hat{u}_n(T) \equiv \hat{u}_n(T; T_1, \dots, T_n)$ as the estimator of the right-hand side of formula (5) based on observations $T_1, T_2, \dots, T_n$. Denote by E_T and E_T the expectations with respect to T and $T_1, \dots, T_n$, respectively. An EB estimator $\hat{u}_n(T)$ can be characterized by its quadratic risk (see [5])

$$R(\hat{u}_n; g) = E_T E_T (\hat{u}_n(T) - u(\tau))^2 \quad (7)$$

which can be partitioned into two components. The first component

$$\begin{aligned} R(g) &= E_T (\hat{u}(T) - u(\tau))^2 \\ &= \inf_v E_T (v(T) - u(\tau))^2 \end{aligned} \quad (8)$$

is the risk of the Bayes estimator, while the second component

$$\hat{R}(\hat{u}_n; g) = E_T E_T (\hat{u}_n(T) - \hat{u}(T))^2 \quad (9)$$

is the error of estimating $\hat{u}(T)$. Robbins [3] called an EB estimator *asymptotically optimal* if $\lim_{n \rightarrow \infty} \hat{R}(\hat{u}_n; g) = 0$.

All techniques for EB estimation can be roughly divided into two classes: *parametric* and *nonparametric*. Parametric methods are based on the assumption that $g(\theta)$ has a known parametric form with one or several unknown parameters. In nonparametric methods, $g(\theta)$ is assumed to be completely unknown and unspecified.

Parametric EB Estimation

Let $g(\theta)$ have a known form with a vector or scale parameter σ , which is unknown: $g(\theta) \equiv g(\theta|\sigma)$. In this case, $\Phi(t)$ and $p(t)$ in formula (6) also have known forms

$$\begin{aligned} \Phi(t) &\equiv \Phi(t|\sigma) = \int_{\Theta} u(\theta) f(t|\theta) g(\theta|\sigma) d\theta, \\ p(t) &\equiv p(t|\sigma) = \int_{\Theta} f(t|\theta) g(\theta|\sigma) d\theta \end{aligned} \quad (10)$$

Hence, the Bayes estimator $\hat{u}(T)$ is a known function $\hat{u}(T|\sigma)$ of σ and T , and the problem reduces to estimation of a parametric function $\hat{u}(T|\sigma)$ on the basis of observations $T_1, \dots, T_n$ with the pdf $p(t|\sigma)$. This is a standard statistical problem. We can, for example, construct the MLE $\hat{\sigma}$ of σ on the basis of observations $T_1, \dots, T_n$ (see **Maximum Likelihood**) and then estimate $\hat{u}(T)$ by $\hat{u}_n(T) = \hat{u}(T|\hat{\sigma})$.

In addition, we can calculate the posterior pdf of θ

$$g(\theta|T, \sigma) = \frac{[f(T|\theta)g(\theta|\sigma)]}{[\int_{\Theta} f(T|\theta)g(\theta|\sigma) d\theta]} \quad (11)$$

and estimate it by $\hat{g}_n(\theta|T) = g(\theta|T, \hat{\sigma})$. We can use $\hat{g}_n(\theta|T)$ for construction of the $(1 - \gamma)$ credibility interval $(\hat{\tau}_l, \hat{\tau}_r)$ for τ by choosing $\hat{\tau}_l$ and $\hat{\tau}_r$ such that

$$\int_{\hat{\tau}_l}^{\hat{\tau}_r} \hat{g}_n(\theta|T, \hat{\sigma}) d\theta = 1 - \gamma \quad (12)$$

(see **Confidence Intervals** for the review of confidence and credibility intervals). The credibility interval for $u(\tau)$ is of the form $(\hat{u}_l, \hat{u}_r)$, where $u(\tau) \in (\hat{u}_l, \hat{u}_r)$ whenever $\tau \in (\hat{\tau}_l, \hat{\tau}_r)$.

Parametric EB estimation is reviewed in [7] and [8] as well as in [4] and [5].

Nonparametric EB Estimation

If $g(\theta)$ is completely unknown, then we can construct a nonparametric estimator of $\hat{u}(T)$. One possible way of doing this is to estimate $g(\theta)$ by $\hat{g}(\theta)$ on the basis of the observations $T_1, \dots, T_n$ and then plug $\hat{g}(\theta)$ into formulae (5) and (6). The first attempt in this direction was to estimate the unobserved parameters θ_i by $\hat{\theta}_i$ on the basis of T_i , $i = 1, \dots, n$, and then use "observations" $\hat{\theta}_i$ for estimation of $g(\theta)$ (see, e.g., [2]). The difficulty, though, is that in order for

$\hat{\theta}_i$ to be consistent estimators of θ_i , the past data should contain series of observations, that is, m_i 's in (4) should be relatively large. This drawback can be avoided by constructing an estimator of $g(\theta)$ directly on the basis of $T_1, \dots, T_n$ (see, e.g., Maritz and Lwin [5] or Walter [9]). However, the shortcoming of any approach that is based on estimating prior density $g(\theta)$ is that estimation of $g(\theta)$ is a much harder task than estimation of $\hat{u}(T)$ itself. Consequently, the risk $\hat{R}(\hat{u}_n; g)$ of the resulting estimator will converge to zero very slowly as $n \rightarrow \infty$.

In order to avoid estimation of $g(\theta)$, one can estimate $u(\tau)$ directly. For example, if $u(\tau) = \tau$ and $f(t|\theta)$ belongs to a one-parameter exponential family

$$f(t|\theta) = w(t)C(\theta)\exp(\theta t) \quad \text{with} \\ u(t) > 0 \quad \text{for } t > a \quad (13)$$

Singh [10] observed that

$$\hat{\tau} = \frac{p'(T)}{p(T)} \quad (14)$$

where $p(t)$ is the marginal density of T defined in formula (6). Note that the gamma and the normal pdfs belong to the family (13).

Estimation of the pdf $p(t)$ and its derivative $p'(t)$ at a point t is a very common problem in statistics and is addressed in a variety of papers and monographs (see, e.g., [11] for a review). One of the possibilities is to construct kernel estimators for $p(t)$ and $p'(t)$

$$\hat{p}_n(t) = \frac{1}{nh} \sum_{i=1}^n K_0\left(\frac{t - T_i}{h}\right), \\ \hat{p}'_n(t) = \frac{1}{nh^2} \sum_{i=1}^n K_1\left(\frac{t - T_i}{h}\right) \quad (15)$$

where h is small and the kernels K_j , $j = 0, 1$, satisfy the conditions $\int t^i K_j(t) dt = 1$ if $i = j$ and $\int t^i K_j(t) dt = 0$ if $i \neq j$, $i = 0, 1, \dots, r$. Here $r \geq 2$ is a positive integer.

The approach of Singh [10] can be generalized as follows (see [6]). Note that in order to construct an EB estimator $\hat{u}_n(T)$ one needs to estimate the numerator and the denominator in formula (5), and then estimate the ratio. The denominator $p(T)$ can be estimated using, for example, the first formula in (15). To estimate the numerator

$\Phi(T) = \int_{\Theta} u(\theta) f(T|\theta) g(\theta) d\theta$ one needs to find a sequence of functions $\varphi_\varepsilon(t)$ such that $\int \varphi_\varepsilon^2(t) dt < \infty$ for every $\varepsilon > 0$ and for every s and θ

$$\int f(t|\theta) \varphi_\varepsilon(s, t) dt - u(\theta) \rightarrow 0 \quad \text{as } \varepsilon \rightarrow 0 \quad (16)$$

Then $\Phi(T)$ can be estimated by

$$\hat{\Phi}_n(T) = n^{-1} \sum_{i=1}^n \varphi_\varepsilon(T, T_i) \quad (17)$$

with small $\varepsilon = \varepsilon_n$. Finding the sequence $\varphi_\varepsilon(s, t)$ may be tricky; however, luckily, in many particular situations, for example, when θ is the location or a scale parameter of the family $f(t|\theta)$,

$$\int_T f(t|\theta) \varphi_\varepsilon(s, t) dt = u(\theta) \quad (18)$$

has an exact solution $\varphi(s, t)$ (see, e.g., [12, 13]).

The last step is estimation of the ratio $\hat{u}(T) = \Phi(T)/p(T)$ by

$$\hat{u}_n(T) = H(\hat{\Phi}_n(T)/\hat{p}_n(T), \delta) \quad (19)$$

where $\delta > 0$ is a small parameter. Here, function $H(a, \delta)$ is chosen to ensure that the overall **mean squared error** tends to zero as n grows: $H(a, 0) = a$ and $H(a, \delta)$ is bounded for any $\delta > 0$. Singh [10] proposed the use of

$$H(a, \delta) = aI(|a|\delta \leq 1) + \delta^{-1}I(a\delta > 1) \\ - \delta^{-1}I(a\delta < -1) \quad (20)$$

where $I(A)$ is the indicator function of the set A . Pensky [6, 12] suggested a smooth version

$$H(a, \delta) = a(1 + \delta|a|^{2\tau})^{-s}, \\ \tau = 1, 2, \dots; \quad s > 0; 2\tau s \geq 1 \quad (21)$$

EB Estimation of the Reliability Function and the Mean Lifetime

Let $T_1, \dots, T_n$ be observations on the lifetimes of the devices of a certain type corresponding to unobserved parameters $\theta_1, \dots, \theta_n$ and let T be a measurement with the pdf $f(t|\tau)$. The general ideas reviewed

4 Empirical Bayes Estimation of Reliability

above can be used for EB estimation of *the mean lifetime*

$$m(\tau) = \int_0^\infty t f(t|\tau) dt \quad (22)$$

and *the reliability function* at a point t_0

$$\begin{aligned} R(\tau, t_0) &= P(T \leq t_0) = \int_0^{t_0} f(t|\tau) dt \\ &= \int_0^\infty I(t \leq t_0) f(t|\tau) dt \end{aligned} \quad (23)$$

For this purpose, one should carry out EB estimation of $u_1(\tau) = m(\tau)$ and $u_2(\tau) = R(\tau, t_0)$ where t_0 is treated as a fixed known constant.

Consider, for example, estimation of the mean lifetime and the reliability function in the case when $f(t|\theta)$ belongs to the family of generalized gamma distributions

$$\begin{aligned} f(t|\theta) &= \frac{\alpha t^{\alpha\nu-1}}{\theta^\nu \Gamma(\nu)} \exp\left(-\frac{t^\alpha}{\theta}\right) I(t > 0), \\ \theta &\in \Theta \end{aligned} \quad (24)$$

where α and ν are known parameters, $\alpha\nu \geq 1$, and the parameter space Θ is a subset of $(0, \infty)$, the positive real line. The family of pdfs (24) includes most of the pdfs used in reliability, for example, the exponential ($\alpha = \nu = 1$), the gamma ($\alpha = 1$), and the Weibull ($\nu = 1$) pdfs. It is easy to calculate that for the family of pdfs (24) we have

$$m(\tau) = \tau^{1/\alpha} [\Gamma(\nu)]^{-1} \Gamma(\nu + \alpha^{-1}) \quad (25)$$

$$R(\tau, t_0) = [\Gamma(\nu)]^{-1} \gamma(\nu, \tau^{-1} t_0^\alpha) \quad (26)$$

where $\gamma(a, x)$ is the incomplete gamma function (see [14, p. 260]).

Parametric EB

Consider an example of parametric EB estimation of $m(\tau)$ and $R(\tau, t_0)$ when $g(\theta)$ is the inverted gamma prior of the form

$$\begin{aligned} g(\theta) &\equiv g(\theta|\sigma) = [\Gamma(\beta)]^{-1} \sigma^\beta \theta^{-(\beta+1)} \\ &\times \exp(-\sigma/\theta) I(\theta > 0) \end{aligned} \quad (27)$$

This prior was considered by many authors, for example, Dey and Kuo [15] and Li [16]. Here, for simplicity, we assume that parameter σ is unknown

while parameter β is known. In this case, the marginal pdf of T is

$$\begin{aligned} p(T) &\equiv p(T|\sigma) = \alpha [B(\nu, \beta)]^{-1} \\ &\times \sigma^\beta T^{\alpha\nu-1} (T^\alpha + \sigma)^{-(\nu+\beta)} I(T > 0) \end{aligned} \quad (28)$$

Hence, the MLE $\hat{\sigma}$ of σ is the value minimizing the likelihood function

$$L(\sigma; T_1, \dots, T_n) = \prod_{i=1}^n p(T_i|\sigma) \quad (29)$$

(see **Maximum Likelihood**) and $\hat{\sigma}$ is the root of the equation

$$\frac{1}{n} \sum_{i=1}^n \frac{T_i^\alpha}{T_i^\alpha + \hat{\sigma}} = \frac{\nu}{\nu + \beta} \quad (30)$$

Therefore, Bayes estimators of $m(\tau)$ and $R(\tau, t_0)$ are given by formulae (5) and (6) with $u_1(\theta) = m(\theta)$ and $u_2(\theta) = R(\theta, t_0)$ (see formulae (25), (26), and (28)):

$$\begin{aligned} \hat{m}(T) &= [\Gamma(\nu)\Gamma(\nu + \beta)]^{-1} \Gamma(\nu + 1/\alpha) \\ &\times \Gamma(\nu + \beta - 1/\alpha) (T^\alpha + \sigma)^{1/\alpha} \end{aligned} \quad (31)$$

$$\begin{aligned} \hat{R}(T, t_0) &= \frac{\Gamma(2\nu + \beta)}{\Gamma(\nu + 1)\Gamma(\nu + \beta)} \frac{t_0^{\alpha\nu}}{(T^\alpha + \sigma)^\nu} \\ &\times F\left(2\nu + \beta, \nu; \nu + 1; -\frac{t_0^\alpha}{T^\alpha + \sigma}\right) \end{aligned} \quad (32)$$

where $F(a, b; c; z)$ is the hypergeometric function (see [14, p. 556]). The EB estimators $\hat{m}_n(T)$ and $\hat{R}_n(T, t_0)$ are obtained by replacing σ with $\hat{\sigma}$ in the expressions (31) and (32). By formula (11), the posterior pdf of τ given T is the inverted gamma

$$\begin{aligned} g(\tau|T, \hat{\sigma}) &= \frac{(T^\alpha + \hat{\sigma})^{\nu+\beta}}{\tau^{\nu+\beta+1} \Gamma(\nu + \beta)} \\ &\times \exp\left(-\frac{(T^\alpha + \hat{\sigma})}{\tau}\right) \end{aligned} \quad (33)$$

and the **confidence intervals** for τ are of the form (12).

Nonparametric EB

Now, let us consider nonparametric EB estimation of $m(\tau)$ and $R(\tau, t_0)$. In this subsection, following [17],

we carry out EB estimation under the assumption that θ belongs to some interval: $\Theta = [a_1, a_2]$, $0 < a_1 < a_2 < \infty$ where a_1 and a_2 are unknown. The problem reduces to estimation of $\hat{m}(T) = \Phi(T)/p(T)$ and $\hat{R}(T, t_0) = \Psi(T)/p(T)$ where $\Phi(T)$ and $\Psi(T)$ are of the form (6) with $u_1(\theta) = m(\theta)$ and $u_2(\theta) = R(\theta, t_0)$ given by formulae (25) and (26), respectively. Denote

$$Q(z, \varepsilon) = [\Gamma(1/\alpha)]^{-1} z^{1/\alpha-1} - (\varepsilon z)^{(1/\alpha-1)/2} J_{1/\alpha-1} \left(2\sqrt{z/\varepsilon} \right) \quad (34)$$

where $J_b(\cdot)$ is the Bessel function of integer order (see [14, p. 360]). Estimate $\Phi(T)$ and $\Psi(T)$, respectively, by

$$\begin{aligned} \hat{\Phi}_n(T) &= n^{-1} \sum_{j=1}^n \varphi_\varepsilon(T, T_j), \\ \hat{\Psi}_n(T, t_0) &= n^{-1} \sum_{j=1}^n \psi(T, T_j, t_0) \end{aligned} \quad (35)$$

with

$$\varphi_\varepsilon(T, T_j) = \begin{cases} \alpha T^{\alpha v-1} \frac{T_j^{\alpha-\alpha v} (T_j^\alpha - T^\alpha)^{1/\alpha-1}}{\Gamma(1/\alpha)} & \text{if } 0 < \alpha < 4/3 \\ T^{\alpha v-1} T_j^{\alpha-\alpha v} Q(T_j^\alpha - T^\alpha, \varepsilon), & \text{if } \alpha \geq 4/3 \end{cases} \quad (36)$$

and

$$\begin{aligned} \psi(T, T_j, t_0) &= \frac{\alpha T_j^{\alpha-\alpha v} T^{\alpha v-1} (T_j^\alpha - T^\alpha)^{2v-1}}{\Gamma^2(v)} \\ &\times B\left(\frac{t_0^\alpha}{T_j^\alpha - T^\alpha}; v, v\right) \\ &\times I(t_0^\alpha \leq T_j^\alpha - T^\alpha) \end{aligned} \quad (37)$$

Here $B(A; v, \mu)$ is the incomplete beta function (see [14, p. 263]), $\varepsilon = 2(\ln n \ln \ln n)^{-1}$. Estimate $p(T)$ by $\hat{p}_n(T)$ (see formula (15)) with

$$K_0(x) = (-x)^{-1/2} J_1 \left(2\sqrt{-x} \right) I(x < 0) \quad (38)$$

with $h = 2(\ln n \ln \ln n)^{-1}$. Finally, choose $H(a, \delta) = aI(0 < a\delta \leq 1) + \delta^{-1}I(a\delta > 1)$ and estimate $\hat{m}(T)$ and $\hat{R}(T, t_0)$ by $\hat{m}_n(T) = H(\hat{\Phi}_n(T)/\hat{p}_n(T), \delta)$ and $\hat{R}_n(T, t_0) = H(\hat{\Psi}_n(T)/\hat{p}_n(T), \delta)$, respectively.

EB Estimation of the Failure Intensity Function

Consider a repairable system and let $T = (t_1, \dots, t_m)$ be its failure times. The number of failures $N(t)$ during time interval $[0; t]$ is commonly described by the Poisson distribution

$$\begin{aligned} P(N(t) = k) &= [\Lambda(t)]^k \exp(-\Lambda(t)) / k!, \\ k &= 0, 1, 2, \dots \end{aligned} \quad (39)$$

Here, $\Lambda(t)$ is the average number of failures during time interval $[0; t]$, and $\lambda(t) = \Lambda'(t)$ is the *failure intensity function*. It is known (see, e.g., [18]) that in this situation the joint pdf of $t_1, \dots, t_m$ is of the form

$$f(t_1, \dots, t_m) = \left(\prod_{i=1}^m \lambda(t_i) \right) \exp\{-\Lambda(t_m)\} \quad (40)$$

The main concern of a practitioner is the behavior of the failure intensity $\lambda(t)$: whether it is increasing, decreasing, or staying constant. It is common in reliability to work with the *power law processes* with the intensity function of the form

$$\lambda(t) = \theta \beta t^{\beta-1}, \quad t > 0 \quad (41)$$

Here, β is the growth parameter, so that $\lambda(t)$ is increasing if $\beta > 1$, decreasing if $\beta < 1$, and equal to a constant failure rate $\lambda(t) \equiv \lambda$ if $\beta = 1$. The value of β can be estimated on the basis of the sample T but the estimator is not reliable when m is small.

Now, consider the situation when one observes failure times of n **repairable systems** $T_i = (T_{i1}, T_{i2}, \dots, T_{im_i})$, $i = 1, \dots, n$, where the i th system has an intensity function of the form equation (41) with parameters θ_i , and a common parameter β . Using EB approach reviewed above, one can use these past data for estimation of the intensity function $\lambda(t) = \tau \beta t^{\beta-1}$ corresponding to data $T = (t_1, \dots, t_m)$. Following [18], we assume that $\theta_1, \dots, \theta_n, \tau$ have the gamma prior pdf

$$\begin{aligned} p(\theta|a, b) &= [\Gamma(a)]^{-1} b^a \theta^{a-1} \exp(-b\theta), \\ \theta &> 0 \end{aligned} \quad (42)$$

where hyperparameters a and b are considered fixed but unknown. Hence, the marginal pdf of

$\mathbf{T} = (T_1, T_2, \dots, T_n)$ is of the form (see [18], p. 223)

$$p(T_1, \dots, T_n | a, b, \beta) = \beta^M \frac{b^{an}}{[\Gamma(a)]^n} \left(\prod_{i=1}^n \prod_{j=1}^{m_i} T_{ij} \right)^{\beta-1} \times \prod_{i=1}^n \frac{\Gamma(n_i + a)}{(T_{im_i}^\beta + b)^{m_i+a}} \quad (43)$$

where $M = \sum_{i=1}^n m_i$. Parameters a , b , and β can be found by the MLE procedure. In the general situation, MLEs of a , b , and β cannot be expressed explicitly; however, if for all n systems the values of T_{im_i} are identical and equal to s , then the MLE $\hat{a}$ of a is the solution of the equation

$$\sum_{i=1}^n \sum_{j=1}^{m_i} \frac{1}{\hat{a} + j - 1} = n \ln \left(1 + \frac{M}{n\hat{a}} \right) \quad (44)$$

while the MLEs $\hat{\beta}$ and $\hat{b}$ of β and b have, respectively, the following forms

$$\hat{\beta} = M \left(\sum_{i=1}^n \sum_{j=1}^{m_i} \ln(s/T_{ij}) \right)^{-1}, \quad \hat{b} = n\hat{a}s^{\hat{\beta}}/M \quad (45)$$

The posterior pdf of τ is then of the form

$$p(\tau | T) = \frac{(t_m + \beta)^{m+\hat{a}} \tau^{m+\hat{a}-1}}{\Gamma(m + \hat{a})} \exp(-\tau(\hat{b} + t_m^{\hat{\beta}})) \quad (46)$$

and can be used for point or interval estimation of τ .

Discussion

EB estimation of reliability was carried out by a number of researchers in the last 40 years. Martz and Waller ([2], Chapter 13) and Attia [19] offer a review of the techniques on the EB approach to reliability which complements this article.

Parametric EB estimation of the mean lifetime and the reliability function was considered in [15, 16, 20–23] among others. Dey and Kuo [15] investigated EB estimation of the mean lifetime for Weibull distribution with type II censored data, while Li [16] treated the case when $f(x|\theta)$ belongs to the exponential family. Both Dey and Kuo [15] and Li [16] employed the (inverted) gamma prior. Abu-Salih *et al.* [20] and Chiou [22] studied EB estimation of the reliability function in the case of the exponential

distribution, Ahsanullah and Ahmed [21] for a special class of distribution families, Jaheen [23] in the case of the Burr type X distribution, and Padgett [24] in the case of the lognormal model.

Various authors were involved in nonparametric estimation of reliability function and mean lifetime, see, for example, [17, 25–28] among others. For instance, Nakao and Liu [27], Papadopoulos [28], and Sarhan [29] estimated θ_i from previous data $T_i = (T_{i1}, T_{i2}, \dots, T_{im_i})$ and then constructed a nonparametric estimator $\hat{g}(\theta)$ of $g(\theta)$ using estimators $\hat{\theta}_i$ as “observations”. Lahiri and Park [25] and Liang and Padgett [26] carried out nonparametric EB estimation of the mean lifetime and the reliability function, respectively, based on **Dirichlet process** prior. The concept of nonparametric Bayes and EB estimators based on Dirichlet process prior is reviewed systematically in, for example, [30]; however, it is too mathematically involved to be treated in this small article. Finally, Pensky and Singh [17] used the approach based on estimating the numerator and the denominator of the Bayes estimator reviewed above. In addition, Liang and Singh [31] carried out EB testing for the value of the mean lifetime [32] and the reliability function [31].

EB estimation of the failure rate was studied in [33–36], and [18]. Gupta and Liang [33] used nonparametric EB approach for selecting the most reliable Poisson population, that is, the population with the smallest failure rate. Li [34] studied admissibility of EB rules for different types of priors. Mazzuchi and Soyer [35] estimated **software failure** rate when the life length of software on each stage of testing is described by exponential distribution and failure rates λ_i have the gamma pdfs. Finally, Rigdon and Basu provided extensive treatment of the EB estimation of the failure intensity function in their book [18].

Acknowledgment

This research was supported in part by the NSF grant DMS 0505133.

References

- [1] Berger, J.O. (1993). *Statistical Decision Theory and Bayesian Analysis*, 2nd Edition, Springer-Verlag, New York.
- [2] Martz, H.F. & Waller, R.A. (1982). *Bayesian Reliability Analysis*, John Wiley & Sons, Chichester.

- [3] Robbins, H. (1964). The empirical Bayes approach to statistical decision problems, *Annals of Mathematical Statistics* **35**, 1–20.
- [4] Carlin, B.P. & Louis, T.A. (1996). *Bayes and Empirical Bayes Methods for Data Analysis*, Chapman & Hall, London.
- [5] Maritz, J.S. & Lwin, T. (1989). *Empirical Bayes Methods*, 2nd Edition, Chapman & Hall, London.
- [6] Pensky, M. (1997). A general approach to nonparametric empirical Bayes estimation, *Statistics* **29**, 61–80.
- [7] Casella, G. (1985). An introduction to empirical Bayes data analysis, *The American Statistician* **39**, 83–87.
- [8] Morris, C.N. (1983). Parametric empirical Bayes inference: theory and applications, *Journal of the American Statistical Association* **78**, 47–65; With discussion.
- [9] Walter, G.G. (1981). Orthogonal series estimator of the prior distribution, *Sankhya* **A43**, 228–245.
- [10] Singh, R.S. (1979). Empirical Bayes estimation in Lebesgue-exponential families with rates near the best possible rate, *Annals of Statistics* **7**, 890–902.
- [11] Silverman, B.W. (1986). *Density Estimation for Statistics and Data Analysis*, Chapman & Hall, London.
- [12] Pensky, M. (1996). Empirical Bayes estimation of a scale parameter, *The Mathematical Methods of Statistics* **5**, 316–331.
- [13] Pensky, M. (1997). Empirical Bayes estimation of a location parameter, *Statistics and Decisions* **15**, 1–16.
- [14] Abramowitz, M. & Stegun, I.A. (1992). *Handbook of Mathematical Functions with Formulas, Graphs, and Mathematical Tables*, Dover Publications, New York. Reprint of the 1972 edition.
- [15] Dey, D.K. & Kuo, L. (1991). A new empirical Bayes estimator with type II censored data, *Computational Statistics and Data Analysis* **12**, 271–279.
- [16] Li, T.F. (1984). Empirical Bayes approach to reliability estimation for the exponential distribution, *IEEE Transactions on Reliability* **33**, 233–236.
- [17] Pensky, M. & Singh, R.S. (1999). Empirical Bayes estimation of reliability characteristics for an exponential family, *The Canadian Journal of Statistics* **27**, 127–136.
- [18] Rigdon, S.E. & Basu, A.P. (2000). *Statistical Methods for the Reliability of Repairable Systems*, John Wiley & Sons, New York.
- [19] Attia, A.F. (1999). Empirical Bayes approach to reliability, *The Egyptian Statistical Journal* **43**, R1–R22.
- [20] Abu-Salih, M.S., Yousef, M.A.Q. & Ali, M.A. (1988). On Bayes and empirical Bayes estimation of parameters and reliability function of two-parameter exponential distribution, *The Journal of Information and Optimization Sciences* **9**, 405–413.
- [21] Ahsanullah, M. & Ahmed, S.E. (2001). *Bayes and Empirical Bayes Estimates of Survival and Hazard Functions of a Class of Distributions*, *Lecture Notes in Statistics*, Vol. 148, Springer, New York, pp. 81–87.
- [22] Chiou, P. (1993). Empirical Bayes shrinkage estimation of reliability in the exponential distribution, *Communications in Statistics: Theory and Methods* **22**, 1483–1494.
- [23] Jaheen, Z.F. (1996). Empirical Bayes estimation of the reliability and failure rate functions of the Burr type X failure model, *Journal of Applied Statistical Science* **3**(4), 281–288.
- [24] Padgett, W.J. (1979). Some Bayes and Empirical Bayes estimators of reliability in the lognormal model, *Decision Information*, Academic Press, New York, London, Toronto, pp. 259–275.
- [25] Lahiri, P. & Park, D.H. (1991). Nonparametric Bayes and empirical Bayes estimators of mean residual life at age t , *Journal of Statistical Planning and Inference* **29**, 125–136.
- [26] Liang, K.Y. & Padgett, W.J. (1981). Nonparametric empirical Bayes estimation of reliability, *Metron* **39**, 143–152.
- [27] Nakao, Z. & Liu, Z.Z. (1990). Empirical Bayesian interval estimation for the reliability function of a bivariate exponential model, *Mathematicae Japonicae* **35**, 949–957.
- [28] Papadopoulos, A.S. (1983). Empirical Bayes confidence bounds for the Weibull distribution, *The Journal of Information and Optimization Sciences* **4**, 43–47.
- [29] Sarhan, A.M. (2003). Empirical Bayes estimates in exponential reliability model, *Applied Mathematics and Computation* **135**, 319–332.
- [30] Ghosh, J.K. & Ramamoorthi, R.V. (2003). *Bayesian Nonparametrics*, Springer-Verlag, New York.
- [31] Singh, R.S. (1998). Empirical Bayes procedures for testing the quality and reliability with respect to mean life, *Quality Improvement Through Statistical Methods* (Cochin, 1996), Statistics for Industry and Technology, Birkhäuser Boston, Boston, pp. 371–379.
- [32] Liang, T. (2005). On empirical Bayes testing for reliability, *Communications in Statistics: Theory and Methods* **34**, 661–670.
- [33] Gupta, S.S. & Liang, T. (2002). Selecting the most reliable Poisson population provided it is better than a control: a nonparametric empirical Bayes approach, *Journal of Statistical Planning and Inference* **103**, 191–203.
- [34] Li, T.F. (2002). Bayes empirical Bayes approach to estimation of the failure rate in exponential distribution, *Communications in Statistics: Theory and Methods* **31**, 1457–1465.
- [35] Mazzuchi, T.A. & Soyer, R. (1988). A Bayes empirical-Bayes model for software reliability, *IEEE Transactions on Reliability* **37**, 248–254.
- [36] Basu, A.P. & Rigdon, S.E. (1986). Examples of parametric empirical Bayes methods for the estimation of failure processes for repairable systems, *Reliability and Quality Control* (Columbia, 1984), North-Holland, Amsterdam, pp. 47–55.

MARIANNA PENSKEY

Maintenance and Markov Decision Models

Introduction

Maintenance modeling and Markov decision chains (MDCs) have a large joint history. For both areas the main research started in the 1950s. They formed a fruitful partnership, with theory allowing applications and applications guiding the need for theoretical developments. In this chapter we sketch the main ideas behind MDCs and indicate how they are applied in the maintenance area. In particular, we treat the area of maintenance of civil structures. Apart from providing ample references, we state the pros and cons of the Markov decision modeling.

Historical Perspective

Intense scientific interest in problems related to maintenance management originated only a few decades ago. The basis for the scientific support of maintenance is found in reliability engineering. The book *Mathematical Theory of Reliability* (Barlow and Proschan [1]) indicates the start of scientific interest in maintenance problems. In the early stages, maintenance was almost exclusively interpreted as preventive replacement of components (*see Replacement Strategies*). The development of some simple but insightful models describing aging phenomena of technical components, as well as the choice of appropriate actions to cope with these phenomena, showed that a scientific approach could be useful. One such model is the MDC. In the 1970s and early 1980s, the basic models were extended in several directions. In the 1980s, scientists made a step forward in bridging the gap between the analytical models and practice by a systematic analysis of systems. Moreover, similar to many other areas of applied (stochastic) optimization, the major advances in information technology brought the use of quantitative analytic models within computational reach.

During the last 15 years, substantial progress has been made in developing quantitative decision support systems for maintenance management of complex systems. Sophisticated decision support systems are in operation nowadays in the oil industry (both

for maintenance of refineries as well as offshore installations), the civil infrastructure sector (road, bridges, and railway), and electric power generation. The number of organizations that make use of some kind of quantitative tools to support inspection and replacement of equipment is increasing rapidly.

Maintenance

Maintenance can be defined as the combination of all technical and associated administrative actions intended to retain an item or system in, or restore it to, a state in which it can perform its required function. The maintenance objectives can be summarized under four headings – ensuring system function (availability, efficiency, and product quality); ensuring system life (asset management), ensuring safety, and ensuring human well-being.

For production equipment, ensuring the system function is often the prime objective. Here, maintenance has to provide the right (but not the maximum) reliability, availability, efficiency, and capability (i.e., producing at the right quality) of production systems in accordance with the need for these characteristics. In principle it is possible to give an economic value to the maintenance results, and a cost balance can be done. Ensuring system life and asset management refers to keeping systems as such in proper condition, while there are only indirect links to a possible production of goods or services. This objective is appropriate for civil structures like buildings, dams, offshore platforms, and roads, as their function is complex and not easy to measure. Often norms have to be set to define failure, and the benefits of maintenance are therefore more difficult to quantify. In this case, one has to minimize maintenance costs to meet the norms or conditions on states. Safety plays a role, as failures can have dramatic consequences, e.g., in the case of airplanes, and nuclear and chemical plants. In this case, testing and inspection activities constitute an important part of the maintenance work. Here, costs of maintenance have to be minimized while keeping the risks within strict limits and meeting statutory requirements. Finally, we refer to human well-being or shine as an objective; though there is no direct economic or technical necessity, it has primarily a psychological one (which indirectly may be economical). An example is painting, which is not done for protective reasons.

2 Maintenance and Markov Decision Models

Many models for maintenance consider replacing of parts or systems. Typical models are the age- and block-replacement models (*see Multivariate Age and Multivariate Renewal Replacement*). These models consider only one state of the system, *viz* working or failed. Next to that there are condition-based models that consider multiple states and these can be modeled by MDCs. For mechanical equipment, deterioration can be quick, and the economic importance of failures can be clear. In the case of civil infrastructures, deterioration can be very slow, and failure does not always have direct economic consequences. As a result, it is much more important to decide when to do maintenance for such systems and to allocate the budget to those assets with the highest needs. Markov decision models can be helpful in this respect.

Markov Decision Models

A Markov model is a special type of dynamic model with which the probabilistic evolution of a system can be modeled in time. The main assumption underlying such a model is that all information about the future behavior is captured in the state description. In other words, the present state provides all relevant information about the future behavior, and knowledge about previous states is not necessary. More formally, a Markov chain is a discrete-time process governed by a discrete state space E (observed at discrete time points) and transition matrix P , for which the Markov property holds, *i.e.*,

$$\begin{aligned} P_{ij} &= P(X_{t+1} = j | X_0 = i_0, X_1 = i_1, \dots, X_t = i) \\ &= P(X_{t+1} = j | X_t = i) \end{aligned} \quad (1)$$

If the transition probabilities do not depend on t , then the Markov chain is said to be stationary. An MDC is an extension of a Markov chain, where the Markov chain can be steered by actions and with which optimal actions can be determined. An alternative name is *Markov decision process* as it basically describes a stochastic process. The term *chain* is used to indicate that it gives rise to a chain of states in time. It is defined by a four-tuple E, A, P , and r – where E denotes the state space; $A(i)$ the action set in state $i \in E$; $P_{ij}(a)$, $i, j \in E$ the transition probabilities; and $r_i(a)$, $i \in E$, $a \in A(i)$, the immediate rewards in state i when action a is chosen.

The control is defined through policies and decision rules. A policy π is a sequence of decision rules $(\pi^1, \pi^2, \dots)$ where π^t is a function that assigns to each action a the probability of that action being taken at time t . For a memoryless or Markov policy, the decision rule at time t is independent of the states and actions before time t . A policy is said to be stationary and deterministic if all decision rules are identical and nonrandomized. For any decision rule π , we denote by $P(\pi)$, $r(\pi)$ the matrix and the vector with $P_{ij}(\pi(i))$ and $r_i(\pi(i))$ as i th element, respectively. Let $P^k(R) = P(\pi^k) \dots P(\pi^1)$ and $P^0(R) = I$, the identity matrix. The evolution of the MDC under a given stationary and deterministic policy f is a Markov chain with transition probability matrix $P(f)$. Let C be the class of all policies and C_M be the class of Markov policies. In policy optimization the issue is to find the best action a in each state such that some criterion is optimized. The expected total α -discounted rewards are defined by

$$v^\alpha(R) = \sum_{k=0}^{\infty} \alpha^k P^k(R) r(\pi^{k+1}) \quad (2)$$

Here α is the discount factor, which is taken from $[0, 1)$. For a stationary policy f , one can also obtain the discounted rewards $v^\alpha(f)$ by solving a set of equations, *i.e.* (in vector notation)

$$v^\alpha(f) = r(f) + \alpha P(f) v^\alpha(f) \quad (3)$$

Again if $\alpha < 1$, this set has a unique solution. Note that the size of this set of equations is equal to the size of the state space. Normally solving such a set takes a number of operations, which increases with the cube of the size of the state space. A policy R_α is called α -discounted optimal if for all $i \in E$, $R \in C$ we have

$$v_i^\alpha(R_\alpha) \geq v_i^\alpha(R) \quad (4)$$

A second criterion is the long-run expected average rewards, defined by

$$g_i(R) = \liminf_{N \rightarrow \infty} \frac{1}{N+1} \sum_{k=0}^N \sum_{j \in E} P_{ij}^k(R) r_j(\pi^{k+1}) \quad (5)$$

Let $g_i^* = \sup_{R \in C} g_i(R)$. In the finite-state and finite-action case, there exists an average optimal policy and the supremum can be replaced by a maximum. In the case of a denumerable state space, an

average optimal policy may exist only under certain conditions (like unichain and a certain recurrence). A policy R is long-run average optimal if $g_i(R) = g_i^*$, for all $i \in E$. For both criteria it can be shown that one can restrict oneself to the class of Markov policies. The first step in optimization is the derivation of the so-called optimality equations. For the α -discounted rewards they are defined as follows:

$$v_i^\alpha = \max_{a \in A(i)} \left\{ r_i(a) + \sum_{j \in E} \alpha P_{ij}(a) v_j^\alpha \right\}, \quad i \in E \quad (6)$$

It appears that there exists a unique solution to them and that any policy that takes maximizing actions in each state is α -discounted optimal. The optimality equations for the long-run average costs consist of two equations, viz

$$g_i = \max_{a \in A(i)} \left\{ \sum_{j \in E} P_{ij}(a) g_j \right\}, \quad i \in E \quad (7)$$

$$g_i + v_i = \max_{a \in A^0(i)} \left\{ r_i(a) + \sum_{j \in E} P_{ij}(a) v_j \right\}, \quad i \in E \quad (8)$$

$$\text{where } A^0(i) = \left\{ a \in A(i) \mid g_i = \sum_{j \in E} P_{ij}(a) g_j \right\}$$

Again a solution g, v to this set of equations is unique up to a constant vector for v and the solution g equals g^* . Any policy that takes maximizing actions in both equations in all states is average optimal. These equations can be simplified in case the MDC is *unichain* (which means that under all policies the Markov chain has one minimal closed set) or even weaker, if it is *communicating*, i.e., for every two pairs of states i and j , there is a policy f such that $P_{ij}^k(f) > 0$ for some $k > 0$. In this case the optimal average reward vector is a constant vector g and the first equation in the optimality equation is automatically satisfied. The next problem is to solve the equations.

Solution Methods

In principle there are three solution methods, i.e., policy improvement (PI), value iteration, and linear

programming (LP). PI works with a starting policy f_0 and determines in each step for each state an improving action. It appears that the stationary policy belonging to these actions also improves the objective function. As there are only a finite number of policies, the algorithm is finite. An algorithm for the α -discounted rewards is as follows:

- Step 0: let $k = 0$ and choose a starting policy f_0 .
- Step 1: policy evaluation: determine $v^\alpha(f_k)$ from equation (3).
- Step 2: PI: determine for each state $i : \max_{a \in A(i)} r_i(a) + \alpha \sum_{j \in E} P_{ij}(a) v_j^\alpha(f_k)$. Let $f_{k+1}(i)$ be the policy consisting of the maximizing action for each state i . Choose $f_{k+1}(i) = f_k(i)$ when possible.
- Step 3: if $f_{k+1} = f_k$, then stop, else increase k by 1 and return to step 1.

The successive approximation algorithm works with value iteration. It tries to find a solution to the optimality equations by repeatedly evaluating them. An algorithm is as follows:

- Step 1: choose a starting vector $v_i^0, i \in E$, e.g., by taking $v_i^0 = \max_{a \in A(i)} r_i(a)$ and set $k = 0$.
- Step 2: let $v_i^{k+1} = \max_{a \in A(i)} r_i(a) + \alpha \sum_{j \in E} P_{ij}(a) v_j^k, i \in E$.
- Step 3: check if $|v_i^{k+1} - v_i^k| < \varepsilon$ for all $i \in E$, if not go back to step 2.
- Step 4: determine the policy that takes maximizing actions in the last step.

The resulting policy has α -discounted rewards, which differ by less than $2\alpha(1 - \alpha)^{-1}\varepsilon$ from the maximal α -discounted rewards. The algorithm may be speeded up by eliminating nonoptimal actions. It is faster than PI if the transition matrix is sparse and only few transitions are possible. Finally the LP approach consists of formulating an LP for the optimality criterion. One can prove the following formulation

$$\begin{aligned} & \min \sum_{j \in E} \beta_j v_j, \\ & \text{s.t. } v_i \geq r_i(a) + \alpha \sum_{j \in E} P_{ij}(a) v_j, \quad a \in A(i), i \in E \end{aligned}$$

where $\beta_j, j \in E$, is just a vector of positive components. The solution equals the optimal discounted

rewards, and the policy consisting of actions that are binding in the LP restrictions is α -discount optimal. The advantage is that standard LP solvers can be used. From the LP formulation it follows that an MDC can be solved in polynomial time. More complex results are given in Papadimitriou and Tsitsiklis [2].

Issues in Applying MDCs

In general, there are two problems with applying the Markov decision model. The first is identifying states and establishing the Markov property and the second, which is related to the first, is the issue that state spaces can be very large, with the consequence that computation times can be prohibitive.

New techniques concentrate on finding a functional form for the value function. Remember that if one has an approximation for the value function, one can apply the PI step from the PI algorithm to obtain a better- or near-optimal policy. Several approaches have been suggested to this end. We would like to mention neurodynamic programming from Bertsekas and Tsitsiklis [3], reinforcement learning from Das *et al.* [4] and simulation-based approaches by Marbach and Tsitsiklis [5].

Extensions of the Markov Decision Chain Model

In a *semi-MDC* the time between observations of the system is not only fixed, but it is also a random variable, with a distribution that depends on the state and action chosen. Most parts of the analysis can be extended to semi-Markov chains. It is also possible to define a corresponding discrete-time Markov chain, by discretizing the transition time distribution, but this does increase the state space considerably.

In a *continuous-time Markov process*, the time between consecutive transitions is exponentially distributed. These **Markov processes** can be converted into discrete-time Markov chains. We refer to **Markov Processes** for an extensive treatment of Markov processes.

In a *partially observable MDC*, the problem is that not all states are observable and hence no specific actions can be taken once the system enters these states. A separate branch of theory is devoted to these types of Markov chains. A problem is that it

is no longer optimal to restrict oneself to Markov policies, because information on the past history also plays a role. This makes both analysis and solution methods inherently more difficult. Overviews on this type of models have been given by Monahan [6] and Lovejoy [7].

MDC Applications to Maintenance Problems

Maintenance problems have been among the first applications of MDCs, as the model allows both the modeling of the deterioration as well as the determination of the structure of optimal policies. Early examples were given by Sasieni [8] and Howard [9] in his car replacement problem. Let us describe the latter.

Consider a car, the age of which is expressed in periods of 3 months; so an age of j indicates an age of $3j$ months. In every period, the car could either age with another period or have a catastrophic failure/accident after which it is replaced with another car. Instead of letting the car age, we can also decide to replace it preventively with another car. In both cases, one needs to choose what kind of age the replacement car should have. The system can be modeled by taking the present age of the car as a state. Transitions to other states occur because of either the aging of the car or a catastrophic failure, in which case we replace the car. From a lifetime distribution, the transition probabilities can be derived. Next we need the cost figures for operating (e.g., due to some small maintenance), for failures, and for acquiring a car of a certain age j , for all values of j . We would like to remark that in the Markov chain modeling with expected rewards criteria, we do not need to bother whether costs were actually incurred before or after a transition. One just takes an expectation over all possible transitions in a state to define the expected immediate rewards. One of the results of MDC theory is that the optimal policy is of the control-limit type, i.e., replace the car if its age is beyond a critical age threshold. This can be proved by showing that the α -discounted costs increase in the states. Other approaches would not give such results.

The Markov decision approach has particularly been followed in condition-based maintenance, where different conditions of equipment are considered and maintenance actions are based on these. Next, there are many studies in maintenance that make use of

Markov models (*see Markov Processes*), but they use other methods to optimize rather than Markov decision theory.

Real applications have been very much in the sector of civil infrastructure, *viz* roads, bridges, and buildings. Nice overviews of applications have been given by Hontelez *et al.* [10] and Frangopol *et al.* [11]. Applications in other maintenance areas are less common, and as an example we would like to mention Amari *et al.* [12] for inspection planning of condition-based maintenance and Stengos and Thomas [13], who consider the replacement of blast furnaces in a steelworks. It appears that preventively replacing both the furnaces at the same time is not beneficial and that a specific cycle should be followed to reduce the probability that both furnaces fail together. In general, there are two different ways to model states. In the first, one takes the age of a structure as central and takes that as the basis to derive transition probabilities. An alternative is that one takes the remaining life span as a state. In the other modeling, one classifies the condition of the object according to some scale. Table 1 below shows a classification given by Scherer and Glagola [14] for bridge components.

Van Winden and Dekker [15] also use such a condition scale, albeit with five levels. These conditions are typically identified on inspection by classified personnel. One problem however with such scales is that the time to transition is not constant and that a semi-Markov modeling should be used. The latter can be converted by taking the residence time as an extra state variable, but it would make the state space two dimensional.

Another practical issue is that often the state of multiple components or different deterioration mechanisms has to be taken into account. Van Winden and Dekker [15] consider various elements in a building,

like window frames, masonry, pointing, and painting. Dekker *et al.* [16] consider for a road segment longitudinal and transversal unevenness, cracking, and raveling. All these aspects can be modeled but they make the state space multidimensional. A decomposition is not always possible because there can be actions which address multiple deterioration aspects or multiple components and which are cheaper than addressing these components separately.

The final issue is that often multiple objects need to be considered, because of either economies of scale in maintenance or budget restrictions that allow only some maintenance to be carried out. We now would like to discuss approaches for roads, bridges, and buildings.

A seminal paper on road maintenance was published by Golabi [17] and it discusses a pavement maintenance system for Arizona's highways. In the system the whole state highway system was modeled. LP was used as the technique; it has the advantage that restrictions on the fraction of road segments in a certain state can be taken into account. A problem is that no costs are involved with bad states; so one has to find a driver to stay out of these states. A special modeling advantage of roads is that they are more or less similar objects and that large sums of money are involved in their maintenance. The disadvantage is that one needs to consider road segments of 100 m and build up a quite large state space. The MDC approach is typically split 1 into several phases. First for a single segment, a Markov decision model is formulated and an optimal policy is derived. From this policy it follows what kind of budget is needed to keep the road in a good condition. Next one considers a whole network of road segments and determines priorities to work on the segments in case the budget is restricted. The advantage of this approach is that the state of all the highways can be

Table 1 Condition state descriptions^(a)

Condition state	Description
9	New condition
8	Good condition: no repairs needed
7	Generally good condition: potential exists for minor maintenance
6	Fair condition: potential exists for major maintenance
5	Fair condition: potential exists for minor rehabilitation
4	Marginal condition: potential exists for major rehabilitation
3	Poor condition: repair or rehabilitation required immediately

^(a)© American Society of Civil Engineers, 1994

taken into account, which gives a much fairer budget allocation. Many papers extended this approach, e.g., Chen *et al.* [18], Abaza and Ashur [19], Bako *et al.* [20], and Dekker *et al.* [16]. Later on, pavement management systems using Markov decision models came into use in many US states and other countries.

For bridges, the first paper was written by Scherer and Glagola [14], who have investigated the applicability and discussed the problems to keep the state space limited. The first real application, Pontis, was again made by a team led by Golabi [21]. Later on more applications followed (see Hawk and Small [22]). A problem with bridges, compared to roads, is that the bridges are more or less unique, which makes the modeling quite time consuming.

Building maintenance can be approached in a way similar to bridge maintenance. One trick with buildings is that one can work with generic elements and determine policies for these. Next an upscaling is made. For example, pointing is measured in square meters and the total costs can be determined by considering the total number of square meters per building multiplied by the cost per square meter. Sometimes a subdivision is necessary, but otherwise most of the same parts have a similar aging. This limits somewhat the number of components in a building, but considering a whole inventory of the building does give a large state space. One of the interesting results of MDC in this respect is that a trade-off between quality and costs can be obtained, which facilitates decision making, as shown in Figure 1, derived by Van Winden and Dekker [15]. Again, this is difficult to achieve with other techniques.

Finally we would like to mention quite a few papers on partially observable MDCs and maintenance. Again this is a natural combination as quite

often deterioration is not visible, but it can be modeled. Frangopol *et al.* [11] give a number of examples for bridge maintenance.

A natural extension of these is the application of MDCs in health maintenance. Here again the modeling is on the evolution of a disease within a human with several stages. The MDCs are used to determine the effectiveness of early screening for diseases and an example is Singer [23].

Conclusions

Markov decision chains have been a very successful tool in modeling maintenance. As such they are the main models in civil infrastructure maintenance. The main problem lies however in the size of the state space. But new techniques allow for much larger state spaces. So far no real alternative models have emerged which allow the determination of optimal policies. Hence, the main future issue is to develop new techniques to tackle the large multidimensional state spaces.

References

- [1] Barlow, R.E. & Proschan, F. (1965). *Mathematical Theory of Reliability*, John Wiley & Sons, New York.
- [2] Papadimitriou, C. & Tsitsiklis, J.N. (1987). The complexity of Markov decision processes, *Mathematics of Operations Research* **12**, 441–450.
- [3] Bertsekas, W. & Tsitsiklis, J.N. (1996). *Neuro-Dynamic Programming*, Athena Scientific, Nashua.
- [4] Das, T.K., Gosavi, A., Mahadevan, S. & Marchallick, N. (1999). Solving semi-Markov decision problems using average reward reinforcement learning, *Management Science* **45**(4), 560–574.
- [5] Marbach, P. & Tsitsiklis, J.N. (2001). Simulation-based optimization of Markov reward processes, *IEEE Transactions on Automatic Control* **46**(2), 191–209.
- [6] Monahan, G.E. (1982). A survey of partially observable Markov decision processes: theory, models and algorithms, *Management Science* **28**(1), 1–16.
- [7] Lovejoy, W.S. (1991). A survey of algorithmic methods for partially observed Markov decision processes, *Annals of Operations Research* **28**, 47–66.
- [8] Sasieni, M. (1956). A Markov chain process in industrial replacement, *Operational Research Quarterly* **7**, 148–155.
- [9] Howard, R.A. (1960). *Dynamic Programming and Markov Processes*, Technology Press of MIT.
- [10] Hontelez, J.A.M., Burger, H.H. & Wijnmalen, D.J.D. (1996). Optimum condition-based maintenance policies

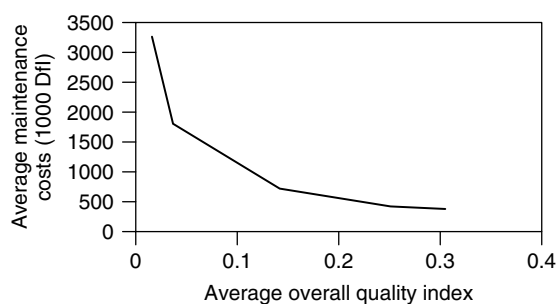


Figure 1 Trade-off between quality and costs

- for deteriorating systems with partial information, *Reliability Engineering and System Safety* **98**(1), 52–63.
- [11] Frangopol, D.M., Kallen, M.-J. & Van Noortwijk, J.M. (2004). Probabilistic models for life-cycle performance of deteriorating structures: review and future directions, *Progress in Structural Engineering and Materials* **6**, 197–212.
- [12] Amari, S.V., McLaughlin, L. & Pham, H. (2006). Cost-effective condition-based maintenance using Markov decision processes, *Proceedings Annual Reliability and Maintainability Symposium, RAMS 2006*, IEEE, 464–469.
- [13] Stengos, D. & Thomas, L. (1980). The blast furnaces problem, *European Journal of Operational Research* **4**, 330–336.
- [14] Scherer, W.T. & Glagola, D.M. (1994). Markovian models for bridge maintenance management, *Journal of Transportation Engineering* **120**(1), 37–51.
- [15] Van Winden, C. & Dekker, R. (1998). Markov decision models for building maintenance, *Journal of the Operational Research Society* **49**(9), 928–935.
- [16] Dekker, R., Plasmeijer, R. & Swart, J. (1998). Evaluation of a new maintenance concept for the preservation of highways, *IMA Journal of Mathematics Applied in Business and Industry* **9**, 109–156.
- [17] Golabi, K. (1982). A statewide pavement management system, *Interfaces* **6**, 5–21.
- [18] Chen, X., Hudson, S., Pajoh, M. & Dickinson, M. (1996). Development of new network optimization model for Oklahoma department of transportation, *Annual Meeting TRB*, Washington DC, No 1524, 103–108.
- [19] Abaza, K.A. & Ashur, S.A. (1999). *Optimum Decision Policy for Management of Pavement Maintenance and Rehabilitation*, *Transportation Research Record 1655*, Journal of Transportation Research Board, 8–15.
- [20] Bako, A., Klafszky, E., Szantai, T. & Gaspar, L. (1995). Optimization techniques for planning highway pavement improvements, *Annals of Operations Research* **58**(1), 55–66.
- [21] Golabi, K. & Shepard, R. (1997). Pontis: a system for maintenance optimization and improvement of US bridge networks, *Interfaces* **27**, 71–88.
- [22] Hawk, H. & Small, E.P. (1998). The BRIDGIT bridge management system, *Structural Engineering International* **8**(4), 309–314.
- [23] Singer, M. (2001). Cost effectiveness of screening for hepatitis C virus in asymptotic, average risk adults, *The American Journal of Medicine* **111**(8), 614–621.

Related Articles

Markov Processes; Replacement Strategies; Multivariate Age and Multivariate Renewal Replacement.

ROMMERT DEKKER, ROBIN P. NICOLAI, AND
LODEWIJK C.M. KALLENBERG

Prior Distribution Elicitation

Introduction

Expert engineering judgment plays a vital role in assessing and modeling the quality and reliability of systems [1, 2], although presenting subjective judgments as formal probability distributions is not without controversy [3]. Like any **data collection** process, a sound methodology for elicitation is a prerequisite for acquiring reliable subjective expert judgment, which will be summarized in the form of a prior distribution.

A prior distribution is a subjective probability distribution describing some unknown quantity of interest that has been elicited from a relevant person(s) prior to observing data. Using Bayes theorem, the prior distribution is updated in the light of data to produce a posterior distribution. Updation requires a conditional probability distribution, known as the *likelihood function*, which describes the probability of observing the data conditional on the state of the unknown quantity of interest.

A subjective probability distribution is one that describes a person's belief in the value of an unknown quantity. Winkler [4] believes that there is no single correct subjective probability distribution for an unknown quantity because each person forms his/her beliefs from personal experience. Cooke [5] points out that there is no way of measuring the correctness of a subjective distribution. In contrast, O'Hagan [6] distinguishes between true and stated probabilities. The former would result if an expert were capable of perfectly accurate assessments, while the latter would result from attempts to specify the underlying true probabilities. The authors agree that a good prior distribution will conform to the laws of probability, although the differing perspective has implications for validation.

The Bayes coherency principal states that uncertainties are described by probabilities and subjective probabilities should be such as to ensure self-consistent betting behavior [7]. This means that a person assesses subjective probabilities in such a way that they would avoid gambles involving certain losses (a so-called Dutch book). In principle, subjective probabilities can be measured

through preference behavior, when a subject declares indifference between the event under discussion and a lottery with a known probability (assuming of course that the "prizes" are comparable).

Even if a prior distribution conforms to the laws of probability, some elicited distributions provide more useful inference than others. Research in experimental psychology has demonstrated that accurate subjective probabilities are unobtainable by simply asking someone to provide a probability number; instead a structured elicitation process is required [8, 9]. These processes consist of an "expert" (or set of experts) and an "analyst". Following the definition of Ferrell [10], an expert is "a person with substantive knowledge about the events whose uncertainty is to be assessed", while an analyst should be a suitably qualified individual who will collect the data from the expert in order to formulate the prior distribution for use in modeling. The aim of an elicitation process is to minimize the impact of biases inherent in surfacing and capturing subjective expert judgment.

After examining aspects of modeling relevant to the elicitation of a prior distribution in the section titled "Model Parameterization and Uncertainties", the section titled "Heuristics and Biases" describes the typical biases encountered during elicitation. The section titled "Assessing Prior Distributions" further considers characteristics of good prior distributions, while the section titled "Improving Probability Assessors" highlights issues associated with improving the ability of assessors to provide useful inference. The section titled "Generic Processes" outlines a generic process for supporting the elicitation of a prior distribution and provides references to specific processes used in reliability. Finally, the section titled "Summary and Conclusions" reflects upon advanced theoretical and practical issues concerning the elicitation of prior distributions.

Model Parameterization and Uncertainties

It is useful to distinguish between different types of uncertainty in modeling: aleatory uncertainty, which corresponds to inherent randomness in the process generating observations; and epistemic uncertainty, which corresponds to state of knowledge of an expert. Within a Bayesian statistical model, the aleatory uncertainty is described by the likelihood

function, while the prior distribution describes the epistemic uncertainty. Typically, as more relevant data are acquired, the epistemic uncertainty is reduced although the aleatory uncertainty is not. As an example, consider the uncertainty associated with determining the length of time a component will operate without failing. The uncertainty associated with the mean time to failure of a component from a particular class is epistemic and will reduce through observing the times to failure for such components. However, even if the mean time to failure of components within this class is known, the aleatory uncertainty would still remain with respect to the actual time to failure for any one component.

Many of the favored model parameters, such as a failure rate and mean time to failure, are not actually observable quantities, but are simply parameters of a model used to make predictions about the future. It has been strongly argued on foundational grounds, for example the discussion in Chapter 2 of Bedford and Cooke [11], that an expert should only be asked for probability assessments on observable quantities. This raises an interesting issue concerning the degree to which the model construct relates to an expert's perception and hence the coverage of the elicitation. For example, Bedford *et al.* [2] propose that elicitation should include qualitative structuring of the model as well as the quantitative assessment of the prior distribution for the quantity of interest to facilitate appropriate parameterization.

Heuristics and Biases

Research in psychology is relevant to the elicitation of a prior distribution. The identification of the major sources of bias within expert assessments motivates the need for supportive elicitation processes and validation of data.

Sunstein [12] discusses two alternative cognitive systems to describe how a person makes decisions: System 1 uses intuition, and results in quick decisions, which can be subject to error; System 2 is more deliberative, calculative, and less likely to produce errors. System 1 is based on heuristics, whereby a person develops shortcut approaches to making decisions. While heuristics can provide insight into a problem, careful deliberation should result in an

improved assessment. This implies that the elicitation of a prior distribution requires an analyst to engage the expert in System-2 thinking. Hammond (see [13]) describes a person's thinking as a continuum between the aforementioned dichotomy.

Tversky and Kahneman [14] were the first authors to make substantial contributions to the identification of heuristics used within probability assessments and to discuss the implications for bias. Three key heuristics – availability, representative, anchoring – are typically used when assessing probabilities.

The *availability heuristic* refers to a person's judgment being influenced by how easily historical events are recalled. As an example, an expert may believe the likelihood of a particular system failing is high because there was a recent failure event. This can lead to a biased assessment as the expert is not considering a full failure history of the system.

The *representative heuristic* refers to the similarity between an outcome and a population. For example, on assessing the probability that a new component design will have a mean lifetime in excess of 1000 operating hours, one may consider how well the existing component population represents the design being evaluated. Care is required in using such a heuristic because an expert may assign an intuitive measure of similarity rather than considering how the probability will change with respect to that of the existing population.

The *anchoring and adjusting heuristic* refers to the situation when an expert provides an initial assessment and all future assessments are derived from an adjustment of the initial assessment. For example, consider the situation where a prior distribution for the mean time to failure for a component is required from an expert. An initial assessment may correspond to a best guess, from which the expert may anchor upon the initial assessment to adjust for upper and lower bounds from which the distribution is determined. Research shows that the adjustments made are typically insufficient [4, 15].

Assessing Prior Distributions

If an expert provides subjective probabilities that are coherent, then these are not incorrect insofar as they reflect the beliefs of the expert. However, it is clear that different experts give different assessments, and so this raises the question of how one might assess

the “quality” or “usefulness” of such distributions. As Cooke [5] pointed out, there is no absolute measure of the correctness of a subjective assessment. However two measures are often used either quantitatively or qualitatively: calibration and information.

To assess the degree to which an expert’s assessments are well calibrated requires the subjective probabilities to be compared with a collection of observed realizations. For example, two experts could assess the same event with one expert providing a probability of 0.9 and the other a probability of 0.2. This event will either be realized or not. Since neither expert assigned a probability of 1 or 0, regardless of the outcome neither will be incorrect. The first expert would be calibrated if 90% of all events assessed by that expert at the 90% level are realized and 10% are not. While for the second expert, 20% of all events assessed by that expert at the 20% level are realized and 80% are not. For this example, both experts could be perfectly calibrated but provide very different probabilities.

Direct assessment of calibration is only possible if we are able to make observations directly on the quantity of interest. In many cases, for example the uncertainty of a failure rate, this is not the case and so the distribution of any directly observable quantities (such as the number of units failing up to a given time) would be a predictive distribution obtained as a mixture of the sampling distribution weighted by the prior. Hence calibration measurements would naturally confound the sampling distribution and the prior.

A second consideration for the usefulness of an expert’s prior distribution is information, which can be assessed through **mean squared error** between the observed realizations and the subjective probabilities, information scores, or a similar such formula. Stronger conclusions can be drawn from an informative prior distribution than from an uninformative one. It is ideal to have an informative, well-calibrated expert. Note that it is possible to have an uninformative, calibrated expert. For example, if upon assessing a mean time to failure an expert provides the entire range of positive real numbers below the median 50% of the time and above the median 50% of the time, then the expert is calibrated but not informative.

Improving Probability Assessors

The quality of subjective probabilities from an expert

is dependent upon both the expert’s experience and the method of elicitation. If the expert lacks experience, the prior distribution will be uninformative or misleading, regardless of the elicitation approach employed. Poorly designed elicitation techniques may degrade the quality of the information provided by experts. Fischhoff [16] proposes four necessary conditions to support improved judgment skills:

1. abundant practice with a set of reasonably homogeneous tasks;
2. clear-cut criterion events for outcome feedback;
3. task-specific reinforcement;
4. explicit admission of the need for learning.

There is extensive evidence that these are often not achieved in practice [10, 17].

Standard techniques for eliciting prior distributions that reflect the uncertainty of the expert are reported in the literature. See, for example, Cooke [5] and Meyer and Booker [18]. The choice of technique depends upon the quantity of interest, the selection of which has been given little attention in the literature.

One important consideration when selecting the quantity of interest is feedback to the expert. This is considered crucial for calibrating the expert and should be event specific [10, 16, 17]. In other words, the feedback must be with respect to assigning probabilities to particular classes of relevant events and not only feedback on the ability of the expert to assign probabilities to any situation. To increase the effectiveness of feedback in terms of learning, conditions that influence the event should reoccur as often as possible [16, 19]. Therefore, the factors on which the measure is conditioned should be as few and as general as possible.

Generic Processes

A necessary first step toward developing a method for eliciting expert judgment is to start with a general approach that possesses the required key features for any elicitation process and to tailor it to the desired situation.

Phillips [20], Clemen and Reilly [21], and Garthwaite *et al.* [22], for example, all describe generic processes. Another generic process, first developed by the Stanford Research Institute (SRI)

4 Prior Distribution Elicitation

for eliciting expert judgment is discussed in Spetzler *et al.* [23], Ferrell [24], and Merkhofer [9]. The issues addressed within the SRI process are common considerations within the alternatives. There are seven stages to the SRI approach: motivate, structure, condition, encode, verify, aggregate, and discretize.

Motivating the expert can be achieved by explaining the elicitation process and how the results will be used. This phase can be used to enhance the comprehension of an expert regarding the process, encourage an expert to provide an accurate assessment of his/her understanding, and determine potential biases, such as motivational biases. These include expert bias where an expert becomes overconfident merely because he/she has been given the title “expert” and management bias where an expert provides goals as opposed to judgment. For example, an expert states the aspiration that there will be no faults within this system by time of manufacture, rather than an assessment of his beliefs.

Structuring involves defining the specific event being considered to ensure there is no ambiguity in the questions and to explore how an expert thinks about the quantity for which probability judgment is to be elicited. The aim is to manage cognitive bias by simplifying the complex task of assigning probabilities by disaggregating the quantity of interest into more elementary variables [25]. However, the unpacking principle [13] may be the consequence. This refers to the situation where the more detailed the description of the event, the greater the likelihood assigned to it. For example, an expert may provide an assessment for the probability of a component failing and subsequently during the elicitation process provide probabilities associated with causes of failure that may result in a system probability exceeding the initial assessment.

Sometimes experts maybe unclear about whether they are being asked for conditional events or conjoined events ($P(A|B)$ or $P(A \text{ and } B)$), as these can be difficult to distinguish in natural language. Judgments can be conditioned upon all relevant information necessary to ensure that an expert’s awareness has been stimulated and to manage biases such as anchoring and availability.

Encoding refers to an expert providing a numerical expression that reflects his/her uncertainty regarding an outcome. There are a number of techniques available for encoding subjective probabilities. Common approaches range from direct **estimation** whereby a

person is simply asked for the probability to inferring probabilities from elicited preferences with specified lotteries. A few devices have been used to support such elicitation, such as probability wheels or visual analog scales. An interesting discussion of alternative methods is given by O’Hagan *et al.* [13].

Encoding a probability distribution is more challenging because additional information is required than for an individual probability. One could partition the distribution into discrete intervals and elicit probabilities for the value of interest belonging to each of the intervals. Evidence exists to suggest that people are better at estimating the median and mode rather than the mean; moreover people are poor at assessing or interpreting the variance [13].

A popular alternative encoding procedure for distributions is the fractile method [5], in which the expert assesses the median value of his/her subjective probability distribution along with the (25th, 75th) and the (5th, 95th) **percentiles**. Once these values have been elicited, a parametric distribution can be sought to maximize fit. The main drawback of this technique is that the elicited data often reflect a central bias [26]. However, after percentiles of the prior distribution have been assessed, graphical techniques can be applied to enhance the quality of the prior distribution [27].

Verification is required to ensure that an expert has provided a reflection of his/her true beliefs. If problems are encountered then the previous stages are to be repeated.

If multiple experts are assessing the same quantity of interest, then their individual probability distributions should be aggregated. There are two approaches to aggregation – mathematical and behavioral. The former implies the experts should not influence each others’ decisions [24]. The latter requires experts to share their judgment and reassess their distributions and includes techniques such as Delphi [24] and nominal group technique [28]. There are several mathematical approaches to aggregation, most of which aim to evaluate a weighted average across the experts. See Cooke [5] for a fuller discussion.

It is often necessary to treat continuous random variables as discrete during the encoding stage. *Discretizing* refers to techniques for fitting continuous distributions to the elicited data while preserving important **moments**. This is accomplished by

dividing the range of all possible values for the uncertain variable into intervals, selecting a representative point from each interval, and assigning that point the probability that the actual value will fall within the corresponding interval. The moments can be preserved through Gaussian quadrature techniques [29].

Two elicitation processes specific to reliability include PREDICT (Kerscher [30, 31] and REMM [32–34]). Both provide modeling frameworks designed to estimate reliability throughout system development beginning with a problem structuring phase, which consists of eliciting a graphical representation of the relationship between the potential failure modes and the system reliability. The graphs form the basis for the elicitation.

Summary and Conclusions

An overview of the key issues associated with the elicitation of a prior distribution has been presented; in particular, the shortcomings due to the prevalence of heuristics employed by an expert when assessing probabilities have been highlighted. Evidence indicates that an expert can improve his/her assessments given feedback about meaningful quantities of interest and this learning can be accelerated under the right conditions. Moreover, good subjective probability assessment requires a supportive process for the experts, whereby key stages are sequenced for managing specific biases. For further details about these issues see O'Hagan *et al.* [13], Meyer and Booker [18], and Cooke [5].

Practically, the design and implementation of an elicitation process requires considered **project management** to ensure that, for example, the correct experts are selected, the choice of elicitation method is relevant to the model requirements and parameterization, the integrity of the probabilities assessed are maintained, and that sufficient resources are obtained to support an effective, and efficient, process. See Meyer and Booker [18] for further discussion.

Theoretically, there is a need to infer from the assessments which probability distributions on model parameters are consistent. This approach has been developed by Cooke [35] with more algorithms and underlying theory for probabilistic inversion in Kraan and Bedford [36]. Taking a more standard Bayesian perspective, Percy [37] discusses the

indirect assessment of parameters from generic prior distributions through the direct assessment of quantiles of the predictive distribution which will have observable outcomes. Gutierrez-Pulido *et al.* [38] take a similar line, considering both moments and quantiles of the time to failure for a system as sources of information from which prior distributions can be fitted. Such methods could also be applicable to other Bayesian contexts where prior distributions on lifetime distribution parameters are to be assessed – for example in **Bayesian accelerated** or proportional hazards life modeling [39].

Elicitation of a prior distribution places a heavy cognitive burden on the subjective probability assessments by an expert. Yet observed data can exist for earlier generations of a system, component or process, although it is not clear how these data can be adapted to inform the assessment of a new design. **Empirical Bayes** provides a framework for integrating relevant historical data and hence can reduce the cognitive burden on an expert. See for example Quigley *et al.* [40] and Bedford *et al.* [41].

References

- [1] Keeney, R.L. & von Winterfeldt, D. (1991). Eliciting probabilities from experts in complex technical problems, *IEEE Transactions on Engineering Management* **38**, 191–201.
- [2] Bedford, T., Quigley, J. & Walls, L. (2006). Expert elicitation for reliable system design, *Statistical Science* **21**(4), 428–450.
- [3] Evans, R.A. (1989). Bayes is for the Birds, *IEEE Transactions in Reliability* **38**, 401.
- [4] Winkler, R. (1967). The assessment of prior distribution in Bayesian analysis, *Journal of the American Statistical Association* **62**, 776–800.
- [5] Cooke, R.M. (1991). *Experts in Uncertainty*, Oxford University Press.
- [6] O'Hagan, A. (1988). *Probability: Methods and Measurement*, Chapman and Hall, London.
- [7] De Finette, B. (1974). *Theory of Probability*, John Wiley & Sons, New York.
- [8] Kahneman, D., Slovic, P. & Tversky, A. (1982). *Judgement Under Uncertainty: Heuristics and Biases*, Cambridge University Press.
- [9] Merkhofer, M. (1987). Quantifying judgemental uncertainty: methodology, experiences, and insights, *IEEE Transactions on Systems, Man, and Cybernetics* **17**, 741–752.
- [10] Ferrell, W. (1994). Discrete subjective probabilities and decision analysis: elicitation, calibration and combination, in *Subjective Probability*, G. Wright & P. Ayton, ed, John Wiley & Sons.

- [11] Bedford, T. & Cooke, R. (2001). *Probabilistic Risk Analysis: Foundations and Methods*, Cambridge University Press, Cambridge.
- [12] Sunstein, C.R. (2005). *Laws of Fear*, Cambridge University Press.
- [13] O'Hagan, A., Buck, C.E., Daneshkhah, A., Eiser, J.R., Garthwaite, P.H., Jenkinson, D.J., Oakley, J.E. & Rakow, T. (2006). *Uncertain Judgements Eliciting Experts' Probabilities*, John Wiley & Sons.
- [14] Tversky, A. & Kahneman, D. (1974). Judgement under uncertainty: heuristics and biases, *Science* **185**, 257–263.
- [15] Alpert, M. & Raiffa, H. (1982). A progress report on the training of probability assessors, in *Judgement Under Uncertainty: Heuristics and Biases*, D. Kahneman, P. Slovic & A. Tversky, eds, Cambridge University Press.
- [16] Fischhoff, B. (1989). Eliciting knowledge for analytical representation, *IEEE Transactions on Systems, Man, and Cybernetics* **19**, 448–461.
- [17] Wright, G. & Bolger, F., eds. (1992). *Expertise and Decision Support*, Plenum Press.
- [18] Meyer, M.A. & Booker, J.D. (2001). *Eliciting and Analysing Expert Judgement: A Practical Guide*, ASA-SIAM Series on Statistics and Applied Probability, 2nd Edition.
- [19] Kadane, J. & Wolfson, L. (1997). Experiences in elicitation, *The Statistician* **47**, 3–20.
- [20] Phillips, L.D. (1999). Group elicitation of probability distributions: are many heads better than one? in *Decision Science and Technology: Reflections on the Contributions of Ward Edwards*, J. Shanteau, B. Mellors & D. Schum, eds, Kluwer Academic Publishers.
- [21] Clemen, R.T. & Reilly, T. (2001). *Making Hard Decisions with Decision Tools*, Duxbury Press.
- [22] Gartwaite, P.H., Kadane, J.B. & O'Hagan, A. (2005). Statistical methods for eliciting probability distributions, *Journal of the American Statistical Association* **100**, 680–701.
- [23] Spetzler, C. & Stael von Holstein, C.A. (1975). Probability encoding in decision analysis, *TIMS* **22**, 340–358.
- [24] Ferrell, W. (1985). In *Combining Individual Judgements, Behaviour Decision Making*, G. Wright, ed, Plenum Press.
- [25] Armstrong, J.S., Denniston, W.B. & Gordon, M.M. (1975). The use of decomposition principle in making judgements, *Organizational Behavior and Human Performance* **14**, 257–263.
- [26] Seaver, D., von Winterfeldt, D. & Edwards, W. (1978). Eliciting subjective probability distributions on continuous variables, *Organizational Behavior and Human Performance* **21**, 379–391.
- [27] Chaloner, K., Church, T., Lois, T. & Mattis, J. (1993). Graphical elicitation of a prior distribution for a clinical trial, *The Statistician* **42**, 341–353.
- [28] Moore, C.M. (1987). *Group Techniques for Idea Building*, Sage.
- [29] Miller, A.C. & Rice, T.R. (1983). Discrete approximations of probability distributions, *Management Science* **29**, 352–362.
- [30] Kerscher, W.J.I., Booker, J.M., Bement, T.R. & Meyer, M.A. (1998). *Characterizing Reliability in a Product/Process Design-Assurance Program*. Annual Reliability and Maintainability Symposium 1998 Proceedings. International Symposium on Product Quality and Integrity (Cat. No.98CH36161), IEEE, Anaheim, pp. 105–112.
- [31] Kerscher, W., Booker, J. & Meyer, M. (2003). Predict: a case study, using fuzzy logic, in *Proceedings of the Reliability and Maintainability Symposium*, Tampa, FL, USA.
- [32] Walls, L., Quigley, J. & Marshall, J. (2006). Modeling to support reliability enhancement during product development with applications in the UK Aerospace Industry, *IEEE Transactions on Engineering Management* **53**(2), 263–274.
- [33] Walls, L. & Quigley, J. (2001). Building prior distributions to support Bayesian reliability growth modelling using expert judgement, *Reliability Engineering and System Safety* **74**, 117–128.
- [34] Hodge, R., Evans, M., Marshall, J., Quigley, J. & Walls, L. (2001). Eliciting engineering knowledge about reliability during design - lessons learnt from implementation, *Quality and Reliability Engineering International* **17**, 169–179.
- [35] Cooke, R.M. (1994). Parameter fitting for uncertain models: modelling uncertainty in small models, *Reliability Engineering and System Safety* **44**, 89–102.
- [36] Kraan, B. & Bedford, T. (2005). Probabilistic inversion of expert judgments in the quantification of model uncertainty, *Management Science* **51**(6), 995–1006.
- [37] Percy, D.F. (2002). Bayesian enhanced strategic decision making for reliability, *European Journal of Operational Research* **139**(1), 133–145.
- [38] Gutierrez-Pulido, H., Aguirre-Torres, V. & Christen, J.A. (2005). A practical method for obtaining prior distributions in reliability, *IEEE Transactions on Reliability* **54**(2), 262–269.
- [39] Bunea, C. & Mazzuchi, T.A. (2005). Bayesian accelerated life testing under competing failure modes. In *Annual Reliability and Maintainability Symposium, 2005 Proceedings. Proceedings Annual Reliability and Maintainability Symposium*, Alexandria, VA, USA, 152–157.
- [40] Quigley, J., Bedford, T. & Walls, L. (2007). Estimating rate of occurrence of rare events with empirical Bayes: a railway application, *Reliability Engineering and System Safety* **92**(5), 619–627.
- [41] Bedford, T., Quigley, J. & Walls, L. (2006). Fault tree inference through combining Bayes and empirical Bayes methods, in *Proceedings of the ESREL Conference*, Estoril, pp. 859–865.

Related Articles

Accelerated Life Tests: Bayesian Models; Bayesian

Robustness: Theory and Computation; Bayesian Reliability Analysis; Bayesian Reliability Demonstration; Empirical Bayes Estimation of Reliability;

Expert Opinion in Reliability; Masked Failure Data: Bayesian Modeling; Repairable Systems: Bayesian Analysis; Software Reliability: Bayesian Analysis.

JOHN QUIGLEY, TIM BEDFORD AND LESLEY
WALLS

Masked Failure Data: Bayesian Modeling

Introduction

Suppose we have a system with J components and a defect in any component causes the system to fail. We test many of these identical systems to assess each component's reliability. In the ideal situation, the exact cause of the system's failure can be identified. Then a detailed **Failure Mode Effect Analysis** can be carried out. However, one often encounters the situation where the exact component causing the system to fail is unknown, nevertheless, it can be narrowed to a subset of components. So we have masked data that is, the data consist of the lifetimes of n systems randomly chosen for life testing and their associated labels identifying the probable components that may cause each system to fail. Masking is usually due to limited resources for diagnosing the causes of failures. More importantly, it is due to the modular nature of present systems. For example, in computer repair shops, the cause of failure is often isolated to a module of several components. The system is brought back into operation by replacing the whole module. The objective in analyzing the masked system **failure data** is to make inferences for the reliability of the components. The **system reliability** is usually derived from the **component reliability** and the system's configuration. Another objective is to make inference for the diagnostic probability that is, the probability of a specific component causing the system to fail given the set of probable causes and the system's failure time.

Recently, several authors studied this problem for the series system using exponential distributions for the component lifetimes. Most authors make an equiprobable assumption for the masking probabilities, where the conditional masking probabilities given each cause of failure in the masked set take the same value. This assumption (called the symmetry assumption by some authors) allows one to consider only the reduced likelihood. The work includes Miyakawa [1], Usher and Hodgson [2], Guess *et al.* [3], Lin *et al.* [4, 5], Doganaksoy [6], Gastaldi [7], Reiser *et al.* [8], Usher [9], Mukhopadhyay and Basu [10], Basu *et al.* [11] and a review article by Sen *et al.* [12].

To relax the equiprobable assumption, Lin and Guess [13], and Guttman *et al.* [14] consider a proportional probability model for a dependent masking probability for a two-component series system. They assume that the masking probability conditional on the first component's failure and the system's failure time is proportional to that of the second component. Although both masking probabilities are functions of the system's failure time, they assume the proportionality is independent of the system's failure time. Craiu and Duchesne [15] propose EM algorithms for models that allow the masking probabilities to depend on time. Craiu and Reiser [16] propose EM based approach for inference for dependent **competing risks** models.

In a similar spirit to relax the equiprobable assumption, Kuo and Yang [17] explore other probability assumptions for the masking probabilities. They explore two different stochastic models for the masking probabilities. One model allows the masking probabilities to have different values depending on the true cause of failure, but they do not depend on the system's failure time. The other model assumes the masking probabilities are further dependent on some monotone functions of the system's failure time. In addition to exponential distributions for the components, they also consider Weibull distributions. To ensure identifiability, they assume the components have independent failure distributions. They show that better predictive power can be achieved by modeling the masking probabilities without the equiprobable assumption. They demonstrate this with a simulated data set and with a real data set based on Dinse [18].

Note that the system's failure occurs at the earliest onset of any component's failure. So we are using the competing risk methodology in our analysis. Furthermore, we are assuming the components are acting independently for identifiability issue. **Markov chain Monte Carlo** (MCMC) methods are typically used for the Bayesian inference for the problems posed here. In particular, an MCMC algorithm augmented with latent variables that simulate the true causes of failures will be used for Bayesian computation of the component reliability. In addition, an **EM algorithm** [19] for this problem will also be illustrated.

Model selection by a predictive approach as in Geisser and Eddy [20] can be developed naturally from the Bayesian framework. Kuo and Yang [17]

2 Masked Failure Data: Bayesian Modeling

provide more details on the model determination and show that improved model can be obtained by modeling the masking probabilities.

The sections titled “Full Likelihood” and “Reduced Likelihood” describe the basic models for the equiprobable assumption. The section titled “Relaxing the Equiprobable Assumption” discusses two variations of the basic models that model the masking probabilities. The MCMC algorithms with auxiliary latent variable augmentation are developed for these models. Some numerical results for a simulated data set are given in the section titled “Simulation Study” to illustrate the techniques. A brief discussion on the model determination is given in the section titled “Conclusion”.

Full Likelihood

In the full likelihood formulation, we not only consider the system lifetimes but also the labels that identify the possible failed components. Let t_{ij} denote the random lifetime of the j th component in the i th system. Assume we observe a sample of n **series systems** each having J components; therefore we observe $t_i = \min(t_{i1}, \dots, t_{iJ})$, for $i = 1, \dots, n$. In addition, for each system, we observe a set of labels of the components that include the failed component. Let S_i denote the set of components (labels) that possibly cause the system i to fail, and let s_i denote the realized set for S_i . The observed data denoted by $\mathbf{d}$ consist of $(t_1, s_1), \dots, (t_n, s_n)$. If each set s_i is a singleton, then this reduces to the usual competing risk study where the exact cause of failure is identified. The set of labels is often masked because the true cause of failure is hiding with other components. Let K_i denote the index of the component actually causing the i th system to fail. It is a singleton because of the series system and continuous lifetimes. Let f_j denote the density for the j th component's lifetime and $\bar{F}_j$ denote its survival function. We assume the component lifetime distributions are independent. For the i th system, the likelihood contribution coming from the datum (t_i, s_i) consists of $p(t_i, s_i)$. This can be written as $\sum_{j=1}^J p(t_i, K_i = j)p(s_i|t_i, K_i = j)$. Note the masking probability $p(s_i|t_i, K_i = j) = 0$, if $j \notin s_i$. Moreover, we have: $p(t_i, K_i = j) = f_j(t_i) \prod_{l=1, l \neq j}^J \bar{F}_l(t_i)$. Therefore the

full likelihood of this data set is

$$L_F = \prod_{i=1}^n \left[\sum_{j \in s_i} \left\{ \left(f_j(t_i) \prod_{l=1, l \neq j}^J \bar{F}_l(t_i) \right) \times P(S_i = s_i | T_i = t_i, K_i = j) \right\} \right] \quad (1)$$

It is often of interest to compute the diagnostic probability that is, the probability that the j th component actually causing the system to fail given the possible malfunctioning subset and the system's failure time. For the i th system and $j \in s_i$, the diagnostic probability for the j th cause is

$$\begin{aligned} p(k_i = j | s_i, t_i) &= \frac{p(k_i = j, s_i, t_i)}{p(s_i, t_i)} \\ &= \frac{P(S_i = s_i | T_i = t_i, K_i = j) f_j(t_i) \times \prod_{l=1, l \neq j}^J \bar{F}_l(t_i)}{\sum_{j \in s_i} P(S_i = s_i | T_i = t_i, K_i = j) \times f_j(t_i) \prod_{l=1, l \neq j}^J \bar{F}_l(t_i)} \end{aligned} \quad (2)$$

Reduced Likelihood

Several authors assume an equiprobable assumption where the masking probabilities take on the same value irrelevant of the particular cause of failure. That is, they make the following assumption for a fixed $j \in s_i$:

$$\begin{aligned} P(S_i = s_i | T_i = t_i, K_i = j') &= P(S_i = s_i | T_i = t_i, \\ &K_i = j) \text{ for all } j' \in s_i \end{aligned} \quad (3)$$

Then they proceed with their analysis with a reduced likelihood that does not depend on the masking probabilities:

$$L_R = \prod_{i=1}^n \left[\sum_{j \in s_i} \left\{ f_j(t_i) \prod_{l=1, l \neq j}^J \bar{F}_l(t_i) \right\} \right] \quad (4)$$

Reiser *et al.* [8] provide a Bayesian analysis with this reduced likelihood assuming exponential distributions $f_j(t_i) = \lambda_j e^{-\lambda_j t_i} I\{t_i > 0\}$, and the noninformative Jeffrey's prior:

$$\pi(\lambda_1, \dots, \lambda_J) \propto \prod_{j=1}^J \lambda_j^{-1} \quad (5)$$

For the system with $J = 2$, let n_1 , n_2 , and n_{12} denote the numbers of systems in the sample with true causes identified to be the first component, the second component, or masked (either component). Then we can summarize the observed data with $T = \sum_{i=1}^n t_i$ (**total time on test**), and n_1 , n_2 , and n_{12} . Then the posterior density of λ_1, λ_2 given the data is

$$\begin{aligned} \pi(\lambda_1, \lambda_2 | \mathbf{d}) &\propto \exp\{-T(\lambda_1 + \lambda_2)\} \lambda_1^{n_1-1} \lambda_2^{n_2-1} \\ &\times (\lambda_1 + \lambda_2)^{n_{12}} \end{aligned} \quad (6)$$

So they show the Bayes estimators (the posterior means) are given by

$$\tilde{\lambda}_1 = \frac{n_1}{T} \frac{n_1 + n_2 + n_{12}}{n_1 + n_2} \quad (7)$$

$$\tilde{\lambda}_2 = \frac{n_2}{T} \frac{n_1 + n_2 + n_{12}}{n_1 + n_2} \quad (8)$$

They are exactly the maximum likelihood estimators given by Usher and Hodgson [2]. An MCMC algorithm for computing the Bayes estimates can be developed easily that will be given in the subsection titled "Gibbs Sampling for the Reduced Model" as a special case of the algorithm. The MCMC algorithm allows us to compute **quantiles**, other functionals of the posterior distribution, and the highest posterior density credible sets.

Lin *et al.* [5] use a step-function prior to λ_j :

$$\pi(\lambda_j) = \sum_{k=1}^{\infty} \alpha_{j,k} I\{\lambda_j \in [a_{k-1}, a_k)\} \quad (9)$$

where I denotes the indicator function and $\sum_{k=1}^{\infty} \alpha_{j,k} (a_k - a_{k-1}) = 1$ satisfying the probability constraint. They provide a numerical example for the $J = 2$ system. An MCMC algorithm with data augmentation similar to Yang and Kuo [21] can be developed easily for this problem.

Relaxing the Equiprobable Assumption

Most of the Bayesian analysis is based on the equiprobable assumption for the masking probability. So reduced likelihood can be used. There are some treatments on relaxing the equiprobable assumption. For example, Kuo and Yang [17] model the masking probabilities from the simple classification point of view. One of the simple assumptions is to make the probabilities independent of the time and having different values conditional on their components. Let us look at the two-components system. Assume $P(S_i = \{1\} | t_i, K_i = 1) = \theta_1$, and $P(S_i = \{2\} | t_i, K_i = 2) = \theta_2$. Consequently, $P(S_i = \{1, 2\} | t_i, K_i = 1) = 1 - \theta_1$, and $P(S_i = \{1, 2\} | t_i, K_i = 2) = 1 - \theta_2$. Then the full likelihood in this case is

$$\begin{aligned} L_F &= \prod_{i|s_i=\{1\}} \theta_1 f_1(t_i) \bar{F}_2(t_i) \prod_{i|s_i=\{2\}} \theta_2 f_2(t_i) \bar{F}_1(t_i) \\ &\times \prod_{i|s_i=\{1,2\}} \{(1 - \theta_1) f_1(t_i) \bar{F}_2(t_i) \\ &+ (1 - \theta_2) f_2(t_i) \bar{F}_1(t_i)\} \end{aligned} \quad (10)$$

Let us note the cardinalities of the indices in each product in equation (10) are exactly n_1 , n_2 , and n_{12} , defined earlier in the section titled "Reduced Likelihood".

The same idea can be generalized to an arbitrary number of components. For example for $J = 3$, we define the following masking probabilities: $P(S_i = \{j\} | t_i, K_i = j) = \theta_j$ for all j , $P(S_i = \{1, 2\} | t_i, K_i = 1) = \theta_{12|1}$, and $P(S_i = \{1, 3\} | t_i, K_i = 1) = \theta_{13|1}$. Consequently, $P(S_i = \{1, 2, 3\} | t_i, K_i = 1) = 1 - \theta_1 - \theta_{12|1} - \theta_{13|1}$. Similarly for causes two and three. These statements can be simplified if we know more about the configuration of the system. For example, components one and two are mounted on the same card separately from component three. Then we can assume $\theta_{13|1} = \theta_{23|2} = \theta_{13|3} = \theta_{23|3} = 0$. There will be fewer parameters to be estimated. Aboul-Séoud and Usher [22] describe an example of the modular system structure.

If the replaced modules are further sent to the factories for identifying the exact component causing the failure, we can produce fairly accurate estimates of the θ 's from the follow-up study. Therefore, we can estimate f_j by assuming the θ 's are known. Alternatively, we can combine the original data

4 Masked Failure Data: Bayesian Modeling

$(t_i, s_i), i = 1, \dots, n$, with the follow-up “autopsy” data to obtain better inferences as done in Flehinger *et al.* [23, 24], and Reiser *et al.* [25].

In addition to exponential distributions for the component lifetimes, we will also consider Weibull distributions for F_j . Gibbs sampling will be developed for Bayesian inference for the θ ’s and the parameters in f_j . Although we concentrate on the two-components system here for illustration, the Gibbs sampler can be extended straightforwardly to any number of components. The method is akin to the Gibbs sampling used in many aggregated models (cf. [21]).

Gibbs Sampling for the Full Model

Let us define the hazard function $h_j(t) = f_j(t)/\bar{F}_j(t)$ for all j . Then equation (10) can be rewritten as

$$L_F = \left(\prod_{i|s_i=\{1\}} \theta_1 h_1(t_i) \prod_{i|s_i=\{2\}} \theta_2 h_2(t_i) \right. \\ \left. \times \prod_{i|s_i=\{1,2\}} \{(1 - \theta_1)h_1(t_i) + (1 - \theta_2)h_2(t_i)\} \right) \\ \times \left(\prod_{i=1}^n \bar{F}_1(t_i) \bar{F}_2(t_i) \right) \quad (11)$$

The summation in the third product of this likelihood prevents us from updating the parameters in F easily. Therefore, we introduce latent variables z_i for the systems where the cause of failure is masked. Let z_i be a Bernoulli variable with the success ($z_i = 1$) probability to be $p_i = (1 - \theta_1)h_1(t_i)/\{(1 - \theta_1)h_1(t_i) + (1 - \theta_2)h_2(t_i)\}$ for i with $s_i = \{1, 2\}$. It is interesting to note that this probability of success p_i is exactly the diagnostic probability $p(k_i = 1|s_i, t_i)$ in equation (2) for the true cause due to the first component conditional on $s_i = \{1, 2\}$ and t_i being observed. Let $z_+ = \sum_{i|s_i=\{1,2\}} z_i$. Then the conditional density of $\theta = (\theta_1, \theta_2)$ and the unknown parameters λ in F ’s given $\mathbf{d}$ and the latent variables $\mathbf{z}$ is

$$f(\theta, \lambda | \mathbf{z}, \mathbf{d}) \propto \theta_1^{n_1} (1 - \theta_1)^{z_+} \theta_2^{n_2} (1 - \theta_2)^{n_{12} - z_+} \\ \times \left(\prod_{i|s_i=\{1\}} h_1(t_i) \prod_{i|s_i=\{2\}} h_2(t_i) \right)$$

$$\times \prod_{i|s_i=\{1,2\}} h_1(t_i)^{z_i} h_2(t_i)^{1-z_i} \\ \times \left(\prod_{i=1}^n \bar{F}_1(t_i) \bar{F}_2(t_i) \right) \pi(\theta, \lambda) \quad (12)$$

where π denotes the prior density of θ and λ . Therefore, updating θ and λ can be done easily and independently for independent prior on θ and λ . We are assuming that readers are familiar with the basics of MCMC algorithms as described in Gelfand and Smith [26] and Tanner and Wong [27]. In the following, we only describe the transitional kernel in the MCMC algorithm. Now we write the Gibbs steps for two special cases, one for each component to have an exponential distribution and the other to have a Weibull distribution.

- (1) The exponential component $\bar{F}_j(t) = \exp(-\lambda_j t)$:

We consider a prior with the following independent components for $j = 1$ and 2 : $\theta_j \sim \beta e(\alpha_j, \beta_j)$ and $\lambda_j \sim G(a_j, b_j)$, where $G(a_j, b_j)$ denotes a gamma density with mean a_j/b_j . The conjugate priors for θ_j and λ_j in the augmented likelihood are chosen so that we can easily specify the conditional densities in the Gibbs algorithm. Gamma and beta prior densities are also versatile enough to incorporate various shapes for the distributions of the unknown parameters.

Step a (data augmentation step): given θ and λ , we generate $z_i \sim B(1, (1 - \theta_1)\lambda_1 / \{(1 - \theta_1)\lambda_1 + (1 - \theta_2)\lambda_2\})$ a Bernoulli distribution, independently for each i with $s_i = \{1, 2\}$. Then we compute $z_+ = \sum_{i|s_i=\{1,2\}} z_i$.

Step b (parameter updating step): given $\mathbf{z}$ and $\mathbf{d}$, we update the parameters by:

$$\theta_1 | \mathbf{z}, \mathbf{d}, \theta_2, \\ \lambda \sim \text{Beta}(\alpha_1 + n_1, \beta_1 + z_+) \quad (13)$$

$$\theta_2 | \mathbf{z}, \mathbf{d}, \theta_1, \\ \lambda \sim \text{Beta}(\alpha_2 + n_2, \beta_2 + n_{12} - z_+) \quad (14)$$

$$\lambda_1 | \mathbf{z}, \mathbf{d}, \lambda_2, \\ \theta \sim G\left(a_1 + n_1 + z_+, b_1 + \sum_{i=1}^n t_i\right) \quad (15)$$

$$\lambda_2|\mathbf{z}, \mathbf{d}, \lambda_1, \boldsymbol{\theta} \sim G \times \left(a_2 + n_2 + n_{12} - z_+, \quad b_2 + \sum_{i=1}^n t_i \right) \quad (16)$$

Note that the parameters $\theta_1, \theta_2, \lambda_1$ and λ_2 are conditionally independent given $\mathbf{z}$, and $\mathbf{d}$ from the above expressions. Therefore, updating the four parameters can be done in parallel.

(2) The Weibull component $\bar{F}_j(t) = \exp(-\lambda_j t^{\rho_j})$:

We consider a prior with the following independent components: $\theta_j \sim \text{Beta}(\alpha_j, \beta_j)$, $\lambda_j \sim G(a_j, b_j)$, and ρ_j has density $\pi(\rho_j)$ that can be arbitrary. The same ideas can be generalized to a system with more than two components. The assumption of independent components in the prior is chosen for convenience. We have chosen the gamma distribution for λ_j for conjugacy. We can choose the prior on ρ_j to be uniformly distributed over $(0, m_j)$ as suggested by Sinha [28, p.160]. We can also explore the choice $\pi(\rho_j, \lambda_j) \propto 1/(\rho_j \lambda_j)$ as given by Sinha and Zellner [29]. Let $\boldsymbol{\rho} = (\rho_1, \rho_2)$.

Step a (data augmentation step): given $\boldsymbol{\theta}, \boldsymbol{\lambda}$ and $\boldsymbol{\rho}$, we generate $z_i \sim B(1, (1 - \theta_1)\lambda_1\rho_1 t_i^{\rho_1-1} / \{(1 - \theta_1)\lambda_1\rho_1 t_i^{\rho_1-1} + (1 - \theta_2)\lambda_2\rho_2 t_i^{\rho_2-1}\})$, independently for i with $s_i = \{1, 2\}$. Then we compute $z_+ = \sum_{i|s_i=\{1,2\}} z_i$.

Step b (parameter updating step): given $\mathbf{z}$ and $\mathbf{d}$, we update the parameters by:

$$\theta_1|\mathbf{z}, \mathbf{d}, \boldsymbol{\lambda}, \boldsymbol{\rho}, \theta_2 \sim \text{Beta}(\alpha_1 + n_1, \beta_1 + z_+) \quad (17)$$

$$\theta_2|\mathbf{z}, \mathbf{d}, \boldsymbol{\lambda}, \boldsymbol{\rho}, \theta_1 \sim \text{Beta}(\alpha_2 + n_2, \beta_2 + n_{12} - z_+) \quad (18)$$

$$\lambda_1|\mathbf{z}, \mathbf{d}, \boldsymbol{\rho}, \boldsymbol{\theta}, \lambda_2 \sim G\left(a_1 + n_1 + z_+, \quad b_1 + \sum_{i=1}^n t_i^{\rho_1}\right) \quad (19)$$

$$\lambda_2|\mathbf{z}, \mathbf{d}, \boldsymbol{\rho}, \boldsymbol{\theta}, \lambda_1 \sim G\left(a_2 + n_2 + n_{12} - z_+, \quad b_2 + \sum_{i=1}^n t_i^{\rho_2}\right) \quad (20)$$

and update the parameters ρ_1 and ρ_2 independently given $\mathbf{z}, \mathbf{d}, \boldsymbol{\lambda}$ and $\boldsymbol{\theta}$, each by the Metropolis law (cf.

[30], and [31, 32]) with the following density as the target density:

$$\begin{aligned} \pi(\rho_1|\mathbf{z}, \mathbf{d}, \boldsymbol{\lambda}, \boldsymbol{\theta}) &\propto \rho_1^{n_1+z_+} \left(\prod_{i|s_i=\{1\}} t_i^{\rho_1-1} \right) \\ &\times \left(\prod_{i|s_i=\{1,2\}} t_i^{(\rho_1-1)z_i} \right) \exp \\ &\times \left(-\lambda_1 \sum_{i=1}^n t_i^{\rho_1} \right) \pi(\rho_1) \end{aligned} \quad (21)$$

$$\begin{aligned} \pi(\rho_2|\mathbf{z}, \mathbf{d}, \boldsymbol{\lambda}, \boldsymbol{\theta}) &\propto \rho_2^{n_2+n_{12}-z_+} \left(\prod_{i|s_i=\{2\}} t_i^{\rho_2-1} \right) \\ &\times \left(\prod_{i|s_i=\{1,2\}} t_i^{(\rho_2-1)(1-z_i)} \right) \\ &\times \exp(-\lambda_2 \sum_{i=1}^n t_i^{\rho_2}) \pi(\rho_2) \end{aligned} \quad (22)$$

Note in the above updating, the generation of $\boldsymbol{\theta}$ and $(\boldsymbol{\lambda}, \boldsymbol{\rho})$ can be done independently.

To generalize the above discussion to three components, we can easily extend the same prior for $\boldsymbol{\lambda}, \boldsymbol{\rho}$ in the component lifetimes to more components. For the masking probabilities $\boldsymbol{\theta}_1 = (\theta_1, \theta_{12|1}, \theta_{13|1}, \theta_{123|1})$ with $\theta_1 + \theta_{12|1} + \theta_{13|1} + \theta_{123|1} = 1$, a Dirichlet prior with parameters $(\alpha_1, \alpha_2, \alpha_3, \alpha_4)$ can be considered. Similarly, we can define independent Dirichlet priors for the other masking probabilities $\boldsymbol{\theta}_2$ and $\boldsymbol{\theta}_3$, where $j = 2$ or $j = 3$ is the true cause of failure. Then we can employ the same data augmentation idea as before and updating the parameters $\boldsymbol{\theta}_1, \boldsymbol{\theta}_2$, and $\boldsymbol{\theta}_3$ independently, each with an updated Dirichlet distribution.

EM Algorithm for the Full Model

If we are only interested in the posterior mode or the mode of the likelihood with their error estimates, we can apply the EM algorithm [19]. The algorithm may be simpler than the Gibbs sampler in implementation, it also may converge faster. The EM algorithm essentially converts a complicated optimization problem to many simpler iterative steps. The iteration steps are similar to that in the Gibbs algorithm, except the data

6 Masked Failure Data: Bayesian Modeling

augmentation step is replaced by an E-step (expectation) where the missing data (latent variables) are replaced by their expectations, and the parameter updating step is replaced by a maximization step. In the following, we make the same prior assumption as in the previous subsection to obtain the posterior mode. The algorithm can be easily modified to obtain the maximum likelihood estimate using the flat prior.

(1) For the exponential component, we make the following changes in the $l + 1st$ iteration:

Step a:

$$z_+^l = \frac{n_{12}(1 - \theta_1^l)\lambda_1^l}{(1 - \theta_1^l)\lambda_1^l + (1 - \theta_2^l)\lambda_2^l} \quad (23)$$

Step b:

$$\theta_1^{l+1} = \frac{\alpha_1 + n_1 - 1}{\alpha_1 + \beta_1 + n_1 + z_+^l - 2} \quad (24)$$

$$\theta_2^{l+1} = \frac{\alpha_2 + n_2 - 1}{\alpha_2 + \beta_2 + n_2 + n_{12} - z_+^l - 2} \quad (25)$$

$$\lambda_1^{l+1} = \frac{a_1 + n_1 + z_+^l - 1}{b_1 + \sum_{i=1}^n t_i} \quad (26)$$

$$\lambda_2^{l+1} = \frac{a_2 + n_2 + n_{12} - z_+^l - 1}{b_2 + \sum_{i=1}^n t_i} \quad (27)$$

(2) For the Weibull component:

Step a: the generation of z_i is replaced by setting $z_+^l = n_{12}(1 - \theta_1^l)\lambda_1^l \rho_1^l t_i^{\rho_1^l - 1} / \{(1 - \theta_1^l)\lambda_1^l \rho_1^l t_i^{\rho_1^l - 1} + (1 - \theta_2^l)\lambda_2^l \rho_2^l t_i^{\rho_2^l - 1}\}$.

Step b: the generation of θ_1 , θ_2 , λ_1 , and λ_2 is replaced by

$$\theta_1^{l+1} = \frac{\alpha_1 + n_1 - 1}{\alpha_1 + \beta_1 + n_1 + z_+^l - 2} \quad (28)$$

$$\theta_2^{l+1} = \frac{\alpha_2 + n_2 - 1}{\alpha_2 + \beta_2 + n_2 + n_{12} - z_+^l - 2} \quad (29)$$

$$\lambda_1^{l+1} = \frac{a_1 + n_1 + z_+^l - 1}{b_1 + \sum_{i=1}^n t_i^{\rho_1}} \quad (30)$$

$$\lambda_2^{l+1} = \frac{a_2 + n_2 + n_{12} - z_+^l - 1}{b_2 + \sum_{i=1}^n t_i^{\rho_2}} \quad (31)$$

and the Metropolis algorithm for generating ρ_1 and ρ_2 is replaced by finding ρ_1 that maximizes equation (21) and finding ρ_2 that maximizes equation (22).

Gibbs Sampling for the Reduced Model

Now we write the Gibbs sampler for the reduced model with the equiprobable assumption. Observe $1 - \theta_1 = 1 - \theta_2 = 1 - \theta$ in this case. We assume the same prior as the one in the full model except θ is $\beta e(\alpha_1, \beta_1)$ independent of other parameters. So the likelihood in this situation reduces to

$$L_R = \theta^{n_1+n_2} (1-\theta)^{n_{12}} \left(\prod_{i=1}^n \bar{F}_1(t_i) \bar{F}_2(t_i) \right) \prod_{i|s_i=\{1\}} h_1(t_i) \times \prod_{i|s_i=\{2\}} h_2(t_i) \prod_{i|s_i=\{1,2\}} \{h_1(t_i) + h_2(t_i)\} \quad (32)$$

So the Gibbs sampling for the exponential model is similar to the full model except z_i is generated by $B(1, \lambda_1/\{\lambda_1 + \lambda_2\})$, and instead of θ_1 and θ_2 , we generate a single $\theta \sim \beta e(\alpha_1 + n_1 + n_2, \beta_1 + n_{12})$. The Gibbs sampling for the Weibull model is the same as the full model except $\theta_1 = \theta_2 = \theta$ and we generate $\theta \sim \beta e(\alpha_1 + n_1 + n_2, \beta_1 + n_{12})$ to replace the generation of θ_1 and θ_2 .

Note the Gibbs sampler for the posterior distribution in equation (6) can be developed similarly. That is, we generate $z \sim \text{Bin}(n_{12}, \lambda_1/\{\lambda_1 + \lambda_2\})$ a binomial distribution, then generate $\lambda_1 \sim G(n_1 + z - 1, T)$, $\lambda_2 \sim G(n_2 + n_{12} - z - 1, T)$ independently.

EM Algorithm for the Reduced Model

The maximization for θ and other parameters can be done separately, where $\hat{\theta} = (\alpha_1 + n_1 + n_2 - 1)/(\alpha_1 + \beta_1 + n_1 + n_2 + n_{12} - 2)$ with no iteration. For other parameters we follow the same procedure as the EM algorithm for the full model in the subsection titled ‘‘EM Algorithm for the Full Model’’ with $\theta_1 = \theta_2$.

Simulation Study

We consider a series system with two exponential components, where $f_j(t) = \lambda_j \exp(-\lambda_j t)$, $j = 1, 2$, and $\lambda_1 = 0.01$ and $\lambda_2 = 0.005$. For each system, we

Table 1 Simulated data

Set of labels	System's failure time
$s_i = \{1\}$	3.4, 10.9, 71.5, 56.9, 41.6, 41.6, 38.6, 4.6, 8.5, 27.8, 2.9, 13.3, 37.4, 57.5, 168.1, 123.9, 136.3, 6.5, 21.4, 135.4, 93.1, 14.3, 10.8, 26.1, 95.3, 68.7, 22.3, 13.3, 73.4, 35.2, 28.0, 21.1
$s_i = \{2\}$	2.5, 212.2, 147.3, 63.6, 45.6, 4.2, 123.1, 10.5, 2.6, 27.4
$s_i = \{1, 2\}$	68.6, 2.3, 5.4, 8.8, 63.7, 325.2, 216.4, 132.0

simulate a pair of failure times from f_1 and f_2 respectively. Then, we take the minimum of the two failure times as the system's life time and note its label, the cause of the system's failure. If the cause is from the first (second) component, we randomly mask 10% (20%) of the observations. That is, the unmasked probabilities are $\theta_1 = 0.9$ and $\theta_2 = 0.8$. The data in Table 1 represent the failure time and cause of failure for a random sample of $n = 50$ systems.

We apply our analysis in the section titled "Relaxing the Equiprobable Assumption" to this simulated data set. We fit the data with the full model in equation (10) and the reduced model in equation (32) with exponentially distributed component lifetimes. Table 2 summarizes the priors, the posterior means, the EM estimates, and the posterior estimates of the diagnostic probability of the cause one for each of the full and reduced models. We choose uniform priors for the unmasked probabilities. We monitor the convergence of the Gibbs sampler using the Gelman and Rubin [33] method that uses the analysis of variance technique to determine whether or not further iterations are needed. We found 3 000 iterations

to be large enough for the priors being considered. All the following numerical results are obtained with 3 000 iterations and 50 replications in the Gibbs sampler. The table shows that the posterior means using the MCMC algorithm and the EM estimates are comparable. Starting at $\theta_1^0 = 0.5$, $\theta_2^0 = 0.5$, $\lambda_1^0 = 0.005$, $\lambda_2^0 = 0.005$, the EM algorithm converges after the 6th iteration when all the iteration error bounds are met, where each error bound for each parameter is set at the E^{-8} level. The diagnostic probability of the first cause conditional on $s_i = \{1, 2\}$ and t_i for such i is: $p(k_i = 1|t_i, s_i = \{1, 2\}) = (1 - \theta_1)\lambda_1 / \{(1 - \theta_1)\lambda_1 + (1 - \theta_2)\lambda_2\}$ for the full exponential model and $p(k_i = 1|t_i, s_i = \{1, 2\}) = \lambda_1 / (\lambda_1 + \lambda_2)$ for the reduced exponential model. Note these diagnostic probabilities are the same independent of t_i as long as i is such that $s_i = \{1, 2\}$. Table 2 also lists the posterior means for these two probabilities evaluated from the Gibbs sampler.

Conclusion

So far, we have presented several models for the masked system life times: full model and reduced model, each with exponential and Weibull component lifetimes. Which model is most desirable? We will use a predictive approach based on cross-validated data to address the model selection issue here. We can define the conditional predictive ordinate (CPO) for the i th data point for a model to be the predictive density evaluated at the i th data point given the remaining data. Then we define the pseudo-marginal predictive likelihood [20] to be the product of the CPOs over i . The pseudo-Bayes factor criterion is to select the model with the largest pseudo-marginal predictive likelihood. The computation of the CPOs

Table 2 Priors, posterior means, and Bayesian diagnostic probability estimates

Model	Full exponential	Reduced exponential
Prior	$\theta_i \sim \text{Beta}(1, 1)$ for $i = 1, 2$ $\lambda_i \sim G(1, 0.001)$ for $i = 1, 2$	$\theta \sim \text{Beta}(1, 1)$ $\lambda_i \sim G(1, 0.001)$ for $i = 1, 2$
Posterior means	$\hat{\theta}_1 = 0.866$, $\hat{\theta}_2 = 0.718$	$\hat{\theta} = 0.827$
Using Gibbs	$\hat{\lambda}_1 = 0.0125$, $\hat{\lambda}_2 = 0.0050$	$\hat{\lambda}_1 = 0.0131$, $\hat{\lambda}_2 = 0.0044$
Estimates	$\hat{\theta}_1 = 0.889$, $\hat{\theta}_2 = 0.714$	$\hat{\theta} = 0.840$
Using EM	$\hat{\lambda}_1 = 0.0121$, $\hat{\lambda}_2 = 0.00471$	$\hat{\lambda}_1 = 0.0128$, $\hat{\lambda}_2 = 0.00401$
$\hat{p}(k_i = 1 t_i, s_i = \{1, 2\})$ for each i with $s_i = \{1, 2\}$	0.53	0.75

follows straightforwardly from the Gibbs samplers. Details of this model determination are given in Kuo and Yang [17].

References

- [1] Miyakawa, M. (1984). Analysis of incomplete data in a competing risks model, *IEEE Transaction on Reliability* **33**, 293–296.
- [2] Usher, J. & Hodgson, T. (1988). Maximum likelihood analysis of component reliability using masked system life-test data, *IEEE Transaction on Reliability* **37**, 550–555.
- [3] Guess, F.M., Usher, J.S. & Hodgson, T.J. (1991). Estimating system and component reliabilities under partial information on cause of failure, *Journal of Statistical Planning and Inference* **29**, 75–85.
- [4] Lin, D.K.J., Usher, J.S. & Guess, F.M. (1993). Exact maximum likelihood estimation using masked system data, *IEEE Transaction on Reliability* **42**, 631–635.
- [5] Lin, D.K.J., Usher, J.S. & Guess, F.M. (1996). Bayes estimation of component-reliability from masked system-life data, *IEEE Transaction on Reliability* **45**, 233–237.
- [6] Doganaksoy, N. (1991). Interval estimation from censored & masked system-failure data, *IEEE Transaction on Reliability* **40**, 280–286.
- [7] Gastaldi, T. (1994). Improved maximum likelihood estimation for component reliabilities with Miyakawa-Usher-Hodgson-Guess' estimators under censored search for the cause of failure, *Statistics & Probability Letters* **19**, 5–18.
- [8] Reiser, B., Guttman, I., Lin, D.K.J., Guess, F.M. & Usher, J.S. (1995). Bayesian inference for masked system lifetime data, *Applied Statistics* **44**, 79–90.
- [9] Usher, J. (1996). Weibull component reliability-prediction in the presence of masked data, *IEEE Transaction on Reliability* **45**, 229–232.
- [10] Mukhopadhyay, C. & Basu, A.P. (1997). Bayesian analysis of incomplete time and cause of failure data, *Journal of Statistical Planning and Inference* **59**, 79–100.
- [11] Basu, S., Sen, A. & Banerjee, M. (2003). Bayesian analysis of competing risks with partially masked cause of failure, *Applied Statistics* **52**, 77–93.
- [12] Sen, A., Basu, S. & Banerjee, M. (2001). Analysis of masked failure data under competing risks, in *Handbook of Statistics*, N. Balakrishnan & C.R. Rao, eds, Elsevier Science, Amsterdam, Vol. 20, pp. 523–540.
- [13] Lin, D.K.J. & Guess, F.M. (1994). System life data analysis with dependent partial knowledge on the exact cause of system failure, *Microelectronics Reliability* **34**, 535–544.
- [14] Guttman, I., Lin, D.K.J., Reiser, B. & Usher, J.S. (1995). Dependent masking and system life data analysis: Bayesian inference for two-component systems, *Lifetime Data Analysis* **1**, 87–100.
- [15] Craiu, R.V. & Duchesne, T. (2004). Inference based on the EM algorithm for the competing risks model with masked causes of failure, *Biometrika* **91**, 543–558.
- [16] Craiu, R.V. & Reiser, B. (2006). Inference for the dependent competing risks model with masked causes of failure, *Lifetime Data Analysis* **12**, 21–33.
- [17] Kuo, L. & Yang, T. (2000). Bayesian reliability modeling for masked system lifetime data, *Statistics and Probability Letters* **47**, 229–241.
- [18] Dinse, G. (1986). Nonparametric prevalence and mortality estimators for animal experiments with incomplete cause-of-death data, *Journal of the American Statistical Association* **81**, 328–336.
- [19] Dempster, A., Laird, N. & Rubin, D. (1977). Maximum likelihood from incomplete data via the EM algorithm, *Journal of Royal Statistical Society. Series B* **39**, 1–38.
- [20] Geisser, S. & Eddy, W. (1979). A predictive approach to model selection, *Journal of the American Statistical Association* **74**, 153–160.
- [21] Yang, T. & Kuo, L. (1999). Bayesian computation for the superposition of nonhomogeneous Poisson processes, *Canadian Journal of Statistics* **27**, 547–556.
- [22] Aboul-Séoud, M.H. & Usher, J.S. (1996). The effect of modular system structures on component reliability estimation, in *5th Industrial Engineering Research Conference Proceedings*, Minneapolis, MN, pp. 464–467.
- [23] Flehinger, B.J., Reiser, B. & Yashchin, E. (1996). Inference about defects in the presence of masking, *Technometrics* **38**, 247–255.
- [24] Flehinger, B.J., Reiser, B. & Yashchin, E. (1998). Survival with competing risks and masked causes of failures, *Biometrika* **85**, 151–164.
- [25] Reiser, B., Flehinger, B.J. & Conn, A.R. (1996). Estimating component-defect probability from masked system success/failure data, *IEEE Transaction on Reliability* **45**, 238–243.
- [26] Gelfand, A. & Smith, A.F.M. (1990). Sampling based approaches to calculating marginal densities, *Journal of the American Statistical Association* **85**, 398–409.
- [27] Tanner, M.A. & Wong, W.H. (1987). The calculation of posterior distributions by data augmentation, *Journal of the American Statistical Association* **82**, 528–550.
- [28] Sinha, S. (1986). *Reliability and Life Testing*, John Wiley & Sons, New York.
- [29] Sinha, B. & Zellner, A. (1990). A note on the prior distributions of Weibull parameters, *SCIMA* **19**, 5–13.
- [30] Metropolis, N., Rosenbluth, A.W., Rosenbluth, M.N., Teller, A.H. & Teller, E. (1953). Equations of state calculations by fast computing machines, *Journal of Chemical Physics* **21**, 1087–1091.
- [31] Kuo, L. & Yang, T. (1995). Bayesian computation for software reliability, *Journal of Computational and Graphical Statistics* **4**, 1–18.
- [32] Kuo, L. & Yang, T. (1996). Bayesian computation for nonhomogeneous Poisson processes in software reliability, *Journal of the American Statistical Association* **91**, 763–773.

- [33] Gelamn, A. & Rubin, D.B. (1992). Inference from iterative simulation using multiple sequences (with discussion), *Statistical Science* **7**, 457–511.

Related Articles

Masked Failure Data: Competing Risks; Masked Failure Data.

LYNN KUO

Bayesian Networks

Introduction

Bayesian networks (BNs), also known as *belief networks* (or Bayes nets for short), belong to the family of probabilistic *graphical models* (GMs). These graphical structures are used to represent knowledge about an uncertain domain. In particular, each node in the graph represents a random variable, while the edges between the nodes represent probabilistic dependencies among the corresponding random variables. These conditional dependencies in the graph are often estimated by using known statistical and computational methods. Hence, BNs combine principles from graph theory, **probability theory**, computer science, and statistics.

GMs with *undirected edges* are generally called *Markov random fields* or *Markov networks*. These networks provide a simple definition of independence between any two distinct nodes based on the concept of a *Markov blanket*. Markov networks are popular in fields such as statistical physics and computer vision [1, 2].

BNs correspond to another GM structure known as a *directed acyclic graph* (DAG) that is popular in the statistics, the machine learning, and the artificial intelligence societies. BNs are both mathematically rigorous and intuitively understandable. They enable an effective representation and computation of the joint probability distribution (JPD) over a set of random variables [3].

The structure of a DAG is defined by two sets: the set of nodes (vertices) and the set of directed edges. The nodes represent random variables and are drawn as circles labeled by the variable names. The edges represent direct dependence among the variables and are drawn by arrows between nodes. In particular, an edge from node X_i to node X_j represents a statistical dependence between the corresponding variables. Thus, the arrow indicates that a value taken by variable X_j depends on the value taken by variable X_i , or roughly speaking that variable X_i “influences” X_j . Node X_i is then referred to as a *parent* of X_j and, similarly, X_j is referred to as the *child* of X_i . An extension of these genealogical terms is often used to define the sets of “descendants” – the set of nodes that can be reached on a direct path from the node, or “ancestor” nodes – the set

of nodes from which the node can be reached on a direct path [4]. The structure of the acyclic graph guarantees that there is no node that can be its own ancestor or its own descendent. Such a condition is of vital importance to the factorization of the joint probability of a collection of nodes as seen below. Note that although the arrows represent direct causal connection between the variables, the *reasoning process* can operate on BNs by propagating information in any direction [5].

A BN reflects a simple conditional independence statement. Namely that each variable is independent of its nondescendants in the graph given the state of its parents. This property is used to reduce, sometimes significantly, the number of parameters that are required to characterize the JPD of the variables. This reduction provides an efficient way to compute the posterior probabilities given the evidence [3, 6, 7].

In addition to the DAG structure, which is often considered as the “qualitative” part of the model, one needs to specify the “quantitative” parameters of the model. The parameters are described in a manner which is consistent with a Markovian property, where the conditional probability distribution (CPD) at each node depends only on its parents. For discrete random variables, this conditional probability is often represented by a table, listing the local probability that a child node takes on each of the feasible values – for each combination of values of its parents. The joint distribution of a collection of variables can be determined uniquely by these local conditional probability tables (CPTs).

Following the above discussion, a more formal definition of a BN can be given [7]. A Bayesian network B is an annotated acyclic graph that represents a JPD over a set of random variables $\mathbf{V}$. The network is defined by a pair $B = \langle G, \Theta \rangle$, where G is the DAG whose nodes $X_1, X_2, \dots, X_n$ represents random variables, and whose edges represent the direct dependencies between these variables. The graph G encodes independence assumptions, by which each variable X_i is independent of its nondescendants given its parents in G . The second component Θ denotes the set of parameters of the network. This set contains the parameter $\theta_{x_i|\pi_i} = P_B(x_i|\pi_i)$ for each realization x_i of X_i conditioned on π_i , the set of parents of X_i in G . Accordingly, B defines a unique JPD over $\mathbf{V}$, namely:

2 Bayesian Networks

$$P_B(X_1, X_2, \dots, X_n) = \prod_{i=1}^n P_B(X_i | \pi_i) = \prod_{i=1}^n \theta_{X_i | \pi_i} \quad (1)$$

For simplicity of representation we omit the subscript B henceforth. If X_i has no parents, its local probability distribution is said to be *unconditional*, otherwise it is *conditional*. If the variable represented by a node is *observed*, then the node is said to be an evidence node, otherwise the node is said to be hidden or latent.

Consider the following example that illustrates some of the characteristics of BNs. The example shown in Figure 1 has a similar structure to the classical “earthquake” example in Pearl [3]. It considers a person who might suffer from a back injury, an event represented by the variable *Back* (denoted by B). Such an injury can cause a backache, an event represented by the variable *Ache* (denoted by A). The back injury might result from a wrong sport activity, represented by the variable *Sport* (denoted by S) or from new uncomfortable chairs installed at the person’s office, represented by the variable *Chair* (denoted by C). In the latter case, it is reasonable to assume that a coworker will suffer and report a similar backache syndrome, an event represented by the variable *Worker* (denoted by W). All variables are binary; thus, they are either true (denoted by “T”) or false (denoted by “F”).

The CPT of each node is listed besides the node.

In this example the parents of the variable *Back* are the nodes *Chair* and *Sport*. The child of *Back* is *Ache*, and the parent of *Worker* is *Chair*. Following the BN independence assumption, several independence statements can be observed in this case. For example, the variables *Chair* and *Sport* are marginally independent, but when *Back* is given they are conditionally dependent. This relation is often called *explaining away*. When *Chair* is given, *Worker* and *Back* are conditionally independent. When *Back* is given, *Ache* is conditionally independent of its ancestors *Chair* and *Sport*. The conditional independence statement of the BN provides a compact factorization of the JPDs. Instead of factorizing the joint distribution of all the variables by the chain rule, i.e., $P(C, S, W, B, A) = P(C)P(S|C)P(W|S, C)P(B|W, S, C)P(A|B, W, S, C)$, the BN defines a unique JPD in a factored form, i.e. $P(C, S, W, B, A) = P(C)P(S)P(W|C)P(B|S, C)P(A|B)$. Note that the BN form reduces the number of the model parameters, which belong to a multinomial distribution in this case, from $2^5 - 1 = 31$ to 10 parameters. Such a reduction provides great benefits from inference, learning (parameter **estimation**), and computational perspective. The resulting model is more robust with respect to **bias**-variance effects [8]. A practical graphical criterion that helps to

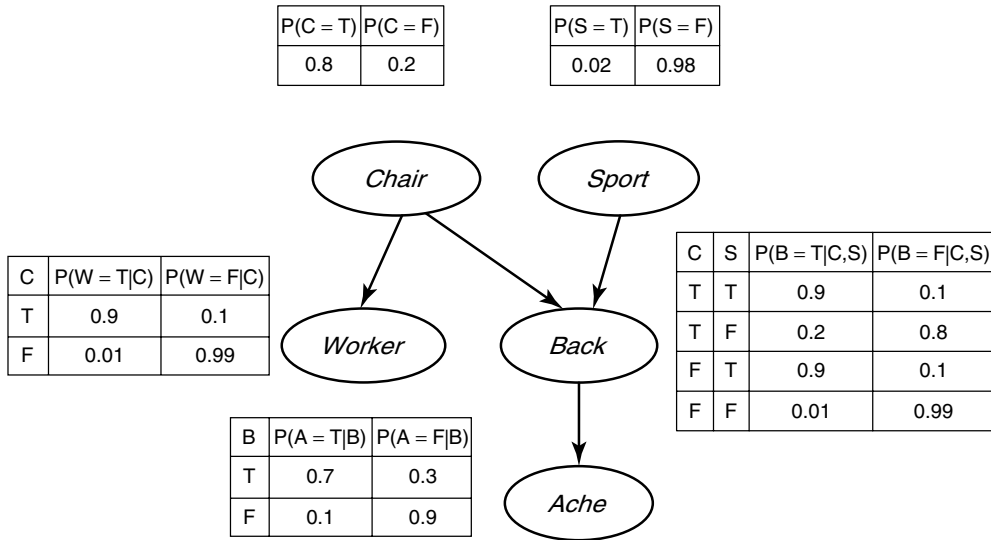


Figure 1 The backache BN example

investigate the structure of the JPD modeled by a BN is called *d-separation* [3, 9]. It captures both the conditional independence and dependence relations that are implied by the Markov condition on the random variables [2].

Inference via BN

Given a BN that specified the JPD in a factored form, one can evaluate all possible inference queries by marginalization, i.e. summing out over “irrelevant” variables. Two types of inference support are often considered: *predictive support* for node X_i , based on evidence nodes connected to X_i through its parent nodes (also called *top-down reasoning*), and *diagnostic support* for node X_i , based on evidence nodes connected to X_i through its children nodes (also called *bottom-up reasoning*). Given the example in Figure 1, one might consider the diagnostic support for the belief on new uncomfortable chairs installed at the person’s office, given the observation that the person suffers from a backache. Such a support is formulated as follows:

$$P(C = T|A = T) = \frac{P(C = T, A = T)}{P(A = T)} \quad (2)$$

where

$$\begin{aligned} P(C = T, A = T) = & \sum_{S, W, B \in \{T, F\}} P(C = T)P(S) \\ & \times P(W|C = T)P(B|S, C = T)P(A = T|B) \end{aligned} \quad (3)$$

and

$$\begin{aligned} P(A = T) = & \sum_{S, W, B, C \in \{T, F\}} P(C)P(S)P(W|C)P(B|S, C) \\ & \times P(A = T|B) \end{aligned} \quad (4)$$

Note that even for the binary case, the JPD has size $O(2^n)$, where n is the number of nodes. Hence, summing over the JPD takes exponential time. In general, the full summation (or integration) over discrete (continuous) variables is called *exact inference* and known to be an *NP-hard problem*. Some efficient algorithms exist to solve the exact inference problem in restricted classes of networks. One of the most popular algorithms is the *message passing algorithm* that solves the problem in $O(n)$ steps (linear in the number

of nodes) for *polytrees* (also called *singly connected networks*), where there is at most one path between any two nodes [3, 5]. The algorithm was extended to general networks by Lauritzen and Spiegelhalter [10]. Other exact inference methods include the *cycle-cutset* conditioning [3] and variable elimination [11].

Approximate inference methods were also proposed in the literature, such as *Monte Carlo* sampling that gives gradually improving estimates as sampling proceeds [9]. A variety of standard *Markov chain Monte Carlo* (MCMC) methods, including the *Gibbs sampling* and the *Metropolis–Hastings algorithm*, were used for approximate inference [4]. Other methods include the *loopy belief propagation* and *variational methods* [12] that exploit the **law of large numbers** to approximate large sums of random variables by their means.

BN Learning

In many practical settings the BN is unknown and one needs to learn it from the data. This problem is known as the *BN learning problem*, which can be stated informally as follows: Given training data and prior information (e.g., **expert knowledge**, **casual relationships**), estimate the graph topology (network structure) and the parameters of the JPD in the BN.

Learning the BN structure is considered a harder problem than learning the BN parameters. Moreover, another obstacle arises in situations of *partial observability* when nodes are hidden or when data is missing. In general, four BN learning cases are often considered, to which different learning methods are proposed, as seen in Table 1 [13].

In the first and simplest case the goal of learning is to find the values of the BN parameters (in each CPD) that maximize the (log)likelihood of the training

Table 1 Four cases of BN learning problems

Case	BN structure	Observability	Proposed learning method
1	Known	Full	Maximum-likelihood estimation
2	Known	Partial	EM (or gradient ascent), MCMC
3	Unknown	Full	Search through model space
4	Unknown	Partial	EM + search through model space

4 Bayesian Networks

dataset. This dataset contains m cases that are often assumed to be independent. Given training dataset $\Sigma = \{\mathbf{x}_1, \dots, \mathbf{x}_m\}$, where $\mathbf{x}_l = (x_{l1}, \dots, x_{ln})^T$, and the parameter set $\Theta = (\theta_1, \dots, \theta_n)$, where θ_i is the vector of parameters for the conditional distribution of variable X_i (represented by one node in the graph), the log-likelihood of the training dataset is a sum of terms, one for each node:

$$\log L(\Theta|\Sigma) = \sum_m \sum_n \log P(x_{li}|\pi_i, \theta_i) \quad (5)$$

The log-likelihood scoring function *decomposes* according to the graph structure; hence, one can maximize the contribution to the log-likelihood of each node independently [14]. Another alternative is to assign a prior **probability density function** to each parameter vector and use the training data to compute the posterior parameter distribution and the Bayes estimates. To compensate for zero occurrences of some sequences in the training dataset, one can use appropriate (mixtures of) conjugate prior distributions, e.g. the Dirichlet prior for the multinomial case as in the above backache example or the Wishart prior for the Gaussian case. Such an approach results in a maximum *a posteriori* estimate and is also known as the *equivalent sample size* (ESS) method.

In general, the other learning cases are computationally intractable. In the second case with known structure and partial observability, one can use the **EM (expectation maximization) algorithm** to find a locally optimal maximum-likelihood estimate of the parameters [4]. MCMC is an alternative approach that has been used to estimate the parameters of the BN model. In the third case, the goal is to learn a DAG that best explains the data. This is an NP-hard problem, since the number of DAGs on N variables is superexponential in N . One approach is to proceed with the simplest assumption that the variables are conditionally independent given a class, which is represented by a single common parent node to all the variable nodes. This structure corresponds to the *naïve BN*, which surprisingly is found to provide reasonably good results in some practical problems. To compute the Bayesian score in the fourth case with partial observability and unknown graph structure, one has to marginalize out the hidden nodes as well as the parameters. Since this is usually intractable, it is common to use an asymptotic approximation to the posterior called *Bayesian information criterion* (BIC) also known as the *minimum description*

length (MDL) approach. In this case one considers the trade-off effects between the likelihood term and a penalty term associated with the model complexity. An alternative approach is to conduct local search steps inside of the M step of the EM algorithm, known as *structural EM*, that presumably converges to a local maximum of the BIC score [7, 13].

BN and Other Markovian Probabilistic Models

It is well known that classic machine learning methods like Hidden Markov models (HMMs), **neural networks**, and Kalman filters can be considered as special cases of BNs [4, 13]. Specific types of BN models were developed to address stochastic processes, known as *dynamic BN*, and counterfactual information, known as *functional BN* [5]. Ben-Gal *et al.* [8] defined a hierarchical structure of Markovian GMs, which we follow here. The structure is described within the framework of DNA sequence classification, but is relevant to other research areas. The authors introduce the *variable-order Bayesian network* (VOBN) model as an extension of the *position weight matrix* (PWM) model, the *fixed-order Markov model* (MM) including HMMs, the *variable-order Markov* (VOM) model, and the BN model.

The PWM model is presumably the simplest and the most common context-independent model for DNA sequence classification. The basic assumption of the PWM model is that the random variables (e.g., nucleotides at different positions of the sequence) are statistically independent. Since this model has no memory it can be regarded as a fixed-order MM of order 0. In contrast, higher fixed-order models, such as MMs, HMMs, and interpolated MMs, rely on the statistical dependencies within the data to indicate repeating motifs in the sequence.

VOM models stand in between the above two types of models with respect to the number of model parameters. In fact, VOM models do not ignore statistical dependencies between variables in the sequence, yet, they take into account only those dependencies that are statistically significant. In contrast to fixed-order MMs, where the order is the same for all positions and for all contexts, in VOM models the order may vary for each position, based on its contexts.

Unlike the VOM models, which are homogeneous and which allow statistical dependences only between

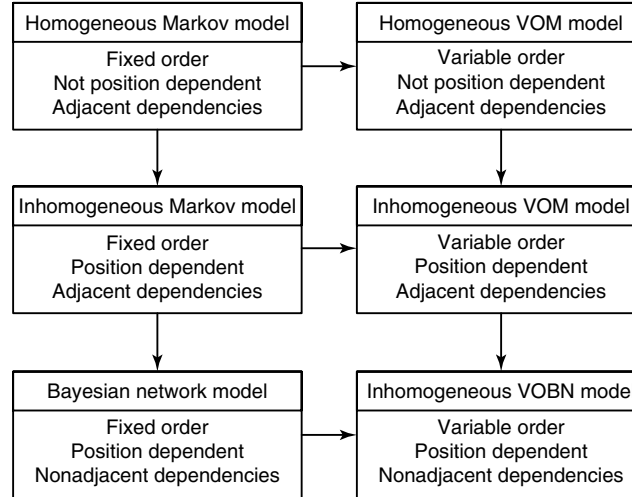


Figure 2 Hierarchical structure of Markovian graphical models [© OUP, 2005.]

adjacent variables in the sequence, VOBN models are inhomogeneous and allow statistical dependencies between nonadjacent positions in a manner similar to BN models. Yet, as opposed to BN models, where the order of the model at a given node depends only on the size of the set of its parents, in VOBN models the order also depends on the context, i.e. on the specific observed realization in each set of parents. As a result, the number of parameters that need to be estimated in VOBN models is potentially smaller than in BN models, yielding a smaller chance for overfitting of the VOBN model to the training dataset. Context-specific BNs (e.g., [15, 16]) are closely related to, yet constructed differently from, the VOBN models [8].

To summarize, the VOBN model can be regarded as an extension of PWM, fixed-order Markov, and BN models as well as VOM models in the sense that these four models are special cases of the VOBN model. This means that in cases where statistical dependencies are insignificant, the VOBN model degenerates to the PWM model. If statistical dependencies exist only between adjacent positions in the sequence and the memory length is identical for all contexts, the VOBN model degenerates to an inhomogeneous fixed-order MM. If, in addition, the CPDs are identical for all positions, the VOBN model degenerates to a homogeneous fixed-order MM. If the memory length for a given position is identical for all contexts and depends only on the number

of parents, the VOBN model degenerates to a BN model. If the context-dependent statistical dependencies in the VOBN model are restricted to adjacent positions, the VOBN model degenerates to the inhomogeneous VOM model. If, in addition, the context-dependent CPDs are identical for all positions, the VOBN model degenerates to a homogeneous VOM model. Figure 2 sketches these relationships between fixed-order MMs, BN models, VOM models, and VOBN models.

Summary

BNs became extremely popular models in the last decade. They have been used for applications in various areas, such as machine learning, text mining, natural language processing, speech recognition, signal processing, bioinformatics, error-control codes, medical diagnosis, weather forecasting, and cellular networks.

The name BNs might be misleading. Although the use of Bayesian statistics in conjunction with BN provides an efficient approach for avoiding data overfitting, the use of BN models does not necessarily imply a commitment to Bayesian statistics. In fact, practitioners often follow frequentists' methods to estimate the parameters of the BN. On the other hand, in a general form of the graph, the nodes can represent not only random variables but also

hypotheses, beliefs, and latent variables [13]. Such a structure is intuitively appealing and convenient for the representation of both causal and probabilistic semantics. As indicated by David [17], this structure is ideal for combining prior knowledge, which often comes in causal form, and observed data. BN can be used, even in the case of missing data, to learn the causal relationships and gain an understanding of the various problem domains and to predict future events.

Acknowledgment

The author would like to thank Prof. Yigal Gerchak for reviewing the final manuscript.

References

- [1] Jordan, M.I. (1999). *Learning in Graphical Models*, MIT Press, Cambridge.
- [2] Stich, T. (2004). Bayesian networks and structure learning, Diploma Thesis, Computer Science and Engineering, University of Mannheim, available at: <http://66.102.1.104/scholar?hl=en&lr=&q=cache:j36KPn-8hWroJ:www.timostich.de/resources/thesis.pdf>.
- [3] Pearl, J. (1988). *Probabilistic Reasoning in Intelligent Systems*, Morgan Kaufmann, San Francisco.
- [4] Griffiths, T.L. & Yuille, A. (2006). A primer on probabilistic inference, *Trends in Cognitive Sciences* Supplement to special issue on Probabilistic Models of Cognition, **10**(7), 1–11.
- [5] Pearl, J. & Russel, S. (2001). Bayesian networks. Report (R-277), November 2000, in *Handbook of Brain Theory and Neural Networks*, M. Arbib, ed, MIT Press, Cambridge, pp. 157–160.
- [6] Spirtes, P., Glymour, C. & Schienens, R. (1993). *Causation Prediction and Search*, Springer-Verlag, New York.
- [7] Friedman, N., Geiger, D. & Goldszmidt, M. (1997). Bayesian network classifiers, *Machine Learning* **29**, 131–163.
- [8] Ben-Gal, I., Shani, A., Gohr, A., Grau, J., Arviv, S., Shmilovici, A., Posch, S. & Grosse, I. (2005). Identification of transcription factor binding sites with variable-order Bayesian networks, *Bioinformatics* **21**(11), 2657–2666.
- [9] Pearl, J. (1987). Evidential reasoning using stochastic simulation of causal models, *Artificial Intelligence* **32**(2), 245–258.
- [10] Lauritzen, S.L. & Spiegelhalter, D.J. (1988). Local computations with probabilities on graphical structures and their application to expert systems (with discussion). *Journal of the Royal Statistical Society. Series B* **50**(2), 157–224.
- [11] Zhang, N.L. & Poole, D. (1996). Exploiting causal independence in Bayesian network inference, *Journal of Artificial Intelligence Research* **5**, 301–328.
- [12] Jordan, M.I., Ghahramani, Z., Jaakkola, T.S. & Saul, L.K. (1998). An introduction to variational methods for graphical models, in *Learning in Graphical Models*, M.I. Jordan, ed, Kluwer Academic Publishers, Dordrecht.
- [13] Murphy, K. (1998). *A brief introduction to graphical models and Bayesian networks*. <http://www.cs.ubc.ca/~murphyk/Bayes/bnintro.html>. Earlier version appears at Murphy K. (2001) The Bayes Net Toolbox for Matlab, Computing Science and Statistics, 33, 2001.
- [14] Aksoy, S. (2006). *Parametric Models: Bayesian Belief Networks*, Lecture Notes, Department of Computer Engineering Bilkent University, available at http://www.cs.bilkent.edu.tr/~saksoy/courses/cs551/slides/cs551_parametric4.pdf.
- [15] Boutilier, C., Friedman, N., Goldszmidt, M. & Koller, D. (1996). Context-specific independence in Bayesian networks, in *Proceedings of the 12th Conference on Uncertainty in Artificial Intelligence*, Portland, August 1–4 1996, pp. 115–123.
- [16] Friedman, N. & Goldszmidt, M. (1996). Learning Bayesian networks with local structure, in *Proceedings of the 12th Conference on Uncertainty in Artificial Intelligence*, Portland, August 1–4 1996.
- [17] David, H. (1999). A tutorial on learning with Bayesian networks, in *Learning in Graphical Models*, M.J. Models, ed, MIT Press, Cambridge, Also appears as Technical Report MSR-TR-95-06, Microsoft Research, March, 1995. An earlier version appears as Bayesian Networks for Data Mining, Data Mining and Knowledge Discovery, 1:79–119, 1997.

Further Reading

- Geman, S. & Geman, D. (1984). Stochastic relaxation, Gibbs distributions, and the Bayesian restoration of images, *IEEE Transactions on Pattern Analysis and Machine Intelligence* **6**, 721–741.
- Metropolis, A.W., Rosenbluth, A.W., Rosenbluth, M.N., Teller, A.H. & Teller, E. (1953). Equations of state calculations by fast computing machines, *Journal of Chemical Physics* **21**, 1087–1092.
- Tenenbaum, J.B., Griffiths, T.L. & Kemp, C. (2006). Theory-based Bayesian models of inductive learning and reasoning, *Trends in Cognitive Science* **10**, 309–318.

Related Article

Bayesian Networks in Reliability.

IRAD BEN-GAL

Bayesian Control Charts

Introduction

Standard control charts proposed by Shewhart (*see Control Charts, Overview*) are based on inspecting a sample of size n at intervals of h time units and issuing an alarm if the result of the sample is k standard deviation worse than what would be expected if the process were in its desired or target state. Usually, k is assumed equal to 3. The alarm signals that the process is deemed to be in an out-of-control state, which is characterized by deteriorated performances. Bayesian approaches for process monitoring have the same main objective as the Shewhart scheme; however, they use Bayes' theorem to update the information on the process state. Among these approaches, two basic schemes can be distinguished. The first is based on the seminal work by Girshick and Rubin [1], who first showed that all the information on the process, accumulated from the beginning of the production run, can be summarized in the posterior probability that the process is out of control. According to this approach, an out-of-control signal should be issued depending on this posterior probability. The set of approaches relying on this assumption are presented in the section titled "Bayesian Control Charting Based on the Posterior Probability that the Process is Out of Control", where it is shown that these Bayes-type control charts result in dynamic or adaptive schemes, that is, control charts in which at least one of the parameters n , h , and k can vary depending on information coming out from the process (*see Variable Sampling Rate Control Charts*). Connections with cumulative sum (CUSUM) control charts and Shewhart control charts are further shown in this section.

A second type of Bayesian control scheme found the out-of-control rule on the posterior distribution of the parameter of interest (e.g., the process mean if the objective is to maintain control over the process level). The section titled "Bayesian Control Charting Based on the Posterior Distribution of Process Parameters" briefly reviews approaches which share this common framework, outlining the main differences among them.

Bayesian Control Charting Based on the Posterior Probability that the Process is Out of Control

Assume that the production process can be in two possible states and let S_t represent the state at time t . The first state ($S_t = 0$) is the desirable or target state, usually referred to as the "in-control state". The second state ($S_t = 1$) is the state in which the process has deteriorated performance (e.g., a higher number of nonconforming items is produced) and is hence referred to as the "out-of-control" state. The shift from the in-control to the out-of-control state is due to an **assignable cause**, which is assumed to occur according to a **Poisson process**. Given this assumption, the time of occurrence of the assignable cause is an exponentially distributed random variable with mean $1/\nu$.

Unfortunately, the state of the process cannot be directly observed. The observable quantity at time t is denoted by Y_t and is related to the state of the underlying or "core" process through the probability densities $f_0(y_t)$ and $f_1(y_t)$ by the following relationships:

$$Y_t|S_t \sim \begin{cases} f_0(y_t) & \text{for } S_t = 0 \\ & \text{(process is in control)} \\ f_1(y_t) & \text{for } S_t = 1 \\ & \text{(process is out of control)} \end{cases} \quad (1)$$

Assume that the observable quantity Y_t can be inspected each h time units, because of technical reasons. Therefore, without loss of generality, the time index t will denote the generic t th inspection time, that is, inspection will occur at time $h, 2h, \dots, th, \dots$.

The control procedure consists in observing the quantity Y_t at each inspection epoch in order to decide whether the state of the core process is 0 or 1. If the process is considered to be in state 1 (the out-of-control state), the right decision should be to stop the process, search for the assignable cause and remove it, thus causing the process to go back to its nominal condition. Otherwise, if the process is considered to be in state 0 (the in-control state), no stop should be performed and the process should continue.

Girshick and Rubin [1] derived the optimal rule required to detect the out-of-control state when the objective is to minimize the expected costs due to false alarms, restorations, and process operation.

2 Bayesian Control Charts

Clearly, the operating costs associated with the out-of-control state are assumed to be greater than the ones associated with the in-control state (otherwise there will be no economical reason to bring back the process to its nominal condition). As showed by Girshick and Rubin, the optimal control policy is Bayesian, since it is based the posterior probability that the process is out of control. In particular, let π_t represent the posterior probability that the process is out of control *after* the t th inspection epoch, that is after observing y_t . The optimal decision consists in issuing an alarm if:

$$\pi_t \geq \pi^* \quad (2)$$

where the threshold π^* depends on the costs associated to repair actions and operating conditions.^a If the alarm is issued, the process is inspected and, if the assignable cause is found, restored. After restoration, the process returns to the initial state (i.e., π_t is reset to π_0). If no alarm is issued, the process continues producing.

The posterior probability, which has a central role in the out-of-control condition (2), can be sequentially updated as summarized in Figure 1. In particular, the posterior probability that the process is out of control *after* observing y_t can be computed through Bayes' theorem:

$$\pi_t = \frac{\pi'_{t-1} f_1(y_t)}{\pi'_{t-1} f_1(y_t) + (1 - \pi'_{t-1}) f_0(y_t)} \quad (3)$$

where π'_{t-1} represents the prior probability that the process is out of control *before* observing y_t and can be computed as:

$$\pi'_{t-1} = \pi_{t-1} + (1 - \pi_{t-1})\gamma \quad (4)$$

where π_{t-1} represents the state of our knowledge after observing the previous $(t-1)$ th sample (see Figure 1) and $\gamma = 1 - e^{-\nu h}$ is the probability that the process shifted to the out-of-control state within the last h time units, given that it started in the in-control state.

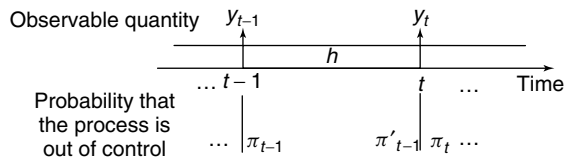


Figure 1 Sequential update of the knowledge about the state of the process

Equations (3) and (4) can be used to sequentially update the process-state knowledge. Obviously, these equations require an initial value for the probability π_0 of starting at $t = 0$ in the out-of-control state. This initial value is usually set equal to zero to represent a perfect process setup.

By manipulating equation (3) slightly, we can obtain a different expression of the out-of-control rule, given by:

$$\ln \frac{f_1(y_t)}{f_0(y_t)} \geq \ln \left[\frac{(1 - \pi'_{t-1})}{\pi'_{t-1}} \frac{\pi^*}{(1 - \pi^*)} \right] \quad (5)$$

Therefore, the likelihood ratio $f_1(y_t)/f_0(y_t)$, which plays a relevant role in other control schemes such as the CUSUM control chart (*see Cumulative Sum (CUSUM) Chart*; [2]), has a central task in the Bayesian control charting, too. Furthermore, the Bayesian approach has another similarity with the CUSUM control chart. Both these monitoring schemes belong to the class of *memory control charts* [3]. Memory control charts, such as the moving-average (MA), the exponentially weighted moving-average (EWMA) (*see Exponentially Weighted Moving Average (EWMA) Control Chart*), and the CUSUM control charts, incorporate both present and past sample information in order to decide whether the process is deemed out of control. Despite this similarity, a key feature distinguishes the Bayesian approach from other memory control schemes. As noted by Basseville and Nikiforov [2], the Bayesian approach requires some knowledge of the process shift mechanism, which is not necessarily required in other memory control charts. In particular, it is based on the assumption of Poisson occurrence of the assignable cause and requires the availability of γ – the probability that the process shifts in the out-of-control state between two inspections epochs – in order to update the knowledge on the process state through equation (4).

In contrast with memory control charts, Shewhart approaches adopt a different scheme: at any given period the process is deemed out of control by considering only the last sample information. This is why Shewhart control charts can be referred to as *memoryless control charts*. Despite this main difference, a relationship between the Bayesian approach and the traditional Shewhart control charting can be outlined, as shown in [4]. This connection will be shown with reference to the $\bar{X}$ control chart (*see Control Charts for the Mean*). In this case, when the process

operates in the in-control state ($S_t = 0$), the quality characteristic X_t is normally distributed with mean equal to the target value μ_0 and variance σ^2 . On the other hand, when the process moves in the out-of-control state ($S_t = 1$), the quality characteristic X_t experiences a sudden shift in the mean level, which moves from μ_0 to $\mu_1 = \mu_0 + \delta\sigma/\sqrt{n}$. The assignable cause does not induce any change in the variance, which remains at the in-control level σ^2 . Each h time units, a sample of fixed size $n \geq 1$ is inspected to compute the sample mean $\bar{X}_t = \sum_{i=1}^n x_{ti}$. This sample mean is the observable quantity Y_t and has a **normal distribution** given by $f_j(y_t) = N(y_t; \mu_j, \sigma^2/n)$, where $j = \{0, 1\}$ represents the state S_t of the core process at time t . Given $f_0(y_t)$ and $f_1(y_t)$, the out-of-control condition (5) can be rewritten as:

$$-\frac{1}{(2\sigma^2/n)} \left[\left(y_t - \mu_0 - \frac{\delta\sigma}{\sqrt{n}} \right)^2 - (y_t - \mu_0)^2 \right] \geq \ln \left[\frac{(1 - \pi'_{t-1})\pi^*}{(1 - \pi^*)\pi'_{t-1}} \right] \quad (6)$$

which, after some easy manipulation, results in:

$$y_t \geq UCL_t = \mu_0 + \frac{k_t\sigma}{\sqrt{n}} \quad (7)$$

where

$$k_t = \delta/2 + 1/\delta \ln \left[\frac{(1 - \pi'_{t-1})\pi^*}{(1 - \pi^*)\pi'_{t-1}} \right] \quad (8)$$

The Shewhart $\bar{X}$ control chart has a quite similar scheme. The only difference is that the upper control limit UCL_t , which is designed to detect an increase in the mean, has a constant $k_t = k$ in equation (7). Therefore, the Bayesian approach can be interpreted as an adaptive Shewhart chart, where the position of the control limit adapts on line according to the information coming out from the process.

Figure 2 shows an example of this adaptive behavior for the simple case of $n = 1$. In this example, the in-control state is characterized by a target mean $\mu_0 = 5$, while the assignable cause induces a shift in the mean level at $\mu_1 = 6$. The assignable cause is supposed to arise at time $t = 21$. In spite of the process state, the variance of the quality characteristic is constant and equal to $\sigma^2 = 1$. The control limit in Figure 2 is computed assuming $\delta = 1$ (i.e.,

the control chart is designed to detect a shift size equal to 1σ), $\gamma = 0.002$, $\pi^* = 0.4$ and assuming that the process is known to start in the in-control state ($\pi_0 = 0$). As shown in Figure 2, the Bayesian control chart adapts the control limit UCL_t depending on the information coming out from the process. In particular, after the assignable cause occurred at time $t = 21$, the upper control limit starts to come closer and closer to the observed quantity until an out-of-control signal is issued at time $t = 27$.

Up to this point, the optimal policy was obtained by minimizing the expected total cost of operation and restorations. Therefore the Bayesian control rule described above is optimal when inspection costs can be neglected. In the last part of their paper [1], Girshick and Rubin revised the optimal policy described above, by including inspection costs in the objective function. In this case, the optimal control scheme can suggest skipping some inspections in order to minimize the expected costs. The resulting control policy can hence be defined as an “adaptive sampling interval” policy. In particular, the optimal policy consists in computing the posterior probability π_t that the process is out of control at time t and

continue production and skip the next inspection
if $\pi_t < \pi_1^*$

continue production and perform the next inspection

if $\pi_1^* \leq \pi_t < \pi_2^*$ (9)

stop production and search for the assignable cause

if $\pi_2^* \leq \pi_t$

where π_1^* and π_2^* are thresholds which can be computed by minimizing the expected costs.

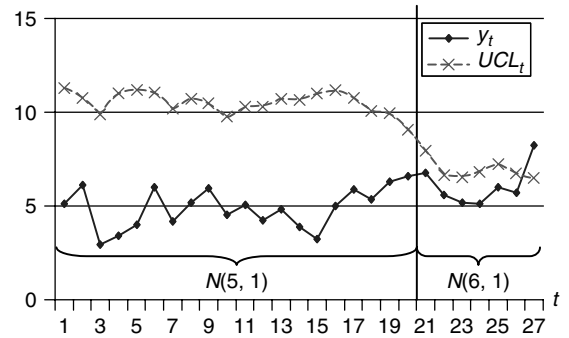


Figure 2 An example of a Bayesian control chart

Extension of Bayesian Control Charting Based on the Posterior Probability that the Process is Out of Control

Starting from the basic scheme proposed by Girshick and Rubin [1], different authors studied this type of Bayesian control charting. In particular, the computation of the optimal thresholds in both the simple (equation 2) and the adaptive (equation 9) policies attracted the main interest of researchers. Taylor [5, 6] focused on the first type of Bayesian procedure, the one which neglects the inspection costs and is based on a unique threshold π^* . The author proposed a simplified approach aimed at deriving this threshold as a function of costs and system parameters, when small shifts are of interest. The approach proposed is based on results reported by Shiryaev [7], who studied a similar decision problem in continuous time.

More recently, different approaches [4, 8–13] focused on proposing a dynamic programming formulation to design the optimal Bayesian policy for finite production runs. Finite production runs arise when parts are processed in lots whose sizes are small and known in advance. After one lot has been processed, a setup (or maintenance) operation is executed and the process is brought back to an initial state. All the approaches presented showed that the Bayesian policy induces significant savings over traditional control charts, even when the same economic objective function is used to design both the approaches (for the economic design of traditional control charts, *see Economic Design of Control Charts*). In particular, Calabrese [8] and Porteus and Angelus [14] dealt with Bayesian control charts for attributes (*see Control Charts for Attributes*), assuming the fraction of nonconforming items as the quality characteristic of interest. Calabrese [8] also showed that the control problem can be formulated as a partially observable Markov decision process [9].

In his long-lasting interest on Bayesian control charts, Tagaras [4, 10–13] focused on control charts for variables (in particular, control charts to detect a shift in the level of the process, that is, $\bar{X}$ control charts). In his first work [4], Tagaras proposed a dynamic programming formulation to derive the optimal control policy when **adaptive sampling** intervals are allowed. In [10], the author further relaxed the assumption of fixed **sample size** and proposed a Bayesian control chart which allows

both adaptive sampling intervals and adaptive sample sizes. In [12], the performances of all the different Bayesian adaptive schemes (with or without adaptive sampling intervals and adaptive sample sizes) were compared. The main conclusion of this study was that most of the economic advantages can be achieved by allowing an adaptive sampling interval, even with a fixed sample size. In a recent paper, Tagaras and Nenes [13] further extended the Bayesian approach to a two-sided shift (increase or decrease) of the mean level.

Bayesian Control Charting Based on the Posterior Distribution of Process Parameters

The Bayesian approaches discussed in the previous section share a common structure, in which the alarm is issued on the basis of the posterior probability that the process is out of control. A different point of view considered in the literature, assumes that the out-of-control rule can be based on the posterior distribution of the parameter of interest. This posterior distribution can be sequentially computed *via* Bayes' theorem, since the posterior probability computed after observing the previous sample can be used as the prior before observing the current one. Approaches proposed in this area present many differences concerning the process parameter of interest, the in-control process model and the specific details of the out-of-control rule. In the following, these approaches will be reviewed outlining these specific features.

Although starting from an **estimation** problem, Chernoff and Zacks [15] formulated a Bayesian scheme to detect shifts in the process mean of a normally distributed quality characteristic. Given some knowledge on the shift mechanism, the posterior distribution of the process level is computed to check whether a shift has already occurred in a set of available samples. Given this specific objective of the proposed procedure, this approach can be adopted in the design stage, usually known as phase I, of control charting.

Wasserman and Sudjianto [16] proposed a control chart for a second-order dynamic linear model whose mean is subject to trends or level changes and attention is devoted to short production runs (*see Control Charts for Short Production Runs*). The

proposed control chart consists in computing mean and variance of the posterior distribution of the process level and then drawing a confidence band as in traditional Shewhart control charts (i.e., posterior mean plus and minus k times the posterior standard deviation). An alarm is issued if the nominal or target process level falls outside the control band.

Tsiamirizis and Hawkins (see **Statistical Process Control, Bayesian**; [17, 18]) proposed a Bayesian scheme for a process whose mean drifts according to a random walk when the process is in control. When an assignable cause occurs, the mean is subject to a random shock. As in [16], the parameter of interest is the process mean and the objective of the control chart is to detect a shift in the process mean. The authors proposed two out-of-control rules. The first is based on computing the posterior probability that the posterior mean crossed a given threshold. An out-of-control signal is issued if this probability is greater than a given value, which can in turn be related to the costs of type I and type II errors using a 0–1 loss function. The second out-of-control rule proposed is based on the Bayes factor, which gives the ratio of the posterior to prior odds.

Alt [19] proposed a Bayesian control chart for monitoring the mean of a multinormal process (see **Multivariate Control Charts Overview**). Again, the parameter of interest is the process mean and the out-of-control rule is based on the posterior mean and **covariance** matrix of the process level. The performance of this Bayesian approach is compared with the ones attained by different multivariate (CUSUM and EWMA) control chart in terms of the average run lengths (see **Average Run Lengths and Operating Characteristic Curves**).

Shiau and Feltz [20] proposed different Bayesian control charts for univariate and multivariate quality characteristics which can be continuous or categorical. Also in these approaches the posterior distribution plays a central role although an empirical Bayes approach is adopted in its computation.

Eventually, Bayarri and Garcia-Donato [21], presented a Bayesian control chart for attributes (see **Control Charts for Attributes**) aimed at monitoring the number of nonconformities (defects) per inspection unit. In this approach, the out-of-control rule is slightly different from the previous ones, since it is based on the posterior predictive distribution of the characteristic of interest.

This brief review shows the lack of a common framework in this second stream of research on Bayesian control charting. Clearly, approaches proposed in this area cannot be directly compared since they start from different assumptions on the in-control process model (with or without dynamic) and on the quality characteristic of interest (variable *vs* attributes, univariate *vs* multivariate variables). However, further effort could be usefully devoted to compare the different out-of-control rules proposed in order to provide an unifying guide for practical applications.

End Note

^a Girshick and Rubin [1] derived the optimal policy with reference to a slightly different function of π_t . However, the function they used is a monotonic increasing function of π_t . This is why the optimal policy can be directly expressed in terms of π_t .

References

- [1] Girshick, M. & Rubin, H. (1952). A Bayes' approach to a quality control model, *Annals of Mathematical Statistics* **23**, 114–125.
- [2] Basseville, M. & Nikiforov, I.V. (1993). *Detection of Abrupt Changes: Theory and Application*, Prentice-Hall.
- [3] Alwan, L.C. (2000). *Statistical Process Analysis*, Irwin McGraw-Hill.
- [4] Tagaras, G. (1994). A dynamic programming approach to the economic design of X-charts, *IIE Transactions* **26**(3), 48–56.
- [5] Taylor, H.M. (1965). Markovian sequential replacement processes, *Annals of Mathematical Statistics* **36**, 1677–1694.
- [6] Taylor, H.M. (1967). Statistical control of a Gaussian process, *Technometrics* **9**(1), 29–41.
- [7] Shiryaev, A.N. (1963). On optimum methods in quickest detection problems, *Theory Probability and Applications* **8**(1), 22–46.
- [8] Calabrese, J.M. (1995). Bayesian process control for attributes, *Management Science* **41**(4), 637–645.
- [9] Monahan, G.E. (1982). A survey of partially observable Markov decision processes: theory, models, and algorithms, *Management Science* **28**(1), 1–16.
- [10] Tagaras, G. (1996). Dynamic control charts for finite production runs, *European Journal of Operational Research* **91**, 38–55.
- [11] Tagaras, G. (1998). A survey of recent developments in the design of adaptive control charts, *Journal of Quality Technology* **30**(3), 212–231.

- [12] Tagaras, G. & Nikolaidis, Y. (2002). Comparing the effectiveness of various Bayesian $\bar{X}$ control charts, *Operations Research* **50**, 878–888.
- [13] Tagaras, G. & Nenes, G. (2006). Two-sided Bayesian $\bar{X}$ control charts for short production runs, in *Bayesian Process Monitoring, Control and Optimization*, B.M. Colosimo & E. del Castillo, eds, Chapman & Hall/CRC Press, Boca Raton.
- [14] Porteus, E.L. & Angelus, A. (1997). Opportunities for improved statistical process control, *Management Science* **43**(9), 1214–1228.
- [15] Chernoff, H. & Zacks, S. (1964). Estimating the current mean of a normal distribution which is subject to change in time, *Annals of Mathematical Statistics* **35**(3), 999–1018.
- [16] Wasserman, G.S. & Sudjianto, A. (1993). Short run SPC based upon the second order dynamic linear model for trend detection, *Communications in Statistics - Computation and Simulation* **22**(4), 1011–1036.
- [17] Tsiamyrtzis, P. & Hawkins, D.M. (2005). A Bayesian scheme to detect changes in the mean of a short-run process, *Technometrics* **47**(4), 446–456.
- [18] Tsiamyrtzis, P. & Hawkins, D.M. (2006). A Bayesian approach to statistical process control, in *Bayesian Process Monitoring, Control and Optimization*, B.M. Colosimo & E. del Castillo, eds, Chapman & Hall/CRC Press, Boca Raton.
- [19] Alt, P. (2006). A Bayesian approach to monitoring the mean of a multivariate normal process, in *Bayesian Process Monitoring, Control and Optimization*, B.M. Colosimo & E. del Castillo, eds, Chapman & Hall/CRC Press, Boca Raton.
- [20] Shiau, J.H. & Feltz, C.J. (2006). Empirical Bayes process monitoring, in *Bayesian Process Monitoring, Control and Optimization*, B.M. Colosimo & E. del Castillo, eds, Chapman & Hall/CRC Press, Boca Raton.
- [21] Bayarri, M.J. & Garcia-Donato, G. (2005). A Bayesian sequential look at U-control charts, *Technometrics* **47**(2), 142–151.

Related Articles

Average Run Lengths and Operating Characteristic Curves; Control Charts for Attributes; Control Charts for Short Production Runs; Control Charts for the Mean; Control Charts, Overview; Cumulative Sum (CUSUM) Chart; Economic Design of Control Charts; Exponentially Weighted Moving Average (EWMA) Control Chart; Multivariate Control Charts Overview; Statistical Process Control, Bayesian; Variable Sampling Rate Control Charts.

BIANCA M. COLOSIMO

Degradation Processes

Introduction

Many failure mechanisms in engineering, medical, social, and economic settings can be traced to an underlying **degradation** process. Most items under study degrade physically over time and a measurable physical deterioration almost always precedes failure. Examples of items that degrade over time are devices subject to wear in reliability and the immune system of HIV infected individuals in biostatistics. In an engineering setting, degradation is the irreversible accumulation of damage throughout life that leads to failure (*cf.* [1]) and may involve corrosion, material fatigue, wearing out, and fracturing. In the biosciences, degradation may be characterized by markers of health status and quality of life data.

Traditionally, statistical reliability **estimation** is mostly based on failure-time data. In practical applications, collecting samples of lifetimes can be time consuming and costly because the systems and devices under study are often highly reliable and expensive. An increasing emphasis on product reliability and short product-development cycles in industry results in few or no failures during reliability tests. Accelerated **life tests** can be used but the acceleration methods may not properly imitate the actual failure process and, even with acceleration, only few failures may occur. In reliability analysis, biostatistics, and econometrics longitudinal data on degradation and covariates are frequently available additionally to time-to-event data. In many cases, collecting degradation data may be difficult and costly but degradation data can provide much more reliability information than traditional failure-time data. When additional longitudinal data are available, an alternative approach to classical reliability and survival analysis is the modeling of degradation and environmental factors and of their failure-generating mechanisms by stochastic processes. This stochastic-process-based approach shows great flexibility and can give rise to new or alternative time-to-failure distributions defined by the **degradation model**. It is particularly useful when the application of traditional reliability models based only on failure and survival data is limited due to rare failures of highly reliable items or due to items operating in dynamic environments.

Models for Degradation Data

In engineering applications, degradation data can be obtained by direct measurements of **physical degradation** such as crack growth and tire wear. If actual physical degradation cannot be observed directly, surrogates like a deteriorating product performance parameter (e.g., power output), usage or exposure measures, and accumulated maintenance costs are usually referred to as *degradation data*, as well. In the biosciences, degradation data can be measurements of a biomarker like CD4 cell counts or quality of life data. As degradation is a complex random mechanism it is best represented by a stochastic process $X = \{X(t) : t \geq 0\}$. In most applications, the monotonicity of degradation paths is a reasonable assumption because most real degradation processes are increasing. However, there are situations in which degradation need not be increasing. For instance, fatigue cracks sometimes show a tendency to heal and CD4 cell counts may fluctuate. A further cause for nonmonotone degradation paths is the fact that degradation is measured with error and the missing monotonicity is due to the error term whereas the actual degradation is monotone. Degradation processes are frequently described by general path models and by processes with stationary and independent increments. More sophisticated models for degradation processes are Markov diffusion processes (e.g., the Ornstein–Uhlenbeck process), **Markov processes**, and marked point processes (*cf.* [2, 3]).

General Path Models

In applied literature degradation processes are frequently described by a general path model $X(t) = g(t, A)$, that is, by a stochastic process that depends on a finite dimensional random variable $A = (A_1, \dots, A_r)$ via a deterministic function g . In [4, 5], and [6], a Paris growth curve path model was used for alloy fatigue crack size data. Moreover, general path degradation models with linear, convex, and concave functions g were applied to tire wear, electronic component degradation, silicon solar cell degradation, and disk storage media (*cf.* [4, 5, 7, 8]). Meeker and Escobar [5] consider general path degradation models in more detail (*see also Degradation and Failure*).

2 Degradation Processes

Lévy Process Models

A Lévy process is a stochastic process with stationary and independent increments (*see Lévy Processes*). The application of such processes as degradation models can be based on the fact that in systems subject to shocks, the order of damage causing shocks is often irrelevant. So, the increments of the accumulated damage process are stationary and exchangeable random variables which is somewhat weaker than stationary and independent increments. The most widely used Lévy processes for degradation modeling are *compound Poisson processes*, the *Wiener process with drift*, and the *gamma process*. When degradation evolves as an accumulation of shocks, each of them resulting in a random degradation increment, a compound **Poisson process** is an appropriate model. It is given by $X(t) = \sum_{i=1}^{N(t)} X_i$ where N is a Poisson counting process jumping at shock times and $\{X_1, X_2, \dots\}$ is a sequence of independent identically distributed (i.i.d.) random variables that represent the random increments of degradation caused by shocks and that are independent of N . The Wiener process with drift is a Gaussian process given by $X(t) = x_0 + \mu t + \sigma W(t)$ where W denotes a standard **Brownian motion**, x_0 is some initial degradation level, and μ and σ are drift and variance coefficients. The Wiener process is the only Lévy process with continuous sample paths, however, its paths are not monotone. The gamma process has increasing sample paths but it is a pure jump process.

A Lévy process Y has linear mean and variance functions. To cover nonlinear degradation behavior a timescale $\tau = \tau(t)$ can be used, that is, a positive, increasing function of real time with $\tau(0) = 0$. The timescale drives the speed of deterioration and describes slowing or accelerating degradation in real time. It measures the physical progress of degradation and is referred to as operational time. For instance, Bagdonavičius and Nikulin [9] consider a linear timescale with an initial acceleration period

$$\tau(t) = \gamma_1 t + \gamma_2 (1 - \exp(-\gamma_3 t)), \quad \gamma_i > 0 \quad (1)$$

for tire protector wear and Whitmore and Schenkelberg [10]

$$\begin{aligned} \tau(t) &= 1 - \exp(-\gamma_1 t^{\gamma_2}) \text{ and} \\ \tau(t) &= t^{\gamma_1}, \quad \gamma_i > 0 \end{aligned} \quad (2)$$

in the context of self-regulating heating cables. The choice of τ depends on, whether degradation is unbounded or approaches a saturation point. A further possible form of the timescale is a **Box-Cox transformation**

$$\tau(t) = \begin{cases} \gamma_0((t+1)^{\gamma_1} - 1)/\gamma_1, & \gamma_1 > 0 \\ \gamma_0 \log(t+1), & \gamma_1 = 0 \end{cases} \quad (3)$$

To model the influence on failure of an item's dynamic operating environment, the timescale may depend on a covariate process Z , for instance on different loads, stress levels, or usage measures, that is, $\tau_Z(t) = \tau(t, Z)$. An example is the model of additive accumulation of damage:

$$\tau_Z(t) = \tau \left(\int_0^t \exp(\beta^T Z(s)) ds \right) \quad (4)$$

The degradation model is given by

$$X(t) = X_0 + Y(\tau_Z(t)) =: X_0 + Y^Z(t) \quad (5)$$

where Y^Z is the timescaled Lévy process Y and X_0 is some possibly random initial degradation level independent of (Y, Z) . Conditioned on $\{Z = z(\cdot)\}$, $Y^{z(\cdot)}$ is a process with independent but not stationary increments. Obviously,

$$\begin{aligned} E[X(t) | Z] &= E[X_0] + \mu \tau_Z(t) \quad \text{and} \\ \text{Var}[X(t) | Z] &= \text{Var}[X_0] + \sigma^2 \tau_Z(t) \end{aligned} \quad (6)$$

with $\mu = E[Y(1)]$ and $\sigma^2 = \text{Var}[Y(1)]$. For example, let the baseline degradation process Y be a Wiener process with drift $Y(t) = \mu t + \sigma W(t)$ so that the timescaled degradation process X is given by $X(t) = X_0 + \mu \tau_Z(t) + \sigma W(\tau_Z(t))$. This model is appropriate if the actual degradation process $\mu \tau_Z(t)$ is latent but observable with some additive error term which has a variance proportional to the actual degradation process.

Among others Wiener process degradation models have found application in [11] and [12]. Whitmore ([13]) applied a Wiener process model with error to degradation data of transistor gain. Wiener processes with a nonrandom timescale were used as degradation models for heating cables and carbon-film resistors in [10] and [14]. A bivariate Wiener process degradation model was considered in [15] and in [16] and applied to equipment degradation data in aluminum

production and to biomarker data of AIDS clinical trials. In these models, one process is a latent or unobservable degradation process and the second one represents a marker process that is correlated with the degradation process and tracks its progress. Degradation models based on the gamma process were considered in [17], [2], and with application to crack growth data in [18]. In [9] a gamma process with a random timescale was applied to tire wear, and in [19], both geometric Brownian motion and gamma processes were considered in a timescaled degradation model.

Models Relating Degradation and Failure

Considerable interest has focused on joint models for the distribution of the failure time T and the degradation process X . Two model classes, *threshold models* and *shock models*, predominate in the engineering context. In a threshold model, X is *failure defining* by $T = \inf\{t \geq 0 : X(t) \geq X^*\}$ where X^* is a critical threshold level which is possibly random. This kind of failure, where an item is regarded as failed and usually switched off when its degradation level first reaches a critical threshold, is called *soft failure* (cf. [5]) or *nontraumatic failure*. In a shock model the item or system is subject to external shocks which may be survived or which lead to failure. This kind of failure is referred to as *hard failure* (cf. [5]) or *traumatic failure*. The shocks usually occur according to a Poisson process whose intensity depends on degradation and environmental factors. In a shock model X is *failure rate determining*, which means, that the conditional failure rate of T

$$\lambda(t, X) = \lim_{\Delta \downarrow 0} \frac{1}{\Delta} P(t \leq T < t + \Delta \mid X, T \geq t) \quad (7)$$

determines the conditional survival function via the exponential formula

$$P(T > t \mid X) = \exp\left(-\int_0^t \lambda(s, X) ds\right) \quad (8)$$

The class of reliability models, which combines the ideas of hard and soft failures by considering them as **competing risks**, is called *degradation-threshold-shock models* (DTS models) (cf. [20]). In a DTS model, additionally to soft failures which are immediately related to degradation, an item may also

fail when a traumatic event like a shock of large magnitude occurs although the degradation process has not yet reached the threshold. So, the observable failure time of an item is the minimum of the moment when degradation first reaches a critical threshold and the moment when a censoring traumatic event occurs. General expressions of the failure-time distribution in a DTS model and of the likelihood function based on the observation of identically distributed degradation processes X_i of n independent items at discrete times, their failure times T_i , and their failure modes (hard or soft failure) were given in [20]. DTS models based on gamma and Wiener degradation processes were considered in [21] and [17]. **Maximum likelihood** and semiparametric estimation in a DTS model with a time scaled gamma degradation process was investigated in [9]. In the context of DTS models, Bagdonavičius *et al.* [22] consider a linear path model with time-dependent covariates and multiple traumatic event modes. Nonparametric estimators in general path DTS models with multiple traumatic failure modes and renewals were analyzed in [23]. A DTS model based on Lévy processes for repairable items was introduced in [24]. Singpurwalla [2] gave a detailed review on stochastic-process-based reliability models including DTS models and Tsatis and Davidian [25] gave an overview on joint modeling of longitudinal and failure data in biostatistics.

References

- [1] Bogdanoff, J.L. & Kozin, F. (1985). *Probabilistic Methods of Cumulative Damage*, John Wiley & Sons, New York.
- [2] Singpurwalla, N.D. (1995). Survival in dynamic environments, *Statistical Science* **10**, 86–103.
- [3] Kahle, W. & Wendt, H. (2006). Statistical analysis of some parametric degradation models, in *Probability, Statistics and Modelling in Public Health*, M. Nikulin, D. Commenges & C. Huber, eds, Springer Science + Business Media, pp. 266–279.
- [4] Lu, C.J. & Meeker, W.Q. (1993). Using degradation measures to estimate a time-to-failure distribution, *Technometrics* **35**, 161–174.
- [5] Meeker, W.Q. & Escobar, L.A. (1998). *Statistical Methods for Reliability Data*, John Wiley & Sons, New York.
- [6] Couallier, V. (2006). Some recent results on joint degradation and failure time modelling, in *Probability, Statistics and Modelling in Public Health*, M. Nikulin, D. Commenges & C. Huber, eds, Springer Science + Business Media, pp. 73–89.

- [7] Carey, M.B. & Koenig, R.H. (1991). Reliability assessment based on accelerated degradation: a case study, *IEEE Transactions on Reliability* **40**, 499–506.
- [8] Meeker, W.Q., Escobar, L.A. & Lu, C.J. (1998). Accelerated degradation test: modeling and analysis, *Technometrics* **40**, 89–99.
- [9] Bagdonavičius, V. & Nikulin, M. (2001). Estimation in degradation models with explanatory variables, *Lifetime Data Analysis* **7**, 85–103.
- [10] Whitmore, G.A. & Schenkelberg, F. (1997). Modelling accelerated degradation data using Wiener diffusion with a time scale transformation, *Lifetime Data Analysis* **3**, 27–45.
- [11] Doksum, K.A. & Hoyland, A. (1992). Models for variable-stress accelerated life testing experiment based on a Wiener process and the inverse Gaussian distribution, *Technometrics* **34**, 74–82.
- [12] Doksum, K.A. & Normand, S.L.T. (1995). Gaussian models for degradation processes – part I: methods for the analysis of biomarker data, *Lifetime Data Analysis* **1**, 131–144.
- [13] Whitmore, G.A. (1995). Estimation degradation by a Wiener diffusion process subject to measurement error, *Lifetime Data Analysis* **1**, 307–319.
- [14] Padgett, W.J. & Tomlinson, M.A. (2004). Inference from accelerated degradation and failure data based on Gaussian process models, *Lifetime Data Analysis* **10**, 191–206.
- [15] Whitmore, G.A., Crowder, M.I. & Lawless, J. (1998). Failure inference from a marker process based on a bivariate Wiener model, *Lifetime Data Analysis* **4**, 229–251.
- [16] Lee, M.L.T., DeGruttola, V. & Schoenfeld, D. (2000). A model for markers and latent health status, *Journal of the Royal Statistical Society. Series B* **62**, 747–762.
- [17] Wenocur, M.L. (1989). A reliability model based on the gamma process and its analytical theory, *Advances in Applied Probability* **21**, 899–918.
- [18] Lawless, J. & Crowder, M. (2004). Covariates and random effects in a gamma process model with application to degradation and failure, *Lifetime Data Analysis* **10**, 213–227.
- [19] Padgett, W.J. & Tomlinson, M.A. (2005). Accelerated degradation models for failure based on geometric Brownian motion and gamma processes, *Lifetime Data Analysis* **11**, 511–527.
- [20] Lehmann, A. (2006). Degradation-threshold-shock models, in *Probability, Statistics and Modelling in Public Health*, M. Nikulin, D. Commenges & C. Huber, eds, Springer Science + Business Media, pp. 286–298.
- [21] Wenocur, M.L. (1986). Brownian motion with quadratic killing and some implications, *Journal of Applied Probability* **23**, 893–903.
- [22] Bagdonavičius, V., Bikelis, A. & Kazakečius, V. (2004). Statistical analysis of linear degradation and failure time data with multiple failure modes, *Lifetime Data Analysis* **10**, 65–81.
- [23] Bagdonavičius, V., Bikelis, A., Kazakečius, V. & Nikulin, M. (2006). Non-parametric estimation in degradation-renewal-failure models, in *Probability, Statistics and Modelling in Public Health*, M. Nikulin, D. Commenges & C. Huber, eds, Springer Science + Business Media, pp. 23–36.
- [24] Lehmann, A. (2004). On a degradation-failure model for repairable items, in *Parametric and Semiparametric Models and its Applications for Reliability, Survival Analysis and Quality of Life*, M. Nikulin, N. Balakrishnan, N. Limnios & M. Mesbah, eds, Birkhauser, Boston, pp. 65–79.
- [25] Tsiatis, A.A. & Davidian, M. (2004). Joint modeling of longitudinal and time-to-event data: an Overview, *Statistica Sinica* **14**, 809–834.

Related Articles

Cumulative Damage Models Based on Gamma Processes; Damage Processes; Degradation and Failure; Degradation Models; Physical Degradation Models.

AXEL LEHMANN

Hierarchical Markov Chain Monte Carlo (MCMC) for Bayesian System Reliability

Hierarchical Models

Hierarchical models are one of the central tools of Bayesian analysis. Broadly, Bayesian models have two parts (*see Bayesian Reliability Analysis*): the likelihood function and the prior distribution. We construct the likelihood function from the sampling distribution of the data, which describes the probability of observing the data before the experiment is performed. After we perform the experiment and observe the data, we can consider the sampling distribution as a function of the unknown parameters. This function is called the likelihood function. The prior distribution describes the uncertainty about the parameters of the likelihood function. We update the **prior distribution** to the posterior distribution after observing data. We use Bayes' theorem to perform the update, which shows that the posterior distribution is computed (up to a proportionality constant) by multiplying the likelihood function by the prior distribution.

Denote the sampling distribution as $f(\mathbf{y} | \theta)$ and the prior distribution as $g(\theta | \alpha)$, where α represents the parameters of the prior distribution (often called *hyperparameters*). It may be the case that we know $g(\theta | \alpha)$ completely, including a specific value for α . However, suppose that we do not know α , and that we choose to quantify our uncertainty about α using a distribution $h(\alpha)$ (often called the *hyperprior*). This is the most general form of a hierarchical model. A hierarchical model has three parts [1]:

1. The *observational model* for the data.

$$(Y_i | \theta) \sim f(y_i | \theta), \quad i = 1, \dots, k \quad (1)$$

2. The *structural model* for the parameters of the likelihood.

$$(\theta | \alpha) \sim g(\theta | \alpha) \quad (2)$$

3. The *hyperparameter model* for the parameters of the structural model.

$$\alpha \sim h(\alpha) \quad (3)$$

While this is the most general form of the hierarchical model, one of its most common applications is in problems where there are multiple parameters that are related by the structure of the problem. Consider the data in Table 1, which are analyzed in [2] and [3]. These data summarize the number of failures of pumps (s_i) over a period of time (t_i) in several systems (i) of the nuclear power plant Farley 1.

What is the appropriate observational model for the data in Table 1? We could assume that each pump has an identical failure rate, λ . We could then model the number of failures as $S_i \sim \text{Poisson}(\lambda t_i)$, $i = 1, \dots, 10$, or equivalently, $\sum_{i=1}^{10} S_i \sim \text{Poisson}(\lambda \sum_{i=1}^{10} t_i)$.

The assumption of a common failure rate, however, is quite strong and may be inappropriate. Suppose that instead of a common failure rate, we consider the failure rate for each pump, λ_i , as a sample from an overall population distribution. We have no information, other than the data, that distinguishes the failure rates, and we have no ordering or grouping of the rates. Formally, this suggests that we treat the rates as *exchangeable*, which implies that their joint distribution $g(\lambda_1, \dots, \lambda_{10})$ is invariant to permutations of the indices of the parameters.

By treating the rates as exchangeable, we propose the following hierarchical model:

1. An *observational model* for the data.

$$(S_i | \lambda_i) \sim \text{Poisson}(\lambda_i t_i), \quad i = 1, \dots, 10 \quad (4)$$

Table 1 Pump failures (t_i in thousands of hours)

System _{i}	s_i	t_i
1	5	94.320
2	1	15.720
3	5	62.880
4	14	125.760
5	3	5.240
6	19	31.440
7	1	1.048
8	1	1.048
9	4	2.096
10	22	10.480

2 Hierarchical MCMC for Bayesian System Reliability

2. A *structural model* for the parameters of the likelihood.

$$(\lambda_i | \alpha, \beta) \sim \text{Gamma}(\alpha, \beta), \quad i = 1, \dots, 10 \quad (5)$$

3. A *hyperparameter model* for the parameters of the structural model.

$$\alpha \sim \text{Uniform}(0.0, 5.0) \quad (6)$$

$$\beta \sim \text{Gamma}(0.1, 1.0) \quad (7)$$

This model allows the pump for each system to have its own failure rate, but it also models each failure rate as coming from a common population distribution. This hierarchical structure allows us to *borrow strength* for the **estimation** of each λ_i . In particular, the estimation of each λ_i is improved by using the failure data from the other pumps.

Markov Chain Monte Carlo (MCMC)

Given the hierarchical model for the data in Table 1, we now need to estimate the posterior distributions for λ_i , $i = 1, \dots, 10$, α , and β . The posterior distribution is proportional to

$$\begin{aligned} & p(\lambda_1, \dots, \lambda_{10}, \alpha, \beta | \mathbf{s}, \mathbf{t}) \\ & \propto \exp\left(-\beta - \sum_{i=1}^{10} (\beta + t_i) \lambda_i\right) \left(\prod_{i=1}^{10} \lambda_i^{s_i + \alpha - 1}\right) \\ & \quad \times \beta^{10\alpha + 0.1 - 1} (\Gamma(\alpha))^{-10} I(\alpha \in (0, 5)) \end{aligned} \quad (8)$$

(For more details on the derivation, see [3].)

In Bayesian analysis, the posterior distribution provides the “answer”; however, it is not immediately obvious how to use this expression for $p(\lambda_1, \dots, \lambda_{10}, \alpha, \beta | \mathbf{s}, \mathbf{t})$ to answer questions of interest. In particular, suppose that we would like to know the expected rate of failure for an arbitrary pump from this population, given the observed data. Mathematically, we need to calculate

$$\int \frac{\alpha}{\beta} p(\lambda_1, \dots, \lambda_{10}, \alpha, \beta | \mathbf{s}, \mathbf{t}) d\lambda_1, \dots, d\lambda_{10} d\alpha d\beta \quad (9)$$

Monte Carlo integration can approximate this integral: we draw a sample $X^{(j)} = (\lambda_1^{(j)}, \dots, \lambda_{10}^{(j)}, \alpha^{(j)},$

$\beta^{(j)})$, $j = 1, \dots, n$, from $p(\lambda_1, \dots, \lambda_{10}, \alpha, \beta | \mathbf{s}, \mathbf{t})$ and then calculate $\sum_{j=1}^n \alpha^{(j)} / \beta^{(j)}$. **MCMC** provides the algorithms to generate the random sample from the posterior distribution – or more generally, from any distribution, $\pi(\cdot)$, of interest. The strategy is to construct a Markov chain with $\pi(\cdot)$ as its stationary distribution.

A Markov chain is a sequence of random variables $\mathbf{X}^{(j)}$, $j = 0, \dots$, so that for each j , $\mathbf{X}^{(j)}$ is sampled from a distribution $P(\mathbf{X}^{(j)} | \mathbf{X}^{(j-1)})$. Given $\mathbf{X}^{(j-1)}$, the next state $\mathbf{X}^{(j)}$ depends only on $\mathbf{X}^{(j-1)}$ and not on the rest of the history of the chain. The sequence is called a *Markov chain*, and $P(\cdot | \cdot)$ is called the *transition kernel* of the chain. Subject to certain regularity conditions [4], the chain, regardless of its starting value $\mathbf{X}^{(0)}$, eventually converges (after a *burn-in* of m samples) to a unique stationary distribution. As j increases, the samples $\mathbf{X}^{(j)}$ become dependent samples approximately from $\pi(\cdot)$.

Constructing a Markov chain with the appropriate stationary distribution is straightforward using the *Metropolis-Hastings algorithm* [5]. Suppose that we want to draw $\mathbf{X}^{(j+1)}$. We sample a *candidate* point $\mathbf{Y}$ from a *proposal distribution* $q(\cdot | \mathbf{X}^{(j)})$. The proposal distribution may depend on the current point $\mathbf{X}^{(j)}$. The candidate point is accepted (i.e., $\mathbf{X}^{(j+1)} = \mathbf{Y}$) with probability $\alpha(\mathbf{X}^{(j)}, \mathbf{Y})$, where

$$\alpha(\mathbf{X}, \mathbf{Y}) = \min\left(1, \frac{\pi(\mathbf{Y})q(\mathbf{X} | \mathbf{Y})}{\pi(\mathbf{X})q(\mathbf{Y} | \mathbf{X})}\right) \quad (10)$$

Otherwise, $\mathbf{X}^{(j+1)} = \mathbf{X}^{(j)}$. The rate of convergence depends strongly on the relationship of the stationary distribution and the proposal distribution. See [4] for strategies on choosing good proposal distributions.

There are canonical choices for proposal distributions. The *Metropolis algorithm* uses symmetric proposals, where $q(\mathbf{Y} | \mathbf{X}) = q(\mathbf{X} | \mathbf{Y})$. A special case of the Metropolis algorithm is *random-walk Metropolis*, where $q(\mathbf{Y} | \mathbf{X}) = q(|\mathbf{X} - \mathbf{Y}|)$. The *independence sampler* [6] has a proposal density that does not depend on the current state of the chain, $q(\mathbf{Y} | \mathbf{X}) = q(\mathbf{Y})$. Independence samplers work well when $q(\cdot)$ is a good approximation to $\pi(\cdot)$ but with heavier tails.

Two of the most widely used choices for proposal distributions are the *single-component Metropolis-Hastings algorithm* and the *Gibbs sampler*. Instead of updating the entire vector $\mathbf{X}$ in one step, these algorithms update smaller groups of the components of $\mathbf{X}$ – often only one component at a time. Let $\mathbf{X} =$

$(X_1, \dots, X_p)$ and let $\mathbf{X}_{-i} = (X_1, \dots, X_{i-1}, X_{i+1}, \dots, X_p)$. Define the *full conditional distribution* of X_i as the distribution of the i th component of $\mathbf{X}$ conditional on all of the remaining components, $\mathbf{X}_{-i}$. If $\mathbf{X} \sim \pi(\cdot)$, the full conditional distribution of X_i is

$$\pi(X_i | \mathbf{X}_{-i}) = \frac{\pi(\mathbf{X})}{\int \pi(\mathbf{X}) dX_i} \quad (11)$$

The single-component Metropolis-Hastings algorithm generates candidate points Y_i from proposal densities $q_i(Y_i, X_i^{(j)}, \mathbf{X}_{-i}^{(j)})$, where $\mathbf{X}_{-i}^{(j)} = (X_1^{(j)}, \dots, X_{i-1}^{(j)}, X_{i+1}^{(j-1)}, \dots, X_p^{(j-1)})$, so that components with indices smaller than i have already been updated j times. The acceptance probability for Y_i is

$$\alpha(X_{-i}, X_i, Y_i) = \min \left(1, \frac{\pi(Y_i | \mathbf{X}_{-i}) q_i(X_i | Y_i, \mathbf{X}_{-i})}{\pi(X_i | \mathbf{X}_{-i}) q_i(Y_i | X_i, \mathbf{X}_{-i})} \right) \quad (12)$$

The Gibbs sampler [7] is a special case of the single-component Metropolis-Hastings algorithm where the proposal distribution is the full conditional distribution,

$$q_i(Y_i, X_i^{(j)}, \mathbf{X}_{-i}^{(j)}) = \pi(Y_i | \mathbf{X}_{-i}) \quad (13)$$

The acceptance probabilities for the Gibbs sampler are all one.

Application to Bayesian System Reliability

For illustration, suppose that we have a simple parallel system (*see Parallel, Series, and Series-Parallel Systems*) that comprises the 10 pumps described in Table 1. The system works if at least one of the pumps is working. We would like to estimate the probability that the system is working at 10 000 h. (Remember that the times in Table 1 are in thousands of hours.) Because we describe the number

of failures as a Poisson process (*see Poisson Processes*), $\text{Poisson}(\lambda_i t_i)$, the time to the first failure of a pump is distributed $\text{Exponential}(\lambda_i)$. Assume for this example that once a pump fails, it is not repaired. (Of course, in a real nuclear power plant, pumps are always repaired.) We would like to estimate the probability, $R_S(t)$, that the system is working at 10 000 h. The probability that the system is working at 10 000 h is $R_S(10) = 1 - \prod_{i=1}^{10} (1 - \exp(-10\lambda_i))$. (See [8] for the derivation of the expression for system reliability.)

We use a combination of the Gibbs sampler and single-component Metropolis-Hastings algorithm to estimate the failure rates for each pump. This requires the full conditional distributions for each parameter.

$$[\lambda_i | \lambda_{-i}, \alpha, \beta] \sim \text{Gamma}(s_i + \alpha, \beta + t_i) \quad (14)$$

$$[\beta | \alpha, \lambda_1, \dots, \lambda_{10}] \sim \text{Gamma} \left(10\alpha + 0.1, 1.0 + \sum_{i=1}^{10} \lambda_i \right) \quad (15)$$

$$[\alpha | \beta, \lambda_1, \dots, \lambda_{10}] \propto \left(\prod_{i=1}^{10} \lambda_i^{s_i + \alpha - 1} \right) \beta^{10\alpha - 0.9} (\Gamma(\alpha))^{-10} \times I(\alpha \in (0, 5)) \quad (16)$$

For the $(j+1)$ st iteration of MCMC, for $i = 1, \dots, 10$, we set $\lambda_i^{(j+1)}$ equal to a random draw from a $\text{Gamma}(s_i + \alpha^{(j)}, \beta^{(j)} + t_i)$ distribution, and we set $\beta^{(j+1)}$ equal to a random draw from a $\text{Gamma}(10\alpha^{(j)} + 0.1, 1.0 + \sum_{i=1}^{10} \lambda_i^{(j+1)})$. To sample $\alpha^{(j+1)}$, we set a candidate point Y equal to a random draw from a proposal distribution: for example, $\text{Normal}(\alpha^{(j)}, 0.25)I(0 < \alpha < 5)$. As a general heuristic, we choose the standard deviation of the proposal distribution so that the candidate acceptance probability is between 0.25 and 0.45 [9]. For this proposal distribution, the candidate acceptance probability is

$$\min \left(1, \frac{\prod_{i=1}^{10} (\lambda_i^{(j+1)})^{s_i + Y - 1} (\beta^{(j+1)})^{10Y - 0.9} \Gamma(Y)^{-10}}{\prod_{i=1}^{10} (\lambda_i^{(j+1)})^{s_i + \alpha^{(j)} - 1} (\beta^{(j+1)})^{10\alpha^{(j)} - 0.9} \Gamma(\alpha^{(j)})^{-10}} \right) \quad (17)$$

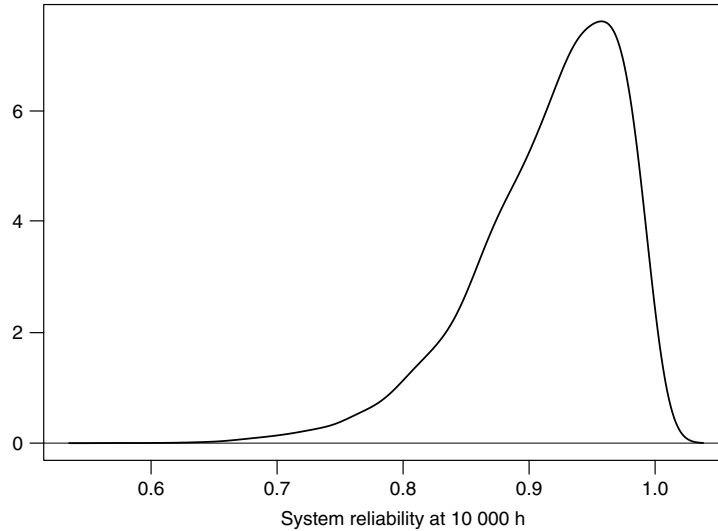


Figure 1 The posterior distribution of $R_S(10)$, the system reliability at 10 000 h

We could also use a MCMC package like WinBUGS [10] to perform the calculations.

Once we have a sample from the posterior distribution, we evaluate $R_S(10)$ for each draw. The posterior mean of $R_S(10)$ is 0.92, with a 95% credible interval of (0.77, 0.99). Figure 1 is a plot of the posterior distribution of the system reliability at 10 000 h. (For more information on creating summaries using MCMC draws, see [11].)

Summary

Hierarchical models offer many advantages including the ability to borrow strength, to estimate individual parameters and the ability to specify complex models that reflect engineering and physical realities. MCMC provides a general set of algorithms that enable Bayesian inference in a variety of complex models.

References

- [1] Draper, D., Gaver, D., Goel, P., Greenhouse, J., Hedges, L., Morris, C., Tucker, J. & Waternaux, C. (1992). *Combining Information: Statistical Issues and Opportunities for Research*, National Academies Press, Washington, DC.
- [2] Gaver, D. & O'Muircheartaigh, I. (1987). Empirical Bayes analyses of event rates, *Technometrics* **29**, 1–15.
- [3] Gelfand, A. & Smith, A. (1990). Sampling-based approaches to calculating marginal densities, *Journal of the American Statistical Association* **85**, 398–409.
- [4] Gilks, W., Richardson, S. & Spiegelhalter, D. (1996). *Markov Chain Monte Carlo in Practice*, Chapman & Hall, London.
- [5] Chib, S. & Greenberg, E. (1995). Understanding the Metropolis-Hastings algorithm, *The American Statistician* **49**, 327–335.
- [6] Tierney, L. (1994). Markov chains for exploring posterior distributions (with discussion), *Annals of Statistics* **22**, 1701–1762.
- [7] Casella, G. & George, E. (1992). Explaining the Gibbs sampler, *The American Statistician* **46**, 167–174.
- [8] Rausand, M. & Høyland, A. (2003). *System Reliability Theory: Models, Statistical Methods, and Applications*, 2nd Edition, John Wiley & Sons, Hoboken.
- [9] Gelman, A., Carlin, J., Stern, H. & Rubin, D. (1995). *Bayesian Data Analysis*, Chapman & Hall, Boca Raton.
- [10] The BUGS Project – WinBUGS. (accessed 26 June 2006). MRC biostatistics unit, Cambridge. <http://www.mrc-bsu.cam.ac.uk/bugs/winbugs/contents.shtml>.
- [11] Smith, A. & Roberts, C. (1993). Bayesian computations via the Gibbs sampler and related Markov chain Monte Carlo methods, *Journal of the Royal Statistical Society. Series B* **55**, 3–23.

Related Article

Bayesian Reliability Analysis; Expert Opinion in Reliability; Fault Trees; Parallel, Series, and Series-Parallel Systems; Poisson Processes.

ALYSON G. WILSON

Life Cycle Costs and Reliability Engineering

Introduction

Life cycle cost (LCC) is the total cost of a system or product over its entire life span. It includes the cost of design, production, operation, maintenance, warranty, and decommission or disposal. This cost is extremely critical as it is one of the most critical factors in decision-making regarding the procurement of a product or system. It is usually calculated using a simple payback measure such as the present or future value of the cost or using complicated analysis such as the analytical hierarchy process and others for comparing alternatives, especially when some of the cost components are intangibles.

LCC analysis became common in 1960s when government agencies and others used it as a criterion in procurement decisions. Since then it has been widely used in industry, business, and government agencies in many applications ranging from simple office equipment to complex engineering structures such as large-scale communications networks. Most of the cost components are common for these applications while others are not. For example, nonrepairable products such as satellites may not include the cost of repair and perhaps the cost of disposal. The first of these costs is a reliability-related cost. Owing to the fact that LCC can indeed encompass many cost components and it is difficult, if not impossible, to discuss all of them, this chapter focuses on LCC elements related to reliability engineering. The use of LCC analysis in reliability engineering is widely used in the automotive industry, which prompted the Society of Automotive Engineering (SAE) to advocate the reduction of LCC for equipment in the automotive sector by demonstrating why reliability and maintainability must be included in up-front decisions for strategic and tactical issues of achieving the lowest long-term cost on ownership as discussed in [1, 2]. In the following sections we intend to address the issues related to the LCC when considering reliability engineering in product or system design and procurement.

Products and systems consist of components arranged or configured in a specific arrangement that produces the desired functions. Therefore, the reliability of the components has a profound effect on

the overall system reliability, life, and LCC. The trade-off between the reliability of the component and its cost must be considered for proper selection of components and while ensuring their reliability. Likewise, once components are selected, they need to be configured in a specific arrangement among many alternatives in order to achieve an optimum system design that meets the requirements of the specification at minimum cost. The cost should of course take into account the cost of the components, maintenance, replacements, operation, and warranty over the expected life of the product. Once an optimum design is achieved, the product goes through the production process and is tested in order to ensure that it will meet its reliability objectives. Normally, all or a subset of accelerated testing, reliability demonstration testing, and **burn-in** testing are performed. The cost of testing could be detrimental to the launching of the product. Indeed, in some situations the cost of testing could exceed the total budget allocation for the product development, which might cause the termination of product entirely. The next step is to ensure that the product will function properly without failures or recalls for a specified period during which inspection and maintenance are implemented. The cost of spares at that time could be higher than the cost of replacement of the product with a new one. During this operation period of the product, **warranty policies** play a major role in extending or shortening the life of the product ownership.

As discussed, there are several phases in the product life cycle where reliability has a major impact on its cost. Therefore, we present these phases with a focus on the reliability and LCC during every phase. They are described subsequently.

LCC during Design Phase

Once a concept design of the product is developed, components are acquired and arranged or configured to meet the product functions. Selection of components has a major impact on the LCC of the product. For example, designing a nonrepairable product requires the selection of highly reliable components produced at stringent requirements and made from “prime” material to ensure that the components will not fail during the intended mission of the product. Such products include undersea fiber-optic cables where the product (cable) should not fail during an expected life of 70–80 years. Clearly, the failure rate

2 Life Cycle Costs and Reliability Engineering

of such components should be constant with very small values in the order of 1.0 FIT (FIT is a measure of failure rate and defined as one failure in 10^9 h). The cost and reliability of such components are relatively high, which results in lower operation cost (no failures and no down time) and the corresponding LCC values justify the use of the components. In this case, the cost of downtime and expected life without failure have indeed dictated the reliability of the components.

When the product or the system is repairable and the cost of down time and repairs are “acceptable”, the reliability of the components need not be as high as in nonrepairable systems. Therefore, the initial cost of the product will be significantly less than that of the nonrepairable product, but its operational costs are higher. The failure rates and the failure rate functions of the components are obtained from historical information or from extensive testing of the components at different operating conditions. Some industries and military standards provide information about the failure rates and design for LCC, see references [3–6].

The system configuration has a direct impact on its LCC. When the system is configured as a series system, i.e., the components are connected in series and the system fails when any of the components fails, then the operation and maintenance cost become more significant over the life of the system (product). Therefore, the total system cost of the system should be optimized to determine the optimum reliability levels of the components. The total cost is composed of the cost of the components (higher reliability components are more costly but decrease the downtime of the system), cost of repair, and downtime during the expected life of the system. Figure 1 shows the relationship between the **component reliability** level and the overall cost of the system. As shown, there is an optimum level that minimizes the overall cost [7].

Improving the series configuration might necessitate the introduction of **redundancy** of the components. When the system reliability is the main objective of the introduction of redundancy, then redundancy should be provided for the “weakest” component (least-reliable component) in the system. This can be achieved by having identical units or nonidentical units depending on the reliability of the component and the associated cost. This becomes an

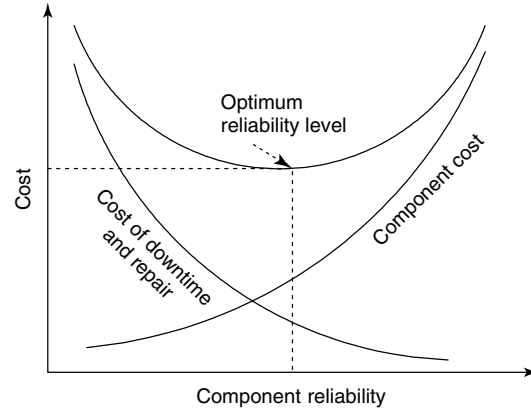


Figure 1 Optimum reliability level

optimization problem with two objectives: maximization of the system reliability and minimization of the LCC, which should incorporate the maintenance cost (inspection, diagnostics, and repair) and the downtime cost. As the system becomes complex, the alternatives for system configuration increase and the computation of LCC is no longer simple. A simple formulation of the optimization problem is given as

$$\begin{aligned} & \text{Max } R(t) \\ & \text{s.t. } \sum_{i=1}^s c_i < C \\ & \sum_{i=1}^s w_i < W \end{aligned}$$

where $R(t)$ is the reliability function, c_i and w_i are the unit cost and unit weight, respectively, and C and W are the total cost and total weight of the system, respectively.

Other designs of the system configurations include **parallel systems**, series–parallel systems, parallel–series systems, and complex network configurations. Optimization of the design of parallel systems requires the determination of the optimum number of parallel components (each has its own reliability level or failure rate function and cost). The associated downtime, repairs, and warranty become difficult to calculate. Approaches have been developed to determine the optimum system configuration that maximizes the reliability (considered as a constraint in

the optimization problem) while minimizing the LCC, see for example [8–10].

In addition to the design of system configuration, designers are always faced with the problem of improving system reliability with minimum increases in cost. Sometimes, system design requires redundancies such as in power generation and number of engines of an aircraft. The latter is dictated by the aviation regulations, which states that no commercial aircraft should be designed with only a single engine. In such cases, both reliability and cost become an integral part of the system design and yet compete with each other.

Optimum allocation of components to the system configuration is indeed a challenging problem that does not have an optimum solution for large-scale systems but has a profound effect on the LCC. Heuristics for determining “good” solutions are developed using **genetic algorithms** and other meta-algorithms. One of the final tasks to be performed during the design phase is the reliability demonstration testing and reliability prediction testing. The costs associated with the test as well as its duration will impact the reliability estimate of the system and consequently its LCC. Longer test times will reveal potential failures of the system and will result in system improvement but it increases its LCC. Therefore, such reliability testing should be optimized with the objective of obtaining the most accurate reliability estimates at a minimum cost. A typical formulation of the test optimization problem is given in [11].

LCC during Production Phase

Once the product (system) is designed and its production begins, its reliability must be demonstrated throughout the production stage. A reliability demonstration test is performed to ensure that the product meets its reliability objectives as it is being manufactured. The test is also used to uncover unexpected failure modes or design errors. It is important that extensive testing be performed to avoid potentially expensive field failures which include all costs associated with product recalls, design fixes, and product repair. If a product failure results in personal injuries or catastrophic failures then the cost of settling lawsuits and litigation costs should also be included in the LCC. More importantly, customer awareness of the product reliability might indeed be detrimental

to the product market share. If any potential failure modes do escape to the field, the presence of a product support team, which includes failure analysis specialists, should be mobilized for rapid containment of any problem and removal of root causes responsible for reliability problems [12]. Clearly these have a significant impact on the system LCC.

There are two additional types of testing to be conducted before shipping the product; they are highly accelerated life testing (HALT) and burn-in testing. HALT is performed so that a product and the process that produces it are mature at the time of the product introduction (minimal, if any, **reliability growth** remains to be achieved), while design and development time are kept to a minimum [13]. HALT requires the application of two types of stresses individually, then simultaneously; they are thermal stresses and vibration step stresses. The stress levels are high in order to “discover” any potential latent defects or unexpected failure modes. The duration of the test and number of test samples add significant cost to the product development but might result in the reduction of potential field failures. This cost indeed adds to the overall LCC and the trade-off between the gain in product reliability due to this testing and the cost of the test as well as the delay in production need to be considered.

The second type of test is burn-in testing. This test is performed in order to “weed out” defective products due to poor assembly. The test is performed at stresses slightly higher than operating conditions so that the early failure rate region (infant mortality region) of the product is substantially reduced. In doing so, the repair cost as well as the recall cost of the product at its early operating hours is minimized. The duration of the test varies among manufacturers and products. Again, similar to HALT, the duration of the test has a major impact on the product life and LCC. Therefore, reliability engineers investigated different approaches for determining the optimum burn-in time. A simple model for estimating the optimum test time is described as follows.

Consider a product that has a Weibull failure time distribution β and θ as the shape and scale parameters of the distribution. Assume that we are interested in determining the optimum burn-in duration t . This can be obtained by expressing the total system cost as a function of the burn-in duration, warranty period, and field failures (if the test duration is too short, the

4 Life Cycle Costs and Reliability Engineering

system will exhibit high field failures). We define the following [14]:

- t : burn-in duration
- w : warranty period
- c_1 : burn-in cost per unit time
- c_2 : cost of a failure in the field during the warranty period
- c_3 : cost of a failure during the burn-in period
- β : the shape parameter of the Weibull distribution that represents the time to fail for the population
- θ : the scale parameter of the Weibull distribution that represents the time to fail for the population

The objective is to determine the optimum burn-in test period that minimizes the total system cost, which consists of the following three cost elements.

The system cost is the total of

1. the cost of conducting a burn-in test (tc_1);
2. the cost of units failed during burn-in (the probability of failing during burn-in, c_3); and
3. the cost of units failing during the warranty period in (the probability of failing during the warranty period, c_2).

Since the product's failure time follows a Weibull model with a **probability density function** of $f(t) = \beta/\theta (t/\theta)^{\beta-1} e^{-(t/\theta)^\beta}$, the reliability function is

$$R(t) = e^{-\left(\frac{t}{\theta}\right)^\beta} \quad (1)$$

The total system cost is

$$c_s = tc_1 + \left[1 - e^{-\left(\frac{t}{\theta}\right)^\beta} \right] c_3 + \left[1 - \frac{e^{-\left(\frac{t+w}{\theta}\right)^\beta}}{e^{-\left(\frac{t}{\theta}\right)^\beta}} \right] c_2 \quad (2)$$

The optimum burn-in duration is found by minimizing equation (2) in terms of t . To illustrate how the optimum burn-in duration is obtained, we use the example given in [14].

Example. Assume that the initial testing of the product shows that its failure time distribution is Weibull with a shape parameter of 0.8 and a scale parameter of 6000 days. The warranty period is 1000 days, and the average cost of a **warranty claim** is \$600. The cost of burn-in is \$0.09 per hour per product, and the cost of a failure during burn-in is \$4.00. Determine the optimum burn-in duration.

Solution. We use equation (2) after substituting the above values in the corresponding parameters. Thus, the system cost is expressed as

$$c_s = 0.09t + 4 \left[1 - e^{-\left(\frac{t}{6000}\right)^{0.8}} \right] + 600 \left[1 - \frac{e^{-\left(\frac{t+1000}{6000}\right)^{0.8}}}{e^{-\left(\frac{t}{6000}\right)^{0.8}}} \right] \quad (3)$$

Plotting the total cost of the system *versus* time results in Figure 2 with an optimum burn-in test duration of 30 h.

Extensive testing as described earlier during this phase of the product life cycle should reduce the overall cost of the product, and trade-off between the test cost and the field failure and repair cost must be carefully performed.

LCC during Operation Phase

The cost of the product during its operation phase includes the operating cost and more importantly its

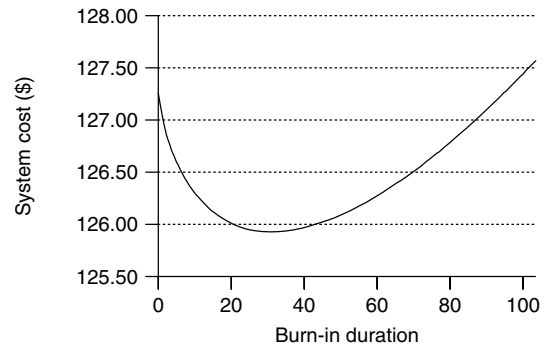


Figure 2 Optimum burn-in duration

maintenance and repair cost. The latter cost requires the determination of optimum **preventive maintenance**, replacements, and inspection (PMRI) schedules. The term *optimum* arises from the fact that high frequency of PMRI increases the total cost of maintenance and reduces the cost due to the downtime of the system, whereas low frequency of PMRI decreases the cost of maintenance but increases the cost due to the downtime of the system. Hence, depending on the type of failure time distribution, an optimum PMRI may exist. Preventive maintenance may imply minimal repairs, replacements, or inspection of the components. Obviously there are systems where minimal repairs or inspections are not applicable, such as a microprocessor of a programmable logical controller. Similarly, the status of some systems can only be determined by inspection or testing—as in “one-shot” devices, such as military explosives. In this case, replacement is the only possible alternative.

The primary function of the preventive maintenance and inspections is to control the condition of the equipment and ensure its availability. Doing so requires the determination of the following:

- frequency of the PMRI,
- replacement rules for components,
- effect of technological changes on the replacement decisions,
- the size of the maintenance crew,
- optimum inventory levels of spare parts,
- sequencing and scheduling rules for maintenance jobs, and
- number and type of machines available in the maintenance workshop.

The above topics are a partial list of what constitutes a comprehensive PMRI [15]. Indeed the comprehensive maintenance during the operation phase of the product may result in added cost that exceeds

the initial acquisition cost of the product, which in turn increases the LCC of the product. The optimum maintenance (preventive or replacement) schedule is determined by minimizing the total maintenance cost of the product during its expected life (or cycle of operation) as demonstrated below.

The constant-interval replacement policy (CIRP) is the simplest preventive maintenance and replacement policy. Under this policy, two types of actions are performed. The first type is the preventive replacement that occurs at fixed intervals of time. Components or parts are replaced at predetermined times regardless of the age of the component or the part being replaced. The second type of action is the failure replacement where components or parts are replaced upon failure. This policy is illustrated in Figure 3 and is also referred to as *block-replacement policy*.

As mentioned earlier, the objective of the PMRI models is to determine the parameters of the preventive maintenance policy that optimize some criterion. The most widely used criterion is the total expected replacement cost per unit time. This can be accomplished by developing a total expected cost function per unit time as follows.

Let $c(t_p)$ be the total replacement cost per unit time as a function of t_p . Then

$$c(t_p) = \frac{\text{Total expected cost interval } (0, t_p]}{\text{Expected length of the interval}} \quad (4)$$

The total expected cost in the interval $(0, t_p]$ is the sum of the expected cost of failure replacements and the cost of the preventive replacement. During the interval $(0, t_p]$, one preventive replacement is performed at a cost of c_p and $M(t_p)$ failure replacements at a cost of c_f each, where $M(t_p)$ is the expected number of replacements (or renewals) during the interval $(0, t_p]$. The expected length of the

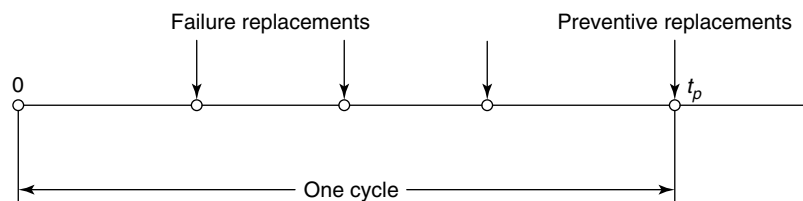


Figure 3 Constant-interval replacement policy

interval is t_p . Equation (4) can be rewritten as

$$c(t_p) = \frac{c_p + c_f M(t_p)}{t_p} \quad (5)$$

The expected number of failures, $M(t_p)$, during $(0, t_p]$ may be obtained using any of the methods discussed in [15].

Other maintenance policies that minimize the total cost over the life of the product are found in [15]. It should be noted that the above policy may include the cost of the spares, logistics, and others. However, these cost elements can also be included separately.

LCC and Warranty

A warranty is a contract or an agreement under which the manufacturer of a product or service must agree to repair, replace, or provide service when the product fails or the service does not meet the customer's requirements before a specified time (length of warranty). This specified "time" may be measured in calendar time units such as hours, months, and years, or in usage units such as miles, hours of operation, number of times the product has been used (number of copies made by a copier, number of pages printed by a printer etc.) or both. Other warranties have no specified "times" and are referred to as *lifetime warranties*.

Three types of warranties are commonly used for consumer goods: the *ordinary free replacement* warranty, the *unlimited free replacement* warranty, and the *pro rata* warranty. Under an ordinary free replacement warranty, if an item fails before the end of the warranty period, it is replaced or repaired at no cost to the consumer. The repaired or replaced item is then covered by an ordinary free replacement warranty with a period equal to the remaining period of the original warranty. Such warranty assures that the consumer will receive as many free repairs or replacements as needed during the original period of the warranty. The ordinary free replacement warranty is the most common type of warranty; it is most often used to cover consumer durables such as cars and kitchen appliances [16].

The second type of warranty, the unlimited free replacement, is identical to the ordinary replacement warranty except that each replacement item carries an identical warranty to the original purchase warranty. Such warranties are only used for small electronic

appliances that have high early failure rates and are usually limited to very short periods of time.

Thus, under the free replacement warranty, whether it is ordinary or unlimited, a long warranty period will result in a very large **warranty cost**. Furthermore, an increase in the warranty period will reduce the number of replacement purchases over the life cycle of the product, which consequently reduces the total profit of the manufacturer. Clearly, the free replacement policy is more beneficial to the consumer than to the manufacturer. Therefore, it is extremely important for the manufacturer to determine the optimal price of the product and the optimal warranty period such that the total cost over the product life cycle is minimized.

The third type of warranty is the *pro rata* warranty. Under this warranty, if the product fails before the end of the warranty period, it is replaced at a cost that depends on the age of the item at the time of failure and the replacement item is covered by an identical warranty. Typically, a discount proportional to the remaining period of the warranty is given on the purchase price of the replacement item. For example, if the period of the warranty is w and the item fails at a time $t < w$, the consumer pays the proportion t/w of the cost of the replacement items (automobile tires represent an ideal product for which the *pro rata* warranty policy is appropriate), and the manufacturer covers the remaining cost of the product replacement. Unlike the free replacement warranty, the *pro rata* warranty is more beneficial to the manufacturer than to the consumer.

As presented above, a pure free replacement warranty favors the consumer whereas a pure *pro rata* warranty favors the manufacturer. Therefore, a mix of these policies may present an alternative warranty, which is fair to both the consumer and manufacturer. For example, a policy that involves an initial free replacement period followed by a *pro rata* period is a sensible one, being fair to both the manufacturer and the consumer [17]. It has a promotional appeal to attract consumers and at the same time keeps the warranty cost for the manufacturer within a reasonable amount. There are other types of warranty policies such as reliability improvement warranty where the manufacturer provides guaranteed mean time between failures (MTBF) or provides support for engineering [18].

Clearly, warranty cost is an added cost to the product life cycle and it should be minimized by

improving the reliability of the product during all phases of its life cycle. Recent effort by one of the leading automakers in reducing its warranty cost and extending the warranty period resulted in a significant increase of its sales.

References

- [1] *Society of Automotive Engineers, Reliability, Maintainability, and Supportability Guidebook*, SAE, 1995.
- [2] Barringer, H.P. (2003). A life cycle cost summary, in *International Conference of Maintenance Societies (ICMS-2003)*, Perth, May 2003.
- [3] Bellcore SR-332 (2001). *Reliability Prediction Procedure for Electronic Equipment*, Telcordia.
- [4] MIL-HDBK-259 (1983). *Military Handbook, Life Cycle Cost in Navy Acquisitions*, Global Engineering Documents.
- [5] MIL-HDBK-276-1 (1984). *Military Handbook, Life Cycle Cost Model for Defense Material Systems*, Data Collection Workbook, Global Engineering Documents.
- [6] MIL-HDBK-276-2 (1984). *Military Handbook, Life Cycle Cost Model for Defense Material Systems Operating Instructions*, Global Engineering Documents.
- [7] Wang, H. & Pham, H. (2006). Availability and maintenance of series systems subject to imperfect repair and correlated failure and repair, *European Journal of Operational Research* **174**, 1706–1722.
- [8] Chiang, C.-H. & Chen, L.-H. (2007). Availability allocation and multi-objective optimization for parallel-series systems, *European Journal of Operational Research* **180**, 1231–1244.
- [9] Bris, R., Châtelet, E. & Yalaoui, F. (2003). New method to minimize the preventive maintenance cost of series-parallel systems, *Reliability Engineering and System Safety* **82**, 247–255.
- [10] Kuo, W., Prasad, V.R., Tillman, F. & Hwang, C.-L. (2001). *Optimal Reliability Design: Fundamentals and Applications*, Cambridge University Press, New York.
- [11] Elsayed, E.A. & Zhang, H. (2007). Design of optimum simple step-stress accelerated life testing plans, in *Recent Advancement of Stochastic Operations Research*, T. Dohi, S. Osaki & K. Sawaki, eds, World Scientific, Singapore.
- [12] Wasserman, G.S. (2003). *Reliability Verification, Testing, and Analysis in Engineering Design*, Marcel Dekker, New York.
- [13] McLean, H.W. (2000). *HALT, HASS & HASA Explained: Accelerated Reliability Techniques*, ASQ Quality Press, Milwaukee.
- [14] Engineered Software, Inc. <http://www.engineeredsoftware.com/burn.in.asp>. Belleville, 2007.
- [15] Elsayed, E.A. (1996). *Reliability Engineering*, Addison-Wesley, Reading.
- [16] Mamer, J.W. (1987). Discounted and per unit costs of product warranty, *Management Science* **33**, 916–930.
- [17] Nguyen, D.G. & Murthy, D.N.P. (1984). Cost analysis of warranty policies, *Naval Research Logistics Quarterly* **31**, 525–541.
- [18] Blischke, W.R. & Murthy, D.N.P. (1994). *Warranty Cost Analysis*, Marcel Dekker, New York.

Further Reading

MIL-HDBK-217 (1995). *Reliability Prediction of Electronic Equipment*, Global Engineering Documents.

Related Articles

Accelerated Life Tests: Bayesian Design; Accelerated Life Tests: Classical Methods for Design and Analysis; Accelerated Life Tests: Designs Comparison within a Bayesian Framework; Analysis of Recurrent Events from Repairable Systems; Block Replacement; Burn-In and Maintenance Policies; General Minimal Repair Models; Inspection Policies for Reliability; Parallel, Series, and Series-Parallel Systems; Reliability Allocation; Reliability Growth Modeling; Reliability of Redundant Systems; Reliability Optimization; Repairable Systems Reliability; Replacement Strategies; Total Productive Maintenance; Warranty Cost Analysis; Warranty Cost Prediction Based on Warranty Data.

ELSAYED A. ELSAYED

Warranty Claims and Costs: Statistical Analysis of

Introduction

Warranty is inherently business related, but the topic, and particularly the data arising from warranties, has generated tremendous interest in the academic literature. This article will refer mainly to books and survey papers, but interested readers should consult the references therein for additional details.

The type of data analysis performed in practice depends on the business purpose and also on the nature of the available data. The data itself generally exists to serve some business purpose, not necessarily identical to the purpose of the proposed analysis. For example a listing of warranty claims for repairs performed for customers may exist primarily to allow the manufacturer to pay dealers for their service, but the same data may be used to estimate the reliability of the product in the field or to estimate the extent of future claim payments. Also, the available data depends on the nature of the warranty coverage. For example, repairs made outside of the warranty coverage period may not be available to the manufacturer.

A warranty has been described as a contract between the buyer and the manufacturer, and it protects both parties. Murthy and Djameludin [1] offer this definition as well as theories for the study of warranty including the theory that warranties provide a signal of reliability to the customer and also the investment theory that a warranty is in effect an insurance policy for the customer. Robinson and McDonald [2] list business facets impacted by warranty including warranty as a product attribute, which seems especially important today. Buyers may value generous warranties in addition to attractive features, and both of these will usually add cost for the producer. Sellers of relatively unknown brands may reduce the uncertainty of a prospective buyer with a generous warranty.

Two basic types of warranty coverage are free replacement and *prorata*. Free repair is similar to free replacement, but may dictate different modeling assumptions for a complex system such as an automobile. Some warranty coverages are two dimensional,

as with age and mileage for automobiles. Warranties between two businesses or those involving government contracts often have specialized features to encourage the seller to improve reliability. See Blischke and Murthy [3] for extensive discussions on these topics.

Regardless of the type of coverage the warranty database will generally have uses beyond its basic purpose as a financial record. The database usually includes a list of records of repairs or other services performed for the customer by the manufacturer or on its behalf. These records constitute a “claim” against the manufacturer. The potential for new claims, either from units in the field or those yet to be sold, represents a future liability. Estimating the extent of the future liability is a business requirement.

The available data may include much more than just a list of claims. Claims may be linked in the database to sales records, customer information, production records, information from surveys, etc. The extent and detail of this auxiliary information will determine the options for statistical analysis. For example, if detailed sales records are available, then the time in the field or current “age” for every unit is known. If the sales records can be linked to the claims records, then the age of any unit at the time of a claim is also known, permitting what has been called *age-based* warranty analysis. Model and option information from sales records as well as production information may allow stratifying the analysis to enable comparisons across models, plants, or suppliers. Other data limitations and specialized analyses are discussed in the last section.

Statistical Models

The age or time of occurrence of warranty claims can be considered a type of reliability data, and many reliability models can be applied to warranty data. One approach is to regard warranty claims as events from a **repairable system**, where counting processes play a prominent role. Let $N_i(t)$ be the number of claims that have occurred through time t on unit i , and let these be modeled by a counting process. A fundamental quantity, both in theory and for practical work, is the *mean cumulative function*, $\Lambda(t) = EN_i(t)$. Means are important in practice for warranty analysis because totals are important. We may prefer nonparametric approaches for estimating the mean

2 Warranty Claims and Costs: Statistical Analysis of

cumulative function, especially if the number of units is large.

Nonhomogeneous **Poisson processes** form a class of models that would naturally apply if the claim corresponds to a “**minimal repair**”, i.e., the unit is assumed to be restored to its state prior to failure. If the claim corresponds to a repair that is assumed to restore the unit to its original state or if the claim corresponds to a replacement of the unit, then the class of **renewal processes** may apply. Note that replacement of a failed component in the unit is not the same as replacement of the unit itself. The former is essentially a repair and would not ordinarily be expected to restore the unit to its original condition.

If the number of claims for a unit is nearly always zero or one, then methods from lifetime data analysis are applicable, either parametric or nonparametric. The lifetime distribution of the time to first repair or replacement is of fundamental interest. See Meeker and Escobar [4], Rigdon and Basu [5], Lawless [6], Kalbfleisch and Prentice [7], and other general references on statistical models for reliability and lifetime data for detailed methods.

One feature that distinguishes warranty data in particular from reliability data in general is the presence of costs, which usually occur along with the time of the claim in a warranty database. This is a complication, but costs can often be modeled separately from the rate of accumulation of claims. If so, then simply note that for a given unit, cost equals number of claims times average cost per claim. Then separate these two elements in the analysis. In addition to the claims arrival process, also model the distribution of costs per claim, which may require some degree of stratification. Another approach is to view the cost as a “mark” associated with the claim and apply the class of marked counting-process models. Finally, nonparametric approaches to estimating the mean cumulative function work equally well for amounts (costs) as they do for counts (number of claims), which eliminates the need for a separate model for costs. See Nelson [8] for details on the nonparametric approach.

Statistical Analysis

As indicated in the last section many reliability models can be applied to warranty data. This section emphasizes issues and analysis approaches that relate

to some of the special characteristics of warranty data. We begin with a brief discussion of warranty data in the context of sampling and warranty claims as distinct from failures. This is followed by a more detailed automotive example to illustrate that warranty analysis should account for the coverage limitations of the **warranty policy**. The section concludes by discussing briefly a few of the other specific problems and approaches that arise with warranty data.

Not Samples and Not Failures

Warranty data do not usually constitute a random sample, so strictly speaking the assumptions required for, say, **confidence intervals** do not apply. However, in practice we may be willing to treat, say, the collection of first units sold as though it were a statistical sample representing all the units that will eventually be sold. But inference should proceed with caution and, if possible, use the data itself to check for evidence of **bias**.

Warranty claims do not necessarily correspond to failures. Manufacturers may cover some warranty repairs or replacements for customers as a goodwill gesture even though in the manufacturer’s view the unit has not failed. If the goal is to estimate warranty costs for example, then we might proceed with analyzing the data as it is with the understanding that claims represent not just failures but also other forms of dissatisfaction that may lead a customer to seek recourse under a warranty agreement. However, if the goal is to estimate the actual failure rate in the field, then some screening or further analysis of the claims may be necessary before submitting the data to an analysis based on methods from reliability theory.

Warranty Coverage Example: Automotive Data

As mentioned in the introduction, warranty data is limited by the nature of the warranty coverage. This is an extremely important feature that must be considered in order to properly analyze the data. Consider for illustration automotive warranty coverage. Since detailed sales information is available, it is reasonable to analyze the data by vehicle age, which is defined by the time that has passed since the vehicle was sold as new. But automotive warranties are typically limited by mileage as well as age, so the warranty data will not include repairs made on vehicles that have

exited warranty coverage due to the mileage limit, even though they are within the nominal age limit. This example is complicated further by the fact that, unlike ages, which are known all the time for all vehicles, mileages are only known for vehicles that generate a claim, and then only at the time of the claim because the odometer reading is only recorded at the time of the claim.

The issue illustrated here is similar to what occurs in epidemiology. We need to know or to estimate the number of units (people in epidemiology) that are at risk. The number at risk is the denominator in a claim rate calculation. In warranty analysis the denominators are often harder to obtain than the numerators, which typically come directly from the claims. Continuing with the automotive example, the number of claims for vehicles of age t can be obtained readily from the claims data. In an analysis that ignores mileages, the corresponding denominator would simply be the number of vehicles that are currently age t or older, which we could write as $\sum I(a_i \geq t)$, where a_i is the current age of vehicle i , $I(\cdot)$ is the indicator function, and the summation is over all vehicles under observation. Ignoring the complication of missing mileage information, a denominator adjusted for mileage coverage would be $\sum I(a_i \geq t)I(M_i(t) \leq L)$, where $M_i(t)$ is the accumulated mileage of vehicle i at age t , and L is the warranty coverage limit for mileage.

Figure 1 illustrates the mileage example. For the details see Chukova and Robinson [9]. The plot displays the estimates of the mean cumulative function of cost as a function of age, first ignoring mileage (bottom curve) and then adjusting the number of vehicles at risk for mileage accumulation (top curve). Note that both estimates may serve a purpose. The bottom curve displays how warranty expenses actually accrue as a function of age. The top curve estimates how they would accrue in the absence of a mileage limit and may be more appropriate for engineering analysis.

Other Special Types of Warranty Analysis

The previous example illustrates one of many special circumstances that can occur in the analysis of warranty data. Some of the others are discussed briefly in this subsection. For more details, see Karim and Suzuki [10] and the collection of papers in Blischke and Murthy [11]. Jerald Lawless and his coauthors

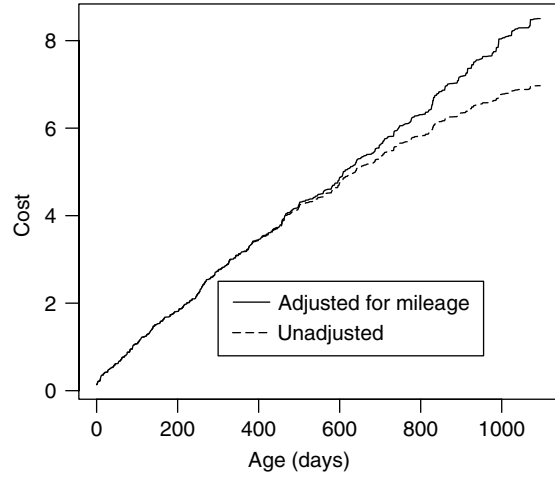


Figure 1 Mean cumulative function of cost by age with and without adjustments for mileage

have made a number of important contributions to warranty analysis; see Lawless [12] and the references therein. Chukova *et al.* [13] provide a survey of warranty analysis and references to work that has appeared in eastern European journals.

Reporting delays can influence warranty data and the resulting analysis. Neither claims nor sales information may be recorded in the database immediately. Solutions involve adjusting the number of units at risk using the observed distribution of delay times.

The automotive mileage example illustrated a case of dual usage measures, age and mileage, but the analysis described before was essentially univariate. However, fully bivariate approaches have been developed, particularly for the case where renewal processes apply. Also, the problem of incomplete mileage information (or any secondary usage measure) can be handled by using supplementary sample data. This would happen in the automotive example if mileage accumulation information were available, say from **survey data**. Not all warranty analysis is age based. Sometimes the data to support it may not exist because sales or warranty claim information is incomplete or cannot be cross-referenced. Nevertheless, some analysis methods have been developed for specific problems where the models are tailored to the specific data that are available. For example, Karim

and Suzuki [10] refer to “marginal counts” analysis, where the available information consists only of the number of units sold and the number of warranty claims by calendar month. This survey article contains a number of important references to articles on this topic and to the others discussed in this section.

Warranty analysis by covariates can be useful. The **covariate** effect may refer to the product, perhaps the model or vintage, or it may refer to the customer who bought the product, or where the item was purchased. In practice, covariate analysis is probably most often accomplished simply by doing separate analyses for the various cohorts of interest, but covariate effects can often be modeled rather easily, particularly with Poisson models.

Since parametric models allow one to extrapolate, they are natural choices to create forecasts of future warranty costs or claim rates. Of course, all the usual warnings associated with extrapolation apply. In practice there may be archives of warranty records for past products that are the same, or similar with respect to warranty trends as the current models. Using the past data can help offset the risk associated with extrapolation. This topic does not appear to have been discussed extensively in the literature. Likewise, **Bayesian approaches**, which would tend to incorporate past data rather naturally, have received little attention. One exception is Chen *et al.* [14], which takes a Bayesian approach and also discusses using past data.

Warranty data analysis can also play a role in **quality monitoring**. Wu and Meeker [15] describe a process for monitoring a warranty database to provide an early warning of reliability problems. They adopt a Poisson model and report warranty claim rates as they would be monitored in calendar time and also as a function of vehicle age, time of production, and the type of service operation. They formulate a multivariate model and a hypothesis test to detect increasing rates, a key feature of which is the management of an overall false alarm rate.

In summary, warranty claim occurrence data can be analyzed as reliability data. Handling cost per claims may require additional analysis. However, the analysis should account for the many special features of the data such as data limitations caused by coverage limits. Due to technology enhancements warranty data are likely to become more extensive, more frequent, more complex, and linked to more

nonwarranty data sources than ever before, thereby offering considerable challenges for future analysis.

References

- [1] Murthy, D.N.P. & Djamaludin, I. (2002). New product warranty: a literature review, *International Journal of Production Economics* **79**, 231–260.
- [2] Robinson, J.A. & McDonald, G.C. (1990). Issues related to field reliability and warranty data, *Data Quality Control: Theory and Pragmatics*, Marcel Dekker, New York, pp. 69–90.
- [3] Blischke, W.R. & Murthy, D.N.P. (1994). *Warranty Cost Analysis*, Marcel Dekker, New York.
- [4] Meeker, W.Q. & Escobar, L.A. (1998). *Statistical Methods for Reliability Data*, John Wiley & Sons, New York.
- [5] Rigdon, S.E. & Basu, A.P. (2000). *Statistical Methods for the Reliability of Repairable Systems*, John Wiley & Sons, New York.
- [6] Lawless, J.F. (2003). *Statistical Models and Methods for Lifetime Data*, 2nd Edition, John Wiley & Sons, New York.
- [7] Kalbfleisch, J.D. & Prentice, R.L. (2002). *The Statistical Analysis of Failure Time Data*, 2nd Edition, John Wiley & Sons, New York.
- [8] Nelson, W.A. (2003). *Recurrent Events Data Analysis for Product Repairs, Disease Recurrences, and Other Applications*, ASA-SIAM, Philadelphia.
- [9] Chukova, S. & Robinson, J. (2005). Estimating mean cumulative functions from truncated automotive warranty data, *Modern Statistical and Mathematical Methods in Reliability, Series on Quality, Reliability and Engineering Statistics*, World Scientific, Hackensack, pp. 121–134.
- [10] Karim, M.R. & Suzuki, K. (2005). Analysis of warranty claim data: a literature review, *International Journal of Quality and Reliability Management* **22**(7), 667–686.
- [11] Blischke, W.R. & Murthy, D.N.P. (eds) (1996). *Product Warranty Handbook*, Marcel Dekker, New York.
- [12] Lawless, J.F. (1998). Statistical analysis of product warranty data, *International Statistical Review* **66**(1), 41–60.
- [13] Chukova, S., Dimitrov, B. & Rykov, V. (1993). Warranty analysis, *Journal of Soviet Mathematics* **67**, 3486–3508.
- [14] Chen, J., Lynn, N.J. & Singpurwalla, N.D. (1996). Forecasting warranty claims, *Product Warranty Handbook*, Marcel Dekker, New York, pp. 803–817.
- [15] Wu, H. & Meeker, W.Q. (2002). Early detection of reliability problems using information from warranty databases, *Technometrics* **44**(2), 120–133.

Related Articles

General Minimal Repair Models; Markov Renewal Processes in Reliability Modeling; Multiattribute Warranty Policies; Repairable Systems

Reliability; Repairable Systems: Statistical Inference; Repairable Systems: Bayesian Analysis; Replacement Strategies; Warranty: Usage and Wear Process for; Warranty Modeling;

Warranty Servicing; Warranty Cost Prediction Based on Warranty Data; Warranty Cost Analysis.

JEFFREY A. ROBINSON

Warranty: Usage and Wear Process for

Introduction

A warranty is a contract that requires the manufacturer of a product to repair or replace it if it fails. Several other articles in this encyclopedia discuss warranties and I refer you to these for a more detailed definition (*see* **Multiattribute Warranty Policies; Warranty Cost Prediction Based on Warranty Data**). This chapter discusses the issue of wear, or usage, of a product subject to a warranty. The goals of the chapter are to motivate why one should consider wear in the context of warranty analysis, describe probability models for wear processes, show how models for the lifetime of a product can be related to the amount of wear it has undergone and finally discuss how to apply these ideas to solve problems related to warranty analysis.

Warranty analysis is concerned with solving the problems surrounding warranties. These include determining the optimal warranty period, estimating the cost of a **warranty policy**, and forecasting future **warranty claims**. The computation of the cost of a warranty depends on the lifetime of the product, which is an uncertain quantity. So an important aspect of warranty analysis is to propose probability models for the lifetime of a product. Warranties are discussed extensively in this encyclopedia (*see* **Warranty Claims and Costs: Statistical Analysis of; Multiattribute Warranty Policies; Warranty Modeling; Warranty Cost Analysis; Warranty Cost Prediction Based on Warranty Data; Warranty Servicing**).

Many products for which warranties are offered operate in a dynamic environment that will vary in a random manner over time and that will affect the lifetime of the product. A convenient way to think of the effect of the environment is that it ages or induces “wear” on the product. Often such wear can have a physical interpretation, for example, on a vehicle, wear may refer to kilometres driven and in this case the warranty period may even be defined in terms of it. So a convenient approach to warranty analysis is to propose models for the accumulation of wear on the product. The wear process defines a model for time to failure. The wear process is used directly

if the warranty period is also defined in terms of use.

Wear processes are naturally modeled as stochastic processes. A warranty pertains to future failures of a product when it is taken out; the exact nature of the wear that the product will experience is unknown and varies with each sample sold. A deterministic model cannot easily capture these properties of wear while random processes can.

We note here that these ideas have much in common with the idea of degradation models and **random environments** and cumulative damage models (*see* **Degradation Models; Cumulative Damage Models Based on Gamma Processes**, respectively).

This chapter is organized as follows. The section titled “Wear Processes” defines different stochastic processes that can be used to model the accumulation of wear on a product. Broadly speaking, these models are described in the order of complexity from the simplest (the **Poisson process**). The section titled “Using Wear Processes in Warranty Analysis” shows how these models are incorporated into a model for the time to failure of a product and gives examples of applications to warranty analysis.

Wear Processes

We will always consider models that describe the cumulative amount of wear that a product accumulates, rather than the rate of wear. The goal is to define a model for how much wear a product has accumulated at any time. We let $W(t)$ denote the amount of wear at time t .

Any process that we define for $W(t)$ should satisfy $W(0) = 0$. At the time the product is bought, it is considered new and has no wear associated with it. Also, $W(t)$ should be positive; it makes no sense to talk about negative wear. One might further restrict $W(t)$ to be nondecreasing. This depends on whether $W(t)$ models the rate or intensity of wear at t , in which case $W(t) \geq 0$ is sufficient, or models the cumulative amount of wear by t , in which case we impose $t_i < t_2$ implies $W(t_1) \leq W(t_2)$. Here we discuss nondecreasing $W(t)$, so that it is the cumulative amount of wear. However, what we have to say in the section titled “Using Wear Processes in Warranty Analysis” is also valid where $W(t)$ is the rate of wear.

2 Warranty: Usage and Wear Process for

Poisson Processes

These are the most common stochastic processes satisfying our two properties for $W(t)$. It is a process that starts at zero and at random intervals of time increases in value. The times between successive increases are independent. Poisson processes are discussed in detail in **Poisson Processes** so we give only a summary here.

All Poisson processes are defined in terms of a function $\lambda(t)$ called the intensity or rate function. At any time t , the probability that the process increases value in a small interval of time δt is approximately $\lambda(t) \delta t$. The process has the property of independent increments: the increase in the process in nonoverlapping time intervals are independent.

In the (simple) Poisson process, each increase in value is always 1. The probability distribution of $W(t)$ is Poisson:

$$P(W(t) = n) = \frac{\Lambda(t)^n}{n!} e^{-\Lambda(t)}, \quad n = 0, 1, \dots \quad (1)$$

where $\Lambda(t) = \int_0^t \lambda(s) ds$. In **Intensity Functions for Nonhomogeneous Poisson Processes**, different intensity functions $\lambda(t)$ are discussed.

Poisson processes are easy to work with because the distribution of $W(t)$ is of a simple form. As a model for wear, they are suitable when the wear takes the form of a count of events; the number of times a light bulb is switched on is one example. However, it is clear that we need more complicated processes to model most types of wear that interest us.

A useful generalization is to allow the value of the process to jump at each event by a random amount, rather than letting it always be 1. If the increases are independent and identically distributed then we have the compound Poisson process. A good reference is Chapter 5 of [1], where the probability distribution function is derived. It is often possible to compute the probability density function of $W(t)$, for example, when the increases have a stable distribution, as described in [2]. If the X_i are exponentially distributed with a mean μ , where $f_X(x) = \mu^{-1} e^{-x/\mu}$, it can be shown that:

$$P(W(t) = 0) = e^{-\Lambda(t)} \quad (2)$$

$$f_{W(t)}(w) = w^{-1} e^{-\Lambda(t) - w/\mu} \sum_{i=1}^{\infty} \left(\frac{w \Lambda(t)}{\mu} \right)^i$$

$$\times \frac{1}{i! (i-1)!}, \quad w > 0 \quad (3)$$

A compound Poisson process can model usage that accumulates as randomly sized shocks, such as satellites subject to solar storms of varying magnitude. Figure 1 shows realizations of the homogeneous, nonhomogeneous, and compound Poisson processes.

Other generalizations of the Poisson process are possible: these include the doubly stochastic Poisson process, which is a Poisson process where the **intensity function** is itself random and the self-exciting point process, where the intensity function depends on the history of the process. The self-exciting point process has practical application to situations where the current usage of an item depends upon its past usage; for example, if a machine undergoes a period of heavy use, the user may then decide to let the machine stay idle for sometime in order to perform maintenance. If the period of heavy use does not occur, the machine is kept running. The reader is referred to [3, Chapter 5] for a detailed description of self-exciting point processes, but it is mentioned that, as with doubly stochastic Poisson processes, the density of $W(t)$ is in general difficult to find.

The Gamma Process

For many types of wear the increases are more continuous in nature than the Poisson process can model. The gamma process is an independent increments process that increases continuously. The idea is that the increase in the value of the process between any two times is gamma distributed, with the mean value of the increase becoming larger as the two times are further apart. Gamma processes are discussed in more detail in **Cumulative Damage Models Based on Gamma Processes**.

Recall that the gamma distribution density function is defined by two parameters: a scale parameter a and a shape parameter b :

$$f(w | a, b) = \frac{a^b}{\Gamma(b)} w^{b-1} e^{-aw}, \quad w \geq 0 \quad (4)$$

where $\Gamma(b)$ is the gamma function; the mean of the gamma distribution is b/a . In a gamma process, the shape is a nondecreasing and continuous function of time $b(t)$ with $b(0) = 0$. For two times $t_2 > t_1$, the distribution of $W(t_2) - W(t_1)$ is gamma with scale a and shape $b(t_2) - b(t_1)$. The mean of $W(t) - W(s)$ is

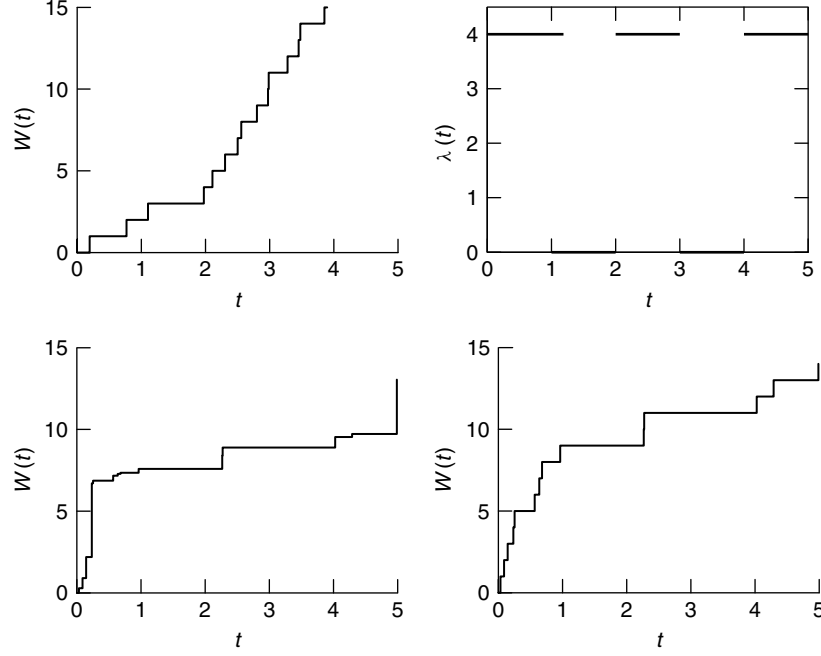


Figure 1 Realizations of three wear processes. Clockwise from top left: a homogeneous Poisson process with $\lambda(t) = 4$; an intensity function $\lambda(t) = 4$ if $\lfloor t \rfloor$ is even, $\lambda(t) = 0$ otherwise; a nonhomogeneous Poisson process with this intensity function; a compound Poisson process with this intensity function and exponentially distributed jumps with a mean of 1

then $(b(t_2) - b(t_1))/a$. The gamma process has been rigorously studied; see, for example, [4] or [5], and generalizations discussed in [6].

The gamma process is clearly suitable for modeling the wear of a system that is in continuous use. The increase in the value of the process varies randomly; that can model the fact that the system is subject to varying amounts of stress as it is used.

Markov Additive Processes

Many systems accumulate significant wear only part of the time rather than always; any system that is not used all the time, or switched on and off, will clearly accumulate wear at a much higher rate when used than not.

It is difficult to model such a situation with a gamma process. However, we can think of a wear process that sometimes acts like a gamma process (system being used) and other times stays constant (system not being used and no wear being accumulated). Çinlar has described a set of processes like this called Markov additive processes [7]. In such

a process, $W(t)$ is defined in terms of a state $S(t) \in \{0, 1\}$, where 0 represents off and 1 represents on. Switches between on and off occur at random times that are independent and exponentially distributed amounts (this is an example of what is called a Markov process). When $S(t) = 1$ then $W(t)$ increases like a gamma process. When $S(t) = 0$ then $W(t)$ stays constant.

To define such a process we need to define: the mean time switched on and the mean time switched off (these are sufficient to define the exponential distribution for the length of each on and off period) and the parameters a and $b(t)$ of the gamma process.

One nice property of this idea is that it generalizes in many useful ways. We can replace the gamma process with any other process for wear that we want. When $S(t) = 0$ we do not have to have that $W(t)$ as a constant but can allow it to increase according to another process, so allowing the system to accumulate some wear even when switched off. Also $S(t)$ can have more than two states; for example, we can define states that represent high, median, and low stress as well as off, then define a suitable wear process for

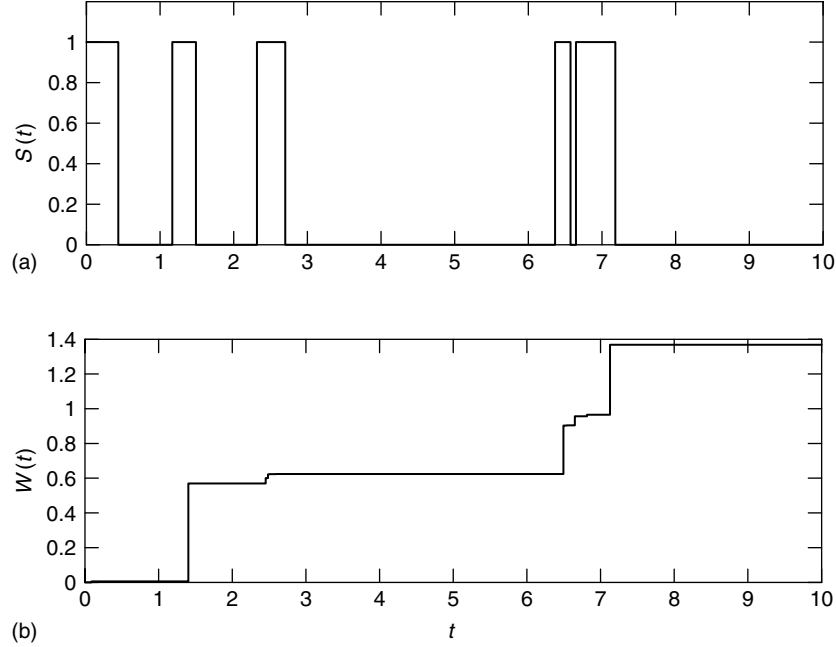


Figure 2 An example of a Markov Additive Process. (a) The state process $S(t)$ and (b) a corresponding process $W(t)$ that is a gamma process when $S(t) = 1$ and does not change when $S(t) = 0$

each. Also allowed is a “shock” to the value of $W(t)$ when it is switched on. For example the wear on a car engine in the first minutes after being switched in is considerable. Figure 2 shows a typical realization for the simple form of this process, where $S(t)$ is 0 or 1, with no wear accumulating when $S(t) = 0$ and wear accumulating as a gamma process when $S(t) = 1$.

Çinlar explores the properties of these processes in detail in [7]. Other references of note are [8] and [9], which describe similar classes of processes that can be used to model usage as a continuous, nondecreasing process.

Although the set of Markov additive processes is very general and contains processes that seem to reasonably describe many types of wear, this is at the expense of tractability, and most of the results concerning these processes are characterization results and not explicit formulations of the density of $W(t)$. In some cases, however, one is able to describe the density function of $W(t)$. For example, let $W(t)$ be a Markov additive process with an active and idle state $S(t) \in \{0, 1\}$. Let $G_1(t)$ be a gamma process with scale parameter a and shape function $b(t) = bt$ and assume no wear occurs when idle, so $G_0(t) = 0$.

We start by defining $W(t | \{S(u), 0 \leq u \leq t\})$ as the process $W(t)$ conditional on knowing the entire history of the state to time t . We show in [10] that $W(t | \{S(u), 0 \leq u \leq t\})$ is gamma distributed with scale a and shape function $b(s) = b \int_0^t S(s) ds$; thus $b(s)$ is simply b multiplied by the amount of time to t that the system is working. We then show how to use this to approximate the **probability density function** of $W(t)$.

Using Wear Processes in Warranty Analysis

In this section, I will discuss one way in which wear processes can be used in warranty analysis, namely to define a joint distribution of the time and usage at failure of the item. All of this is described in much more detail in [11] and [10], which we refer you to for details that we skip below.

This is relevant for warranties that are defined in terms of the wear as well as time, for example, in automobiles, if wear is defined as distance traveled. The earliest use of this idea is [12], who used it to

look at the optimal replacement policies for items subject to wear.

Let T and W be the time and wear at failure of the item. Let $f_{T,W}(t, w)$ be the probability density function of time and wear at failure. Then we can write by the multiplication law of probability $f_{T,W}(t, w) = f(t | w) f(w)$. Now w must be the wear at time t and so:

$$f_{T,W}(t, w) = f(t | W(t) = w) f_{W(t)}(w) \quad (5)$$

where $f(t | W(t) = w)$ is the probability density of failure at time t given the wear at that time and $f_{W(t)}(w)$ is the density function of the amount of wear at time t . The density function $f_{W(t)}(w)$ comes from the wear process model i.e., it is a Poisson distribution if $W(t)$ is a Poisson process, etc.

A simple model for $f(t | W(t) = w)$ is to say that the failure rate $r(t)$ of the item is an increasing function of the amount of wear that it suffers; simple forms for $r(t)$ might be $r(t) = r_0(t) + aW(t)$ or $r(t) = r_0(t)e^{aW(t)}$, with $r_0(t)$ the baseline failure rate under no wear and $a > 0$. I take the first form in what follows. Recall that the reliability function is $P(T > t) = \bar{F}(t) = \exp(-\int_0^t r(s) ds)$ and $f(t | W(t) = w) = -d/dt \bar{F}(t | W(t) = w)$, hence:

$$\begin{aligned} \bar{F}(t | W(t) = w, W(s), 0 \leq s \leq t) \\ = \exp\left(-\int_0^t r_0(s) - a \int_0^t W(s) ds\right) \end{aligned} \quad (6)$$

So the reliability depends on the integrated wear $\int_0^t W(s) ds$ given $W(t) = w$, which we denote as $\mathcal{W}(t)$. This is a random quantity just like $W(t)$. The reliability $\bar{F}(t)$ unconditional on $\mathcal{W}(t)$ is what we want and in [10] we show how this can be done. One example that produces a nice mathematical form for $f(t | W(t) = w)$ is if $W(t)$ is a Poisson process with intensity function $\lambda(t)$. Defining $R_0(t) = \int_0^t r_0(s) ds$ to be the cumulative baseline hazard, we show that:

$$\begin{aligned} f(t | W(t) = w) \\ = (r_0(t) + aw) e^{-R_0(t)} \Lambda(t)^{-w} \\ \times \left[\int_0^t \lambda(s) e^{-a(t-s)} ds \right]^w \end{aligned} \quad (7)$$

and hence the joint density function of time and wear at failure is:

$$\begin{aligned} f_{T,W}(t, w) = \frac{r_0(t) + aw}{w!} e^{-R_0(t) - \Lambda(t)} \\ \times \left(\int_0^t \lambda(s) e^{-a(t-s)} ds \right)^w \end{aligned} \quad (8)$$

If $W(t)$ is homogeneous (with $\Lambda(t) = \lambda t$) and $r_0(t) = r$ then we have

$$f_{T,W}(t, w) = \frac{r + aw}{w!} \left(\frac{\lambda}{a} (1 - e^{-at}) \right)^w e^{-(r+\lambda)t} \quad (9)$$

for $t \geq 0$ and $w = 0, 1, 2, \dots$

Figure 3 shows two examples of equation (9) for different values of the wear rate parameter λ , where we see that as λ increases there is higher probability on failures at an earlier time but with more wear.

Finally, we note that by integrating out w in equation (8) we obtain the density of time to failure alone, which can be used in applications where time is the only warranty measure. In the case of equation (9), it turns out that this gives us a reliability function:

$$\bar{F}(t) = \exp\left(-(r + \lambda)t + \frac{\lambda}{a}(1 - e^{-at})\right), \quad t \geq 0 \quad (10)$$

Applications

For warranties, the derivation of $\bar{F}(t)$ or $f_{T,W}(t, w)$ is only half the story. These distributions must then be used to solve a warranty-related problem. Space does not permit a description of these here but the reader is referred to [10], where there is a discussion on how to produce an optimal warranty size for a product, given data on when it failed and with how much use. A similar problem is solved in [13], although the wear process model there is not a stochastic process. The reader is also referred to [14] which describes several problems associated with warranties where a probability model for time and usage to failure is part of the solution.

References

- [1] Ross, S.M. (2003). *Introduction to Probability Models*, 8th Edition, Academic Press, San Diego.

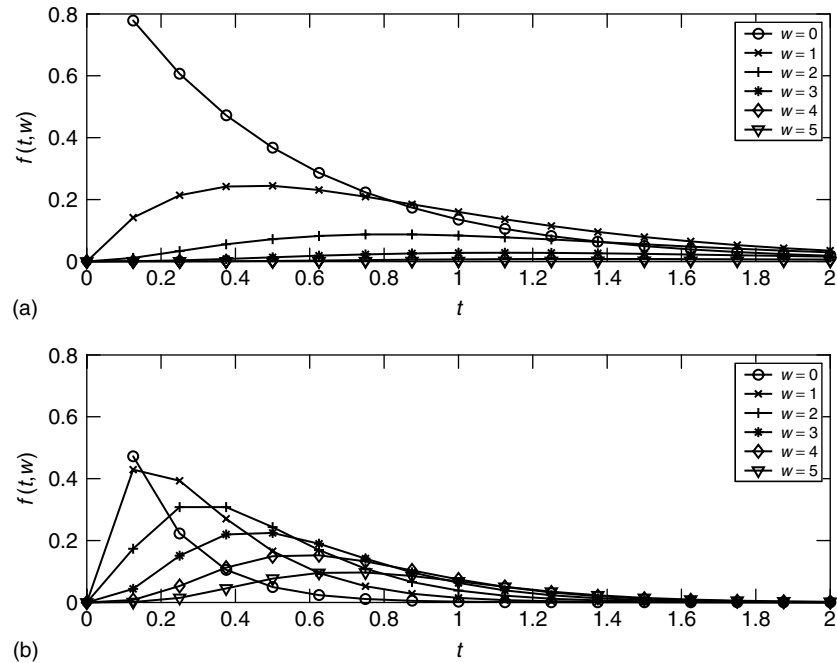


Figure 3 The joint probability density of time and wear at failure of equation (9) for $r = 1$, $a = 0.5$ with (a) $\lambda = 1$ and (b) $\lambda = 5$

- [2] Nolan, J.P. (2002). *Stable Distributions*, Birkhäuser, New York.
- [3] Snyder, D.L. & Miller, M.I. (1991). *Random Point Processes in Time and Space*, 2nd Edition, Springer-Verlag, New York.
- [4] Çinlar, E. (1980). On a generalization of the gamma process, *Journal of Applied Probability* **17**, 467–480.
- [5] Ferguson, T.S. (1973). A Bayesian analysis of some non-parametric problems, *Annals of Statistics* **1**, 209–230.
- [6] Dykstra, R.L. & Laud, P. (1981). A Bayesian nonparametric approach to reliability, *Annals of Statistics* **9**, 356–367.
- [7] Çinlar, E. (1972). Markov additive processes II, *Zeitschrift für Wahrscheinlichkeitstheorie und Verwandte Gebiete* **24**, 94–121.
- [8] Çinlar, E. (1984). *Markov and Semi-Markov Models of Deterioration*, Academic Press, Orlando, pp. 3–41.
- [9] Esary, J.D., Marshall, A.W. & Proschan, F. (1973). Shock models and wear processes, *Annual Proceedings* **1**, 627–649.
- [10] Wilson, S.P. & Singpurwalla, N.D. (1998). Failure models indexed by two scales, *Advances in Applied Probability* **30**, 1058–1072.
- [11] Wilson, S.P. (1993). *Failure modeling with multiple scales*, Ph.D. thesis, Department of Operations Research, The George Washington University.
- [12] Mercer, A. (1961). Some simple wear-dependent renewal processes, *Journal of the Royal Statistical Society. Series B* **23**, 368–376.
- [13] Eliashberg, J., Singpurwalla, N.D. & Wilson, S.P. (1997). Calculating the reserve for a time and usage indexed warranty, *Management Science* **43**, 966–975.
- [14] Singpurwalla, N.D. & Soyer, R. (1992). Non-homogeneous autoregressive processes for tracking (software) reliability growth, and their Bayesian analysis, *Journal of the Royal Statistical Society. Series B* **54**, 145–156.

Related Articles

Warranty Claims and Costs: Statistical Analysis of; Multiattribute Warranty Policies; Warranty Modeling; Warranty Cost Analysis; Warranty Cost Prediction Based on Warranty Data; Warranty Servicing

SIMON P. WILSON

Multiattribute Warranty Policies

Introduction

Nearly everything sold in the market is covered by a warranty – expressed or implied. Various types of warranties serve a number of purposes, including protection for the buyer against low quality or service and limitation of product liability for the seller. In some instances, warranties are mandated and regulated by state and local governments. Because of this diversity of purpose, warranties have been studied by researchers in many different disciplines, such as engineering, management, marketing, economics, law, and accounting. For more details of warranty applications, see **Warranty Servicing**.

A warranty is a seller's guarantee given to the buyer stating that a product is reliable and free from known defects and that the seller will repair or replace the product if it fails to conform to satisfactory performance. Consequently, this obligation generates costs to the producer associated with failures of the product. Thus, one of the decision problems faced by producers in warranty analysis is how to compare various types of warranty plans and estimate the total cost associated with each plan.

A formal development of warranty models can be traced back to a paper by Arrow [1]. Since then, a wide variety of warranty policies has been proposed and various statistical models for determining warranty costs have been developed. Now one can say that warranty analysis constitutes a "field" of study within engineering–statistics–marketing. As further evidenced by a series of review papers by Blischke and Murthy [2] and Murthy and Blischke [3, 4], and a handbook edited by Blischke and Murthy [5], this particular field of study has experienced a rapid growth and extensive applications to various managerial decision problems.

In this note, we first classify various types of warranty plans proposed in the literature, and then review several methods of modeling two-attribute warranty policies. For further development of warranty models and analysis methods, readers are referred to Thomas and Rao [6] and Murthy and Djamaludin [7]. For articles in this encyclopedia that are related to the

analysis of warranty policies see the section titled "Related Articles".

Warranty Policies

A warranty can be thought of as a contractual *obligation* or a potential *liability* to the seller. However, the limit of the seller's obligation or liability is usually specified at the time of product sale and the seller is responsible only for the product failures satisfying the prespecified warranty criteria. On the basis of the number of warranty attributes used in determining the warranty eligibility, warranty plans can be classified as (a) *single-attribute* warranty policies and (b) *multiattribute* warranty policies.

In *single-attribute* warranty policies, we use only one of the warranty characteristics, such as the time, the hours of operation, the mileage, or the frequency of product usage. The most popular criterion is the *warranty period* that is measured from the time of sale of the product to a certain point in the product's lifetime. The popularity of the warranty period is twofold. In most instances, the total number of product failures is a function of time. Besides, time is easily measurable and thus tends to minimize potential conflicts between a seller and a buyer on the warranty eligibility of a failed product.

In *multiattribute* warranty policies, on the other hand, more than one warranty characteristics are employed simultaneously as criteria in judging the warranty eligibility of a failed product. Under the 5-year, 50 000-mi protection plan used in the automobile industry, for instance, the warranty eligibility of an automobile is determined by its age and mileage at the time of breakdown. Since its introduction in the late 1960s, the two-attribute warranties have become a critically important segment of the automobile business environment.

It is reported that aircraft manufacturers also offer this type of two-attribute warranty policy, with usage corresponding to the number of flying hours. In the next section, we review various types of single-attribute warranty policies.

Single-Attribute Warranty Policies

Under the typical situation of a single-attribute warranty transaction, a seller provides some kind of repair or replacement service for any product

2 Multiattribute Warranty Policies

failures occurring during some specified time. For the warranty service, a buyer pays either no cost or a pre-determined portion of the full repair or replacement cost. Therefore, two of the major decision variables in a single-attribute warranty policy are the *warranty price* and the *warranty period*: (a) *how much* a buyer should pay to receive the warranty service and (b) *how long* a seller should provide the warranty service.

On the basis of the nature of a buyer's payment of repair or replacement cost upon a product failure, warranty policies can be divided into (a) the free-replacement warranty, (b) *pro rata* warranty, and (c) lump-sum warranty. First, under the *free-replacement* warranty policy shown in Figure 1(a), a producer provides as many replacements or repairs as needed free of charge to a consumer as long as the product failures satisfy the warranty criteria. In many cases, the anticipated cost of honoring the free warranty service has already been included in the selling price of the product or put into reserve to meet future **warranty claims**.

Second, under the *pro rata* warranty (or rebate) policy, a consumer must pay a fraction of the full repair cost. As shown in Figure 1(b), the fraction is usually determined by the percentage of the warranty period in use by the old item prior to the failure. An automobile tire typically carries such a *pro rata* warranty, in which a buyer pays a small portion of the full replacement cost on the basis of the remaining tread depth or the mileage driven.

Third, under the *lump-sum* warranty policy shown in Figure 1(c), a customer receives a fixed or lump-sum rebate for any product failures occurring within the warranty period. Particularly when it is difficult to measure the exact time of product failure, this policy is employed to minimize potential conflicts

between a producer and a customer about the amount of compensation due for product failure.

We can also imagine a hybrid form of the aforementioned warranty policies. As shown in Figure 1(d) for example, a total warranty period can be divided into two periods: an initial free-replacement warranty period, followed by a *pro rata* period.

On the basis of how long a producer should provide the warranty service, warranties can be classified as (a) fixed-period and (b) renewable policies. First, under the *fixed-period* warranty policy, the warranty period is specified at the time of product sale and fixed regardless of a producer's repairs or replacements. The warranty terms expire automatically at the end of the warranty period. This type of warranty policy is the most popular in practice.

Second, in the *renewable* warranty policy, the original warranty period is extended whenever a repair or replacement occurs during the current warranty period. When the current warranty period expires without any product failures, the warranty terms expire completely. Renewal theory, described in **Renewal Theory**, can be applied to describe this failure process if we assume that successive failure times form a renewal reward process.

We can imagine different types of multiattribute warranty policies but, in the next section, we restrict our attention to the two-attribute warranty policies, which is the most common in practice. For more technical presentations of warranty models, readers are referred to **Warranty Modeling**.

Multiattribute Warranty Policies

As shown in Figure 2, various types of two-attribute warranty policies have been considered in literature: (a) rectangular, (b) triangular, (c) L-shaped, and (d) parabolic warranty policies. The 5-year, 50 000-mi automobile warranty can be displayed in a *rectangular* form as in Figure 2(a); any breakdowns of an automobile will be eligible for the warranty service as long as its age is less than x^* years and its mileage is less than y^* mi. In terms of set notation, the region covered by the warranty is denoted by $\Omega = \{(X, Y) \text{ for } X < x^* \text{ and } Y < y^*\}$.

In the *triangular* warranty policy, on the other hand, the warranty region is expressed as $\Omega = \{(X, Y) \text{ for } aX + bY < c\}$, where a, b , and c are

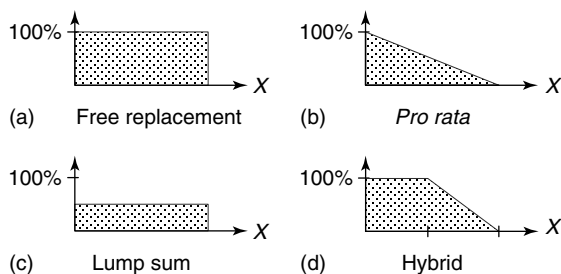


Figure 1 Various types of single-attribute warranty policies

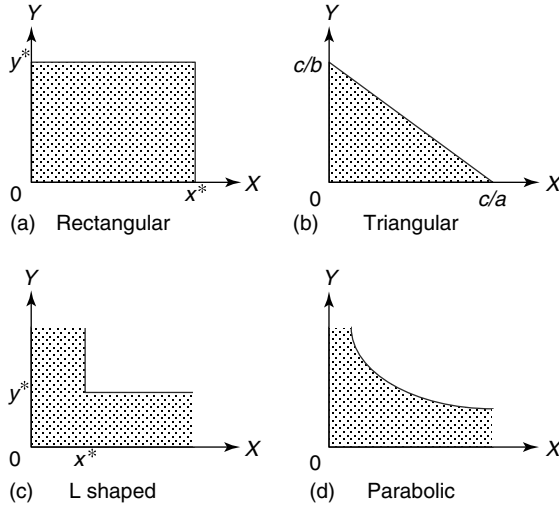


Figure 2 Various types of two-attribute warranty policies

constants specified by a seller at the time of product sale. We can also imagine the *L-shaped* warranty policy as shown in Figure 2(c), in which the product is eligible for the warranty service if X is less than x^* or Y is less than y^* ; i.e. $\Omega = \{(X, Y) \text{ for } X < x^* \text{ or } Y < y^*\}$.

Since the product failure is a function of X and Y , the ISO **warranty cost** curve can be drawn on the plane for a given cost as shown in Figure 2(d). Moskowitz and Chun [8] proposed such a *parabolic* warranty policy, in which the producer offers a set of warranty plans for the same price and each customer can choose any plan in the set. For example, the producer may offer $\Omega = \{(2 \text{ years, } 100\,000 \text{ mi}), (5 \text{ years, } 50\,000 \text{ mi}), (8 \text{ years, } 30\,000 \text{ mi})\}$. In such a case, a retired grandmother would prefer the 8-year, 30 000-mi warranty policy, whereas the 2-year, 100 000-mi plan is appealing to a traveling salesman. While any other type of two-attribute warranty policy tends to favor one group of consumers over another, the parabolic warranty policy is shown to be more equitable to all consumers.

Chun and Tang [9] analyzed the relationships among the aforementioned two-attribute warranty policies. They also showed that the single-attribute warranty policies are special cases of two-attribute policies.

Modeling Two-Attribute Warranties

Offering warranty coverage leads to additional costs to the producer, associated with claims under warranty. Thus, a producer needs to estimate the expected total cost of offering a two-attribute warranty policy. Various approaches have been proposed in modeling two-attribute warranty policies: (a) two-dimensional, (b) one-dimensional, and (c) Markovian approaches.

First, in two-attribute warranty policies, product failures are dependent on the age X and the usage Y of a product. Thus, the sequence of product failures can be modeled as a counting process on the two-dimensional plane of time and mileage. In the counting process, the time and the usage between successive failures can be described as a bivariate **probability density function**. This method of modeling is called a *two-dimensional approach* to a two-attribute warranty plan. As a bivariate probability density function, the bivariate exponential distribution is a natural choice if we assume that the times and the usages between successive failures are independent and identically distributed random variables. The beta-Stacy distribution and the multivariate Pareto distribution of the second kind have been also considered by other researchers [10].

Second, in many practical situations, the failure process on the two-dimensional plane can be represented by a univariate renewal process. This formulation method is referred to as the *one-dimensional approach*, which takes advantage of the linear relationship between X and Y . For most drivers, for example, the average mileage driven per year is constant; that is, the product usage Y is a linear function of the time X .

Note that not all products of a kind have the same usage rate; a high-intensity user like a traveling salesman for example, has a high usage rate, while a low-intensity user like a retired grandmother presumably has a low rate of use. Thus, the usage rate is assumed to be a nonnegative random variable of the continuous type. The distribution of the usage rate could be modeled satisfactorily by a gamma distribution as in Moskowitz and Chun [8]; by changing its parameter values, we can represent a wide variety of distributions with different location, dispersion, shape, and the like. Other types of distributions, such as the uniform distribution, the β -weighted two-point distribution, and the

4 Multiattribute Warranty Policies

exponential distribution, have also been considered in literature.

Third, in most practical situations, the age x^* and the usage y^* in two-attribute warranty policies take discrete values; that is, it is unimaginable to offer a 5.2-year, 71 534-mi protection plan, even though those are the optimal values for a given warranty cost. Thus, two-attribute warranty policies can be modeled *via* a Markovian approach [11], in which the usage rate per year is discretized as a *state* and the time is also discretized as a *transition*. Figure 3 for example, shows a transition probability matrix for a 5-year, 50 000-mi warranty plan.

In the transition probability matrix in Figure 3, the state i is the mileage of an automobile, while the state ∞ represents the automobiles with more than 50 000 mi. The transition probability p_{ij} is the probability that the state of an automobile with the mileage in the range i will be changed to the state j over a period of 1 year. The state ∞ is an *absorbing* state, while other states are *transient*. Thus, two-attribute warranty policies can be easily modeled as an absorbing Markov chain, and we can find many interesting facts about the absorbing chain such as the expected number of times that each state will be entered before we reach an absorbing state. We can also consider the discounting factor per transition and the repair cost for each state.

Concluding Remarks

Warranties have been important elements in business planning and decision making for manufacturers and producers of services. In this note, we have investigated various two-attribute warranty policies and reviewed three different modeling approaches, but no research works have been reported in the

$$P = \begin{matrix} & \begin{matrix} 0-10 & 10-20 & 20-30 & 30-40 & 40-50 & \infty \end{matrix} \\ \begin{matrix} 0-10 \\ 10-20 \\ 20-30 \\ 30-40 \\ 40-50 \\ \infty \end{matrix} & \left\{ \begin{array}{cccccc} 0.1 & 0.2 & 0.3 & 0.3 & 0.1 & 0 \\ 0 & 0.1 & 0.4 & 0.3 & 0.1 & 0.1 \\ 0 & 0 & 0.2 & 0.4 & 0.2 & 0.2 \\ 0 & 0 & 0 & 0.2 & 0.4 & 0.4 \\ 0 & 0 & 0 & 0 & 0.3 & 0.7 \\ 0 & 0 & 0 & 0 & 0 & 1 \end{array} \right\} \end{matrix}$$

Figure 3 An example of the transition probability matrix in the Markovian approach

open literature that compare the effectiveness of those approaches. There may be some cases where the more mathematically onerous two-dimensional method should be used. There may be other cases in which a simpler method, such as the one-dimensional and Markovian approaches, is sufficient enough to model the aggregate behavior of entire products. But our past experience in decision analysis and on other management science problems suggests that the differences among them will be nominal. An empirical study based on automobile warranty data could answer some questions as to which approaches are more effective in modeling the two-dimensional failure process and in estimating the parameters of warranty models.

Since its appearance in literature in the late 1980s, the two-attribute warranty policy has been extended and generalized in many different directions by releasing some of the basic assumptions in original models and proposing other new policies. But few research models have been successfully applied to the automobile industry. The new research direction in the area of two-attribute warranty policy should be not just to propose another decision model that is based on an artificial fact, but to collect empirical data from the industry to address real problems that are empirically grounded – not only the source that came from the field, but also the solution is of interest to the practitioners.

References

- [1] Arrow, K. (1963). Uncertainty and the welfare economics of medical care, *The American Economic Review* **53**, 941–973.
- [2] Blischke, W.R. & Murthy, D.N.P. (1992). Product warranty management – I: a taxonomy for warranty policies. *European Journal Operational Research* **62**, 127–148.
- [3] Murthy, D.N.P. & Blischke, W.R. (1992). Product warranty management – II: an integrated framework for study, *European Journal of Operational Research* **62**, 261–281.
- [4] Murthy, D.N.P. & Blischke, W.R. (1992). Product warranty management – III: a review of mathematical models, *European Journal of Operational Research* **63**, 1–34.
- [5] Blischke, W.R. & Murthy, D.N.P. (1995). *Product Warranty Handbook*, Marcel Dekker, New York.
- [6] Thomas, M.U. & Rao, S.S. (1999). Warranty economic decision models: a summary and some suggested directions for future research, *Operations Research* **47**, 807–820.

- [7] Murthy, D.N.P. & Djamaludin, I. (2002). New product warranty: a literature review, *International Journal of Production Economics* **79**, 231–260.
- [8] Moskowitz, H. & Chun, Y.H. (1994). A Poisson regression model for two-attribute warranty policies, *Naval Research Logistics* **41**, 355–376.
- [9] Chun, Y.H. & Tang, K. (1999). Cost analysis of two-attribute warranty policies based on the product usage rate. *IEEE Transactions on Engineering Management* **46**, 201–209.
- [10] Murthy, D.N.P., Iskandar, B.P. & Wilson, R.J. (1995). Two-dimensional failure-free warranty policies: two-dimensional point process models, *Operations Research* **43**, 356–366.
- [11] Balachandran, K.R., Maschmeyer, R.A. & Livingstone, J.L. (1981). Product warranty period: a Markovian approach to estimation and analysis of repair and replacement costs, *The Accounting Review* **56**, 115–124.

Related Articles

Renewal Theory; Repairable Systems Reliability; Warranty Claims and Costs: Statistical Analysis of; Warranty Cost Analysis; Warranty Cost Prediction Based on Warranty Data; Warranty Modeling; Warranty Servicing; Warranty: Usage and Wear Process for.

YOUNG H. CHUN

Warranty Modeling

Introduction

A warranty is a commitment by a manufacturer or producer to provide products or services of acceptable quality for some specified period of time. During the warranty period the producer assumes some portion of the expenses from the repair or replacement of products that are defective or fail to meet customer expectations. Warranties have been in existence since people have exchanged goods and services [1], but it has only been during the past 50 years that quantitative methods have evolved for treating warranty decisions in the production of consumer products. Early efforts in warranty modeling focused on accounting for product costs and in establishing market position for manufactured products, but by the 1970s the onset of global competitiveness and needs for productivity improvement led to a major shift to a quality focus in manufacturing. Consequently, interest grew in developing methods for analyzing and evaluating economic decisions pertaining to the development and planning of warranties [2].

This article summarizes quantitative models for examining warranty decisions. Although the methods presented can also be applied to service systems, we will restrict attention to the view of a manufacturer producing consumer products. Most products are covered by some warranty either directly or indirectly through consumer protection laws. Following a list of notation and acronyms, in the next section we discuss the basic types of warranties that are applied to consumer products. We then present models for predicting the warranty costs, the amount of reserve for covering future costs, and determining optimum warranty periods, followed by some concluding remarks.

Notation

- A : fixed cost for administering a warranty program
- b_0, b_1 : parameters for warranty benefit function
- $B(w)$: expected cost–benefit function for warranty over $(0, w)$
- c_0 : initial unit cost for warranty failure
- c_1 : unit repair/replacement cost

- $C(w; \phi)$: expected cost for warranty of length w under policy ϕ
- δ : discounting factor, $0 \leq \delta \leq 1$
- $f(t), F(t)$: probability density, cumulative distribution function
- $F^{(n)}$: n -fold convolution of F with itself
- K : production lot size
- λ : parameter for the exponential distribution
- $M(t), m(t)$: renewal function, renewal density function
- $N(t)$: number of warranty failures in $(0, t)$
- π : expected price change per period
- warranty policy**
- ϕ :
- $R(w)$: reserve resources for warranty of length $w > 0$
- θ : expected rate of return
- $X(t)$: unit **warranty cost**
- $X_D(t)$: unit warranty cost with discounting.

Types of Warranties

A warranty is specified by a period of coverage and a set of conditions for compensation. It can be renewing so that items replaced because of failures during the warranty period are granted new warranties, or it can be nonrenewing whereby the period of coverage is fixed. Renewing warranties generally only apply to irreplaceable items that are relatively inexpensive. Most warranties are based on free replacement or *pro rata* rebate schemes.

Free-Replacement Warranty (FRW)

Under a free-replacement policy, the manufacturer absorbs all costs associated with the repair or replacement of products that fail during the warranty period. The free-replacement warranty (FRW) policy is common for repairable items like televisions, washing machines, and lawn mowers that can be brought back to a good-as-new operating condition through repairs. It also applies to nonreparable items like flashbulbs, disposable medical products, and inexpensive electronic components.

Let T be a random variable representing the time to failure of a product that has a repair or replacement cost to the manufacturer of c_1 dollars per unit. Under the FRW policy, the unit cost to the manufacturer

2 Warranty Modeling

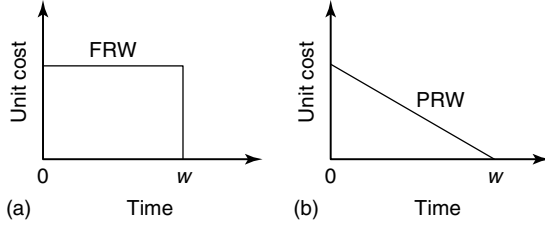


Figure 1 Unit warranty cost for FRW and PRW policies

at time t for a warranty of length $w > 0$, shown in Figure 1(a), is

$$X(t) = \begin{cases} c_1, & 0 < t \leq w \\ 0, & t > w \end{cases} \quad (1)$$

Therefore, the expected cost of a single failure is

$$E[X(t)] = \int_0^w c_1 f(t) dt = c_1 F(w) \quad (2)$$

where $f(t)$ and $F(t)$ are the probability density and cumulative distribution functions for T .

Pro rata Replacement Warranty

A *pro rata* policy provides customers with an amount of rebate that is a function of the age of the item when it fails within the warranty period. This policy is applied to nonrepairable items. Products like tires are generally sold with a *pro rata* replacement warranty (PRW) policy.

Under a PRW policy with a linear *pro rata* function, an item that costs the manufacturer c_0 dollars per unit that is sold with a warranty of length $w > 0$ has a unit cost at time t , as shown in Figure 1(b), that is given by

$$X(t) = \begin{cases} c_0 \left(\frac{w-t}{w} \right), & 0 < t \leq w \\ 0, & t > w \end{cases} \quad (3)$$

The customer must be willing to pay the complement portion of the cost, i.e., $c_0 t/w$, to execute the warranty under a PRW policy. The expected unit cost is

$$E[X(t)] = \int_0^w c_0 \left(\frac{w-t}{w} \right) f(t) dt \quad (4)$$

or

$$E[X(t)] = c_0 F(w) - \frac{c_0}{w} \int_0^w t f(t) dt \quad (5)$$

Example 1 A turbine engine with a manufacturer's cost of \$1000 is sold under a 15 000-h PRW policy. The failure rate for the engine is 2×10^{-5} failures per h of operation. Assume the lifetime of the engine is exponentially distributed. Find the expected unit warranty cost.

Solution Substituting into equation (5) for the exponential distribution

$$\begin{aligned} E[X(w)] &= c_0[1 - e^{-\lambda w}] + \frac{c_0}{w} \int_0^w t \lambda e^{-\lambda t} dt \\ &= c_0(1 - e^{-\lambda w}) + \frac{c_0}{\lambda w} [1 - (1 + \lambda w)e^{-\lambda w}] \end{aligned} \quad (6)$$

For the case where $c_0 = \$1000$; $w = 15\,000$ h; and $\lambda = 2 \times 10^{-5}$ failures per h,

$$\begin{aligned} E[X(15\,000)] &= 1000[1 - e^{-(2 \times 10^{-5})(15\,000)}] \\ &\quad + \frac{1000}{(2 \times 10^{-5})(15\,000)} \\ &\quad \times \left\{ 1 - [1 + (2 \times 10^{-5})(15\,000)] \right. \\ &\quad \left. \times e^{-(2 \times 10^{-5})(15\,000)} \right\} = \$382.30 \quad (7) \end{aligned}$$

Combined FRW/PRW

Many products are sold under a combined policy whereby items that fail during an initial period $(0, w_1)$ are covered in full by an FRW policy, followed by a PRW policy during (w_1, w_2) . Let c_0 be the initial cost. The unit cost, as shown in Figure 2,

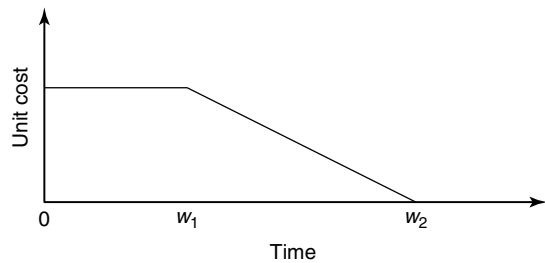


Figure 2 Unit warranty cost for combined FRW/PRW policy

is given by

$$X(t) = \begin{cases} c_0, & 0 < t \leq w_1 \\ c_0 \left(\frac{w_2 - t}{w_2 - w_1} \right), & w_1 < t \leq w_2 \\ 0, & t > w_2 \end{cases} \quad (8)$$

The expected cost for a single claim, i.e., single failure, is therefore,

$$\begin{aligned} E[X(t)] &= c_0 \int_0^{w_1} f(t) dt + \frac{c_0}{w_2 - w_1} \\ &\quad \times \int_{w_1}^{w_2} (w_2 - t) f(t) dt \\ &= c_0 F(w_1) + \left(\frac{c_0 w_2}{w_2 - w_1} \right) [F(w_2) - F(w_1)] \\ &\quad - \frac{c_0}{w_2 - w_1} \int_{w_1}^{w_2} t f(t) dt \end{aligned} \quad (9)$$

Integrating the last term by parts and simplifying,

$$E[X(t)] = \left(\frac{c_0}{w_2 - w_1} \right) \int_{w_1}^{w_2} F(t) dt \quad (10)$$

Example 2 A tire manufacturer produces a particular brand of tire at a cost of \$50 per tire that is sold under a 24-month warranty with a 90-day FRW policy followed by a PRW policy for the remaining period. Failure times are exponentially distributed with a mean time to failure of 30 months. Assume customers are willing to pay their prorated portion of replacement costs when it is required.

Here we have the parameter values as follows:

$$\begin{aligned} c_0 &= \$50, w_1 = 3 \text{ months}, \\ w_2 &= 24 \text{ months}, \lambda = 1/30 \end{aligned}$$

Substituting into equation (10), the expected unit cost is

$$\begin{aligned} E[X(24)] &= \left(\frac{50}{24 - 3} \right) \int_3^{24} (1 - e^{-t/30}) dt \\ &= 2.38[21 + 30(e^{-24/30} - e^{-3/30})] \\ &= \$17.46 \end{aligned} \quad (11)$$

Expected Discounted Unit Cost

For cases where warranty repair and replacement costs are high and warranties are of relatively long

duration, cost models should include consideration for the time value of money. Denoting the discount factor by δ , $0 < \delta < 1$, the expected discounted unit cost is

$$E[X_D(t)] = E[e^{-\delta t} X(t)] \quad (12)$$

FRW with Discounting. From equation (2) it follows that

$$E[X_D(t)] = c_1 \int_0^w e^{-\delta t} f(t) dt \quad (13)$$

For failure time distributed exponentially,

$$E[X_D(t)] = c_1 \int_0^w \lambda e^{-(\lambda+\delta)t} dt \quad (14)$$

or

$$E[X_D(t)] = \frac{\lambda c_1}{\lambda + \delta} [1 - e^{-(\lambda+\delta)w}], \quad w \geq 0 \quad (15)$$

PRW with Discounting. Modifying equation (5) for discounting,

$$\begin{aligned} E[X_D(t)] &= c_1 \int_0^w e^{-\delta t} \left(\frac{w - t}{w} \right) f(t) dt \\ &= c_1 \int_0^w e^{-\delta t} f(t) dt - \frac{c_1}{w} \int_0^w e^{-\delta t} t f(t) dt \end{aligned} \quad (16)$$

and integrating by parts it is easily shown that

$$E[X_D(t)] = \frac{c_1}{w} \int_0^w \int_0^t e^{-\delta \tau} f(\tau) d\tau \quad (17)$$

For exponentially distributed warranty failures,

$$\begin{aligned} E[X_D(t)] &= \frac{c_1}{w} \int_0^w \int_0^t \lambda e^{-(\lambda+\delta)\tau} d\tau \\ &= \frac{c_1}{w} \int_0^w \frac{\lambda(1 - e^{-(\lambda+\delta)\tau})}{\lambda + \delta} d\tau \end{aligned} \quad (18)$$

which can be simplified to

$$\begin{aligned} E[X_D(t)] &= \frac{\lambda c_1}{\lambda + \delta} - \frac{\lambda c_1}{w(\lambda + \delta)^2} \\ &\quad \times [1 - e^{-(\lambda+\delta)w}], \quad w \geq 0 \end{aligned} \quad (19)$$

Example 3 A special device costs the manufacturer \$2500 to produce and is sold under a 5-year PRW policy. Failure times are exponentially distributed

4 Warranty Modeling

with a mean time to failure of 4 years. Given that the value of money is discounted at a rate of 6%, determine the expected cost for warranty failures. Assume that customers are willing to pay their share of costs when warranty failures occur.

Substituting the parameter values $\lambda = 1/4$, $c_1 = \$2500$, and $\delta = 0.06$ into equation (19),

$$E[X_D(t)] = \frac{(0.25)(2500)}{0.25 + 0.06} - \frac{(0.25)(2500)}{5(0.25 + 0.06)^2} \times [1 - e^{-(0.25+0.06)5}] = \$991.48 \quad (20)$$

Warranty Cost Models

Warranty decisions can be critical in product planning, particularly for new product development. The manufacturer needs to know enough about the failure characteristics and related costs to predict the overall product cost. There is risk in the revenues from sales and of course the warranty costs. The manufacturer's total cost is the initial cost of producing the product plus the accumulation of costs for **warranty claims** from failures that occur during the warranty period. The following models are useful in examining the costs for given warranty policies, estimating the amount of reserve to cover warranty costs, and determining minimum cost warranties.

Predicting Manufacturer Cost

Consider a product with initial production cost c_0 that is sold under a warranty policy ϕ , such as FRW, PRW, or some other policy. We shall assume that the times between warranty failures occur as independent random variables $F(t)$ that are identically distributed and each such event $i = 1, 2, \dots$ results in a cost $X_i(t)$. The manufacturer's expected cost is then given by

$$C(w; \phi) = c_0 + E\left(\sum_{i=1}^{N(w)} X_i(t)\right) \quad (21)$$

From renewal theory (*see* for example [3] and **Renewal Theory**) it follows that with $N(w)$ being a stopping time

$$E\left(\sum_{i=1}^{N(w)} X_i(t)\right) = E[X_i(t)]E[N(w)] \quad (22)$$

where $E[X_i(t)]$ is the expected unit cost for a warranty failure given by equations (2) and (5) for FRW and PRW policies. Further, with $X_i(t)$ properly behaved over time $t > 0$, $E[X_i(t)]$ can be treated as a constant c_1 then from equation (19)

$$C(w; \phi) = c_0 + c_1 M(w) \quad (23)$$

where

$$M(w) = E[N(w)] = \sum_{n=1}^{\infty} F^{(n)}(t) \quad (24)$$

is the renewal function and $F^{(n)}(t)$ is the n -fold convolution of $F(t)$ with itself (*see* **Renewal Theory**).

Fixed-Period Warranties. For most consumer products warranty coverage expires after some fixed period of time $w > 0$. Under ideal conditions an item is placed in service functioning during its useful life whereby the failure times are exponentially distributed with

$$F(t) = 1 - e^{-\lambda t}, \quad \lambda > 0, \quad t \geq 0 \quad (25)$$

Therefore the counting distribution for the number of failures in $(0, t)$ is Poisson distributed with rate $\lambda > 0$ and hence $M(w) = \lambda w$ and in equation (22)

$$C(w; \phi) = c_0 + c_1 \lambda w \quad (26)$$

For non-Poisson failures, which are more typical, it will be necessary to compute $M(w)$.

Example 4 An item is produced at a cost of \$100 and sold with a complete warranty under an FRW policy for w months during which the average unit repair cost to the manufacturer is \$30. The **probability density function** for the time between warranty failures is

$$f(t) = \frac{16}{9} t e^{-\frac{4}{3}t}, \quad t \geq 0 \quad (27)$$

which is the Erlang probability density function. It can be shown (e.g., see [4, Appendix C]) that

$$M(t) = \frac{2}{3}t - \frac{1}{4} + \frac{1}{4}e^{-\frac{4}{3}t}, \quad t \geq 0 \quad (28)$$

Substituting the costs $c_0 = 100$ and $c_1 = 30$, and $M(w)$ from equation (28), the expected cost from

equation (23) is given by

$$C(w; FRW) = \frac{185}{2} + 20w + \frac{15}{2}e^{-\frac{4}{3}w}, \quad w \geq 0 \quad (29)$$

Computations of $M(w)$ often require numerical methods for solution. For cases where failure rates are relatively small (i.e., $P(N(w) > 1) \approx 0$), $M(t)$ can be approximated by $F(t)$.

Renewing Warranties. For the case of a renewing warranty, the length of the warranty period is random, since it renews itself to an additional warranty of length w each time a failure occurs at a time $t < w$. So the warranty period is sum of independent random variables. Denoting the time to warranty i by T_i , $i = 1, 2, \dots$, the time to a failure n is then, $S_n = T_1 + T_2 + \dots + T_n$. It follows that for the number of failures in $(0, w)$, $N(w) = n$, then $S_{n+1} > w$. Therefore,

$$P[N(w) = 0] = P[S_1 > w] = 1 - F(w) \quad (30)$$

and

$$\begin{aligned} P[N(w) = 1] &= P[S_2 > w] \\ &= P[T_1 < w]P[T_2 > w] \\ &= [1 - F(w)]F(w) \end{aligned} \quad (31)$$

and further by induction,

$$\begin{aligned} P(N(w) = n) \\ = [1 - F(w)][F(w)]^{n-1}, \quad n = 1, 2, \dots \end{aligned} \quad (32)$$

which is the geometric distribution with mean given by

$$E[N(w)] = \frac{1}{1 - F(w)} \quad (33)$$

It therefore follows from equation (23) that

$$C(w; \phi) = c_0 + \frac{E[X_1(w)]}{1 - F(w)} \quad (34)$$

and since in general for items under renewing warranty $c_1 = c_0$, for an FRW policy from equation (2) the total expected cost is given by

$$C(w; FRW) = \frac{c_0}{1 - F(w)} \quad (35)$$

For the PRW policy with renewing warranty, from equations (5) and (22),

$$C(w; PRW) = \frac{c_0 + \frac{c_0}{w} \int_0^w t f(t) dt}{1 - F(w)} \quad (36)$$

Example 5 Suppose that the turbine engine of Example 1 is under a renewing PRW policy. With $c_0 = \$1000$; $w = 15\,000$ h; and $\lambda = 2 \times 10^{-5}$ failures per h; substituting into equation (36)

$$\begin{aligned} C(15\,000; PRW) \\ = \frac{1000 + \frac{1000}{15\,000} \int_0^{15\,000} (2 \times 10^{-5}) t e^{(-2 \times 10^{-5})t} dt}{e^{(-2 \times 10^{-5})w}} \\ = \$1516 \end{aligned} \quad (37)$$

For the comparable fixed warranty the unit cost of failure is \$382.30 from Example 1, and substituting this value into equation (26) we have for the total cost,

$$\begin{aligned} C(w; PRW) &= 1000 + (383.30)(2 \times 10^{-5})(15\,000) \\ &= \$1,115 \end{aligned} \quad (38)$$

which, as expected, is much less cost to the manufacturer. Obviously, renewing warranties are normally not used for products with costs and failure characteristics like these.

Determining Warranty Reserve

Warranty reserve is the amount of money the manufacturer should set aside to cover forthcoming costs for warranty expenses. Proper warranty reserve management can be critical since too much reserve can result in lost opportunity for investing in other sources of revenue generation, and under **estimation** can result in a false profit plan. The following method first proposed by Menke [5] and later generalized by others [6, 7] is based on predicting reserve as a fraction of the total sales of a product.

Consider a product that is produced and sold in a lot of size K units at price c_1 dollars per unit under a PRW policy over the period $(0, w)$. The failure time probability density function $g(t)$, $t \geq 0$, is presumed to be known. The unit warranty cost from

6 Warranty Modeling

equation (5) is

$$X(t) = \begin{cases} c_0 \left(\frac{w-t}{w} \right), & 0 < t \leq w \\ 0, & t > w \end{cases} \quad (39)$$

Denoting the rate of return and expected price change per period by θ and π respectively, the future amount of reserve is given by

$$d(R) = X(t)d(K \cdot G(t)) \quad (40)$$

or

$$d(R) = c_1 \left(1 - \frac{t}{w} \right) (1 + \theta)^{-t} (1 + \pi)^{-t} \quad (41)$$

For $0 < \theta < 1$ and $0 < \pi < 1$,

$$(1 + \theta)^{-t} (1 + \pi)^{-t} \simeq (1 + \theta + \pi)^{-t} = b^{-t} \quad (42)$$

and it follows that

$$R(w) = \left(\frac{Kc_1}{w} \right) \int_0^w (w-t)b^{-t}g(t)dt \quad (43)$$

Example 6 Consider an example from [4] where an item sold under a 1-year PRW policy has failure times that are exponentially distributed at a rate of 0.4 failures per year. Interest is 5% per period and price changes at a rate of 4% per period.

Substituting the exponential probability density function into equation (43)

$$\begin{aligned} R(w) &= \left(\frac{Kc_1}{w} \right) \int_0^w (w-t)b^{-t}\lambda e^{-\lambda t} dt \\ &= Kc_1 \int_0^w e^{-t(\ln b + \lambda)} dt \\ &\quad - \left(\frac{Kc_1\lambda}{w} \right) \int_0^w te^{-t(\ln b + \lambda)} dt \end{aligned} \quad (44)$$

On simplifying these two terms it can be shown that

$$R(w) = \frac{Kc_1\lambda}{\ln b + \lambda} \left[1 - \frac{1 - b^{-w}e^{-\lambda w}}{w(\ln b + \lambda)} \right] \quad (45)$$

For $\lambda = 0.4$ failures per year, $w = 1$ year, $\theta = 0.05$, $\pi = 0.04$, and hence $b = 1.09$,

$$\begin{aligned} R(1) &= \frac{Kc_1(0.4)}{\ln(1.09) + 0.4} \left[1 - \frac{1 - (1.09)^{-1}e^{-0.4(1)}}{1(\ln(1.09) + 0.4)} \right] \\ &= 0.3669Kc_1 \end{aligned} \quad (46)$$

therefore $R(1)/Kc_1 = 0.3669$. This indicates that for each item sold, 37% of the unit cost should be set aside to cover warranty costs.

Optimum Warranty Policies

The determination of optimum warranty coverage requires knowledge of the product failure characteristics, the cost for repairs and replacements, and some measure of the benefit to the manufacturer for providing the warranty. Warranty benefit is indeed difficult to quantify but there are indirect means of capturing relative effects through marketing, consumer surveys, and other studies that can be used to estimate benefit. Here we will appeal to the argument that benefit will vary in a decreasing manner with the cost of marketing and advertising [8].

Denoting the expected benefit by $B(w)$, and the expected cost by $C(w)$, $w \geq 0$ the expected total cost is given by

$$T(w) = A + B(w) + C(w) \quad (47)$$

where A is the fixed administrative costs for the warranty program. $C(w)$ is an increasing function of w , whereas $B(w)$ will decrease with increasing warranty intervals. We assume that $B(w)$ and $C(w)$ are monotonic and $T(w)$ is convex. Suppose we consider an FRW policy with quadratic benefit given by

$$B(w) = (b_0 - b_1w)^2, \quad 0 \leq b_0/b_1 \leq w \quad (48)$$

The total expected cost is then given by

$$\begin{aligned} T(w) &= (A + c_0) + c_1M(w) + (b_0 - b_1w)^2, \\ 0 &\leq b_0/b_1 \leq w \end{aligned} \quad (49)$$

To find the optimum value of w ,

$$\frac{d}{dw}T(w) = c_1m(w) + 2(b_0 - b_1w)(-b_1) \quad (50)$$

where $m(w) = d/dwM(w)$ is the renewal density function and it follows that the optimum choice of warranty length w^* will satisfy the equation

$$c_1m(w) + 2b_1^2w - 2b_0b_1 = 0 \quad (51)$$

For the case of exponentially distributed failure times, the number of failures in $(0, w)$ is Poisson distributed

and $M(w) = \lambda w$, thus from equation (51) we have the optimum warranty interval given by

$$w^* = \frac{2b_0b_1 - c_1\lambda}{2b_1^2} \quad (52)$$

Example 7 A nonreparable item costs the manufacturer \$500 and failure times are distributed exponential with a rate of 0.75 failures per year. Without a warranty it is estimated that the manufacturer will have to incur a cost of \$22 500 to market the item, and with warranty the cost would decline as $(150 - 15w)^2$. So $c_1 = 500$, $\lambda = 0.75$, $b_0 = 150$, and $b_1 = 15$, and in equation (52)

$$w^* = \frac{2(150)(15) - 500(0.75)}{2(15)^2} = 9.2 \text{ years} \quad (53)$$

Concluding Remarks

The models presented in this paper are useful in predicting costs and addressing warranty economic trade-offs for developing policies. The FRW and PRW policies considered are fundamental and the most common policies that are applied to consumer products. Details on other warranty policies can be found in [9] and [10], and for an extensive review of economic issues pertaining to warranty analysis and planning decisions see [2] and [11].

The development of quantitative methods for warranty decisions has evolved through predicting and accounting for the costs for examining marketing strategies, developing and analyzing warranty policies, and most recently methods for integrating warranty feedback information for quality improvement. The modern view of quality is much broader now that the customer base is much more informed and critical of the value of their purchases. Quality is determined by a set of several elements that characterize the way a product is designed, developed, produced, and accepted by its users. Included among these elements will be some measure of performance capability, conformance quality of how well the design and manufacturing standards are achieved, product life, and other measures that pertain to customer acceptance. Since the cost and frequency of warranty claims are affected by all of these elements of quality, warranty performance provides an indirect measure for tracking the relative quality of products. More

details on using warranty information feedback and to develop quality improvement strategies can be found in [4].

End Notes

^a The views expressed in this article are those of the author and do not reflect the official policy or position of the Air Force, Department of Defense, or the US Government.

References

- [1] Loomba, A.P.S. (1996). *Historical Perspective on Warranty, Product Warranty Handbook*, W.R. Blischke & D.N.P. Murthy, eds, Marcel Dekker, New York, Chapter 2.
- [2] Thomas, M.U. & Rao, S.S. (1999). Warranty economic decision models: a summary and some suggested directions for future research, *Operations Research* **47**(6), 807–820.
- [3] Ross, S.M. (1983). *Stochastic Processes*, John Wiley & Sons, New York.
- [4] Thomas, M.U. (2006). *Reliability and Warranties: Methods for Product Development and Quality Improvement*, Taylor & Francis Group, Boca Raton, Chapter 4.
- [5] Menke, W.W. (1969). Determination of warranty reserve, *Management Science* **15**(4), B542–B549.
- [6] Amato, H.N. & Anderson, E.E. (1976). Determination of warranty reserves: an extension, *Management Science* **22**(12), 1391–1394.
- [7] Thomas, M.U. (1989). A prediction model for manufacturer warranty reserves, *Management Science* **35**(12), 1515–1519.
- [8] Thomas, M.U. (1983). Optimum warranty policies for nonreparable items, *IEEE Transactions on Reliability* **R-32**(3), 282–288.
- [9] Blischke, W.R. & Murthy, D.N.P. (1993). Product warranty management I: a taxonomy of warranty policies, *European Journal of Operational Research* **62**, 127–148.
- [10] Blischke, W.R. & Murthy, D.N.P. (1993). *Warranty Cost Analysis*, Marcel Dekker, New York.
- [11] Thomas, M.U. (2005). Engineering economic decisions and warranties, *Engineering Economist* **50**(4), 307–326.

Related Articles

Multiattribute Warranty Policies; Warranty Claims and Costs; Statistical Analysis of; Warranty Cost Analysis; Warranty Cost Prediction Based on Warranty Data; Warranty Servicing.

MARLIN U. THOMAS

Inspection Policies for Reliability

Introduction

Regular inspections can reduce the risk of equipment failure and so increase reliability. We may think of industrial plants, production lines, pipelines and other infrastructure such as buildings and roads, or merely of the annual inspection of domestic gas boilers, or the regular servicing of our cars. In healthcare, dental inspections and breast cancer screening (*see **Early Detection Programs in Medical Care***) reduce the risk of painful or life-threatening outcomes.

In every field, there are detailed engineering-based requirements of what inspections should look for, and how they should be carried out. Reliability-centered maintenance gives a general strategy. This article is concerned rather with inspection *policy*. When inspections are performed periodically, we ask how often they should be performed. Infrequent inspection does not prevent expensive failures, while too frequent inspection is costly, and may require extensive **system downtime**. We hope to find the golden mean.

All inspection models necessarily include parameters, and one can find the optimum policy to guarantee a specified reliability or to minimize cost per unit time as a function of the parameter values. Without **estimation** of these values, the modeling is useless. Operational researchers often content themselves with deriving an optimal policy given a model with known parameter values. This ignores two crucial steps necessary for practical use of the model: checking that the model is appropriate, and estimating its parameter values. One should also carry out a sensitivity analysis.

Two extreme cases for modeling may be distinguished. Equipment such as a production line is in use all the time, while standby equipment such as an alarm system, a fire extinguisher, an electrical relay, a life vest or an airbag “waits to work”. The item may have failed and be nonfunctional. Inspection reveals that the item is in a failed state. The simplest policy is to replace failed items, although it could be that a “minimal repair” would be carried out, to restore the item to a working state that is not however “good as new”.

When the system is in constant use, clearly an inspection cannot simply reveal that it is in a failed state—it was still functioning when shut down for inspection. The modeling of the benefit of an inspection must be more subtle. The inspection can reveal that the equipment is at risk of failure soon, that is, before the next scheduled inspection. Maybe the pipeline has worn dangerously thin in spots, the motorway bridge is cracked, or a tooth is decaying. Remedial action must be taken to prevent failure. Some components of a system may be known to need regular replacement, in which case the inspection is also an overhaul, as in replacing the oil and the air filters in a car. There are also intermediate cases, for example a 2-component system where one is in “warm standby”, and the system fails if both components fail.

Some mathematical models of inspection are essentially “wear models” (*see **Damage Processes***), that model the thickness of a pipeline or the growth of a crack. If we can specify an acceptable minimum reliability, or assign a cost to failure, we can find the total cost of inspections plus unscheduled failures, and so find the optimum policy, for example “inspect every x years and replace pipe segments with spots less than y -mm thick”. These models are naturally application-specific. Baker [1] gives this type of model for breast cancer screening. The processes of tumor origination and growth are modeled, and the probability of tumor detection at screening as a function of tumor size. Finally, presentation with cancer and subsequent survival are modeled as a function of tumor size. Dynamic programming is used to find the optimal screening ages, with life-years as the criterion to be maximized (*see **Dynamic Programming Methods in Repair and Replacement***).

In contrast, a class of models that can be applied to a wide variety of situations is Christer’s delay-time model. Defects arise randomly (in a Poisson process, *see **Poisson Processes***) and lead to failure after a random time, the delay time. This is the sojourn time in the defective state, and the probability distribution of sojourn time need not be exponential. The function of inspection is to detect and eliminate defects before they cause failure. Figure 1 illustrates this. This is an attractively simple model, because the condition of an item is either OK or defective, and exactly what constitutes a defect is left as an operational decision for maintenance engineers.

2 Inspection Policies for Reliability

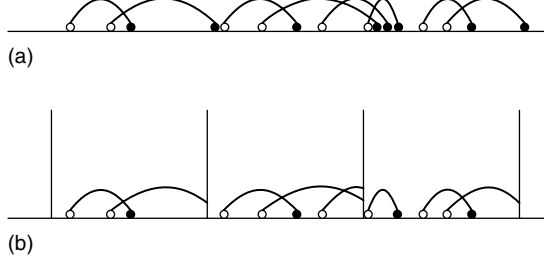


Figure 1 How inspections prevent failures for the delay-time model. The horizontal axis represents time and the open circles represent the origination of defects, the closed circles represent failures and the vertical lines inspections. With periodic inspections, as in (b) of the figure, the second, fourth, fifth and eighth defects have now been detected and repaired, and so have not caused failures

Thomas *et al.* [2] give a taxonomy of inspection models. It is not of course possible to give an exhaustive account of inspection models here, but the following sections give simple examples of policies for standby and continuous operation systems, together with some useful references.

Standby Systems

The simplest scenario is that we wish to increase reliability by replacing failed items such as life vests at regular intervals T . We assume that inspection and replacement take negligible time. For $nT < t < (n+1)T$ the availability (*see System Availability*) $A(t)$ is given by

$$A(t) = \sum_{i=0}^n p_i S(t - iT) \quad (1)$$

where $S(t)$ is the survival function of the failure time distribution of the item, p_i is the probability that the item was replaced at time iT , and $p_0 = 1$. Since $A\{(n+1)T\} = 1 - p_{n+1}$, equation (1) gives the recursive scheme

$$p_{n+1} = 1 - \sum_{i=0}^n p_i S(t - iT) \quad (2)$$

for the calculation of the replacement probabilities.

Figure 2 shows the availability of an item under such a scheme. The process soon settles down to its asymptotic state, where the p_i are constant, so that $p_i \rightarrow p$. Equations (1) and (2) then yield

$$p = \frac{1}{\sum_{i=0}^{\infty} S(iT)}$$

$$A(nT + u) \longrightarrow \frac{\sum_{i=0}^{\infty} S(u + iT)}{\sum_{i=0}^{\infty} S(iT)} \quad (3)$$

where $0 \leq u < T$. For an exponential distribution of failure time, with mean μ , $A(nT + u) = \exp(-u/\mu)$ for all n .

The availability decreases from 1 just after an inspection to $1 - 1/\sum_{i=0}^{\infty} S(iT)$ just before the next.

Choice of a suitable inspection interval would proceed by graphing $A(t)$ as in Figure 2 for various inspection intervals, and choosing one for which minimum availability and inspection cost was acceptable. There is likely to be little data available on item lifetime, but after some experience Weibull parameters could be estimated by a maximum likelihood fit to replacement data. The maximum likelihood method can cope with the interval censoring arising from our knowing only that the item failed during some period T , and the censoring arising from some items not having failed at all when the data came to hand. For example, if m_i items were regularly inspected and found functioning after $i - 1$ periods but failed after i periods, and $n_i - m_i$ items had not failed after their latest inspection at the i th period, the

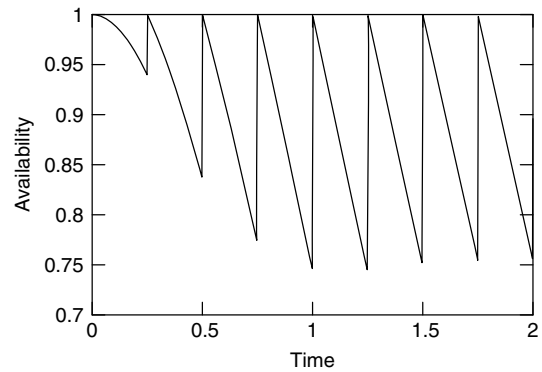


Figure 2 Availability for a standby system with inspection period $T = 1/4$ and replacement of failed items, where the failure time follows a Weibull distribution with survival function $S(t) = \exp(-t^2)$

likelihood function to be maximized for the Weibull parameters would be

$$\mathcal{L} = \prod_{i=1}^{\infty} \{S((i-1)T) - S(iT)\}^{m_i} S(iT)^{n_i - m_i} \quad (4)$$

When costs can be assigned, let c_1 be the cost of inspection, c_2 the cost of replacement, and c_3 the cost of nonavailability of the item per unit time. Then the mean cost C per unit time is given by

$$C(T) = \frac{c_1 + c_2 p + c_3 \int_0^T (1 - A(nT + u)) du}{T}$$

which reduces to

$$C = \frac{c_1 + \frac{c_2 - c_3 \mu}{\sum_{i=0}^{\infty} S(iT)}}{T} + c_3 \quad (5)$$

The optimum period T is found by minimizing $C(T)$. From equation (5), a necessary condition for inspections to be worthwhile is that $c_3 \mu > c_2$, so that the expected gain from operating the item exceeds its cost.

A comprehensive reference to the nonstatistical aspects of inspection models for standby systems is Chapter 8 of [3], where an inspection policy for a standby electric generator is derived, when repair time is not negligible. The earliest work is by [4], who derived inspection times T_i to minimize the total cost of checking and replacing a unit, plus the cost of undetected failure, taken as proportional to the “down” time, over the life of the unit. Another important early reference is [5].

The T_i can also be chosen to minimize cost per unit time if the unit is replaced on failure over an infinite time horizon.

There are very many modeling variations; for example, inspection may be imperfect and detect a failed item with some probability $q < 1$, the item may be minimally repaired rather than replaced, and so on (see **Imperfect Repair**).

Continuous Operation Systems

This example of the use of the delay-time model is based upon [6], an early and influential case study for a high-speed canning line. It is further discussed in [7].

Inspection was carried out every $T = 24$ hours. In general, inspection was found to result on average in a downtime of $65.5/24 = 2.73$ hours per week. Assuming that the inspection downtime period is independent of inspection frequency, downtime due to inspection is $65.5/T$ per week. The mean number of defects appearing per week was measured as 16.6, and the fraction $b(T)$ of them which resulted in failures caused an average downtime of 0.698 h, so that the downtime per week due to failures was modeled as $11.59b(T)$. The downtime per week as a function of T was therefore

$$\frac{65.5}{T} + 11.59b(T)$$

This formula is approximate, and uses the fact that in this case manpower for inspection was not a problem: flexible manpower at inspection meant that an inspection and repair took the same time however many defects had to be fixed. To optimize T , we need to know b as a function of T .

Assuming a stationary process, a defect arising in the interval $(0, T)$ is assumed to arrive at time u from the last inspection with **pdf**. $g(u) = 1/T$. The probability that a random defect causes a failure before it is caught at the next inspection is

$$\begin{aligned} b(T) &= \int_0^T g(u) F(T - u) du \\ &= (1/T) \int_0^T (1 - \exp\{-\eta u\}) du \end{aligned} \quad (6)$$

assuming an exponential distribution for the delay time, when the distribution function $F(h) = 1 - \exp(-\eta h)$. Hence

$$b(T) = 1 - (1 - \exp\{-\eta T\})/\eta T \quad (7)$$

this is the required formula for $b(T)$.

The observed value of b was 0.396, for the current practice with $T = 24$ hours. Solving equation (7) gives $\eta = 0.0463$, which provides a delay-time distribution estimate that is tuned to current practice.

Figure 3 shows estimated downtime per week as a function of T , and it can be seen that daily inspection is about optimal.

There has been quite an industry in developing the delay-time model from this simple beginning. The concept of a two-stage failure process also arose at about the same time in the context of cancer

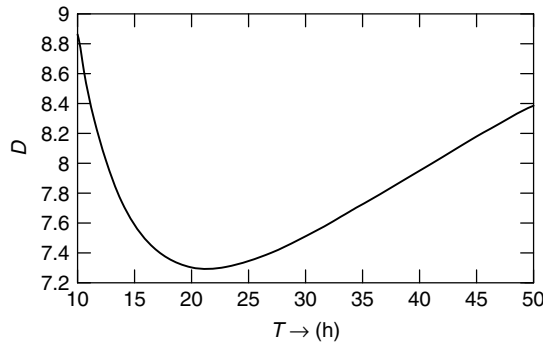


Figure 3 Downtime per week against the interval between inspections for a case study of a canning line from [6]

screening [8]. In Operational Research, Christer's delay-time model was then developed by extending it to individual components [9], making the model more complex (e.g., consideration of fault injection at maintenance [10]), and estimating parameters by maximum likelihood rather than by the method of moments [11, 12].

The estimation of model parameters is a major problem. In this simple example, the method of moments was used to equate observed variables to model predictions. Estimation requires some subtlety if the only data available is the historical record of failures and results found at inspection under some existing inspection policy that we hope to optimize. Likelihood-based statistical estimation can usually cope with problems of censoring that may arise, and some method of choosing the best model such as minimum Akaike information criterion (AIC) is needed. Moving to longer intervals between inspections sometimes requires extrapolation of the tail behavior of fitted distribution functions. For example, the delay-time distribution of a defect can never be known at times larger than the interval between inspections. We rely on the smoothness of the distribution function to extrapolate into the tail.

Experience was gained with a wide variety of applications, for example, building maintenance [13], civil engineering structures [14], industrial plant [15–17], gearboxes [18], pumps [19], concrete structures [20] and much else. Difficulties in obtaining adequate data for making good model predictions led to interest in more efficient statistical inference (e.g. [21]), and also to the elicitation of subjective data (see **Expert Opinion in Reliability**) [17, 22, 23].

Bayesian methods are nowadays increasingly being used by statisticians, and where there is insufficient “objective” data, it can be eked out with subjective data gleaned from maintenance engineers. A Bayesian approach to this model in the context of cancer screening is given in [24]. Reviews of delay-time modeling are found in [21, 25, 26].

References

- [1] Baker, R.D. (1998). What can industrial inspection models learn from medical analogues? Determining the optimum screening policy for breast cancer, *IMA Journal of Mathematics in Business and Industry* **9**, 305–324.
- [2] Thomas, L.C., Gaver, D.P. & Jacobs, P.A. (1991). Inspection models and their application, *IMA Journal of Mathematics Applied in Business and Industry* **3**, 283–303.
- [3] Nakagawa, T. (2005). *Maintenance Theory of Reliability*, Springer, London, pp. 201–229.
- [4] Barlow, R.E. & Proschan, F. (1965). *Mathematical Theory of Reliability*, John Wiley & Sons, New York.
- [5] Keller, J.B. (1973). Optimum checking schedules for systems subject to random failure, *Management Science* **21**, 256–260.
- [6] Christer, A.H. & Waller, W.M. (1984). Delay time models of industrial maintenance problems, *Journal of the Operational Research Society* **35**, 401–406.
- [7] Christer, A.H. & Redmond, D.F. (1992). Revising models of maintenance and inspection, *International Journal of Production Economics* **24**, 227–234.
- [8] Day, N.E. & Walter, S.D. (1984). Simplified models of screening for chronic disease: estimation procedures from mass screening programmes, *Biometrics* **40**, 1–14.
- [9] Christer, A.H. & Wang, W. (1995). A delay-time based maintenance model of a multi-component system, *IMA Journal of Mathematics in Business and Industry* **6**, 205–222.
- [10] Carr, M.J. & Christer, A.H. (2003). Incorporating the potential for human error in maintenance models, *Journal of the Operational Research Society* **54**, 1249–1253.
- [11] Baker, R.D. & Wang, W. (1991). Estimating the delay-time distribution of faults in repairable machinery from failure data, *IMA Journal of Mathematics Applied in Business and Industry* **3**, 259–281.
- [12] Baker, R.D. & Wang, W. (1993). Developing and testing the delay-time model, *Journal of the Operational Research Society* **44**, 361–374.
- [13] Christer, A.H. (1982). Modelling inspection policies for building maintenance, *Journal of the Operational Research Society* **33**, 723–732.
- [14] Christer, A.H. (1988). Condition-based inspection models of major civil-engineering structures, *Journal of the Operational Research Society* **39**, 71–82.
- [15] Christer, A.H. (1991). Prototype modelling of irregular condition monitoring of production plant, *IMA Journal*

- of *Mathematics Applied in Business and Industry* **3**, 219–232.
- [16] Christer, A.H., Wang, W., Sharp, J. & Baker, R.D. (1997). Stochastic maintenance modelling of hi-tech steel production plant, *Stochastic Modelling in Innovative Manufacturing, Lecture Notes in Economics and Mathematical Systems*, Vol. 445, Springer-Verlag, Berlin, pp. 196–214.
 - [17] Christer, A.H., Wang, W., Sharp, J. & Baker, R.D. (1998). A case study of modelling preventive maintenance of a production plant using subjective data, *Journal of the Operational Research Society* **49**, 210–219.
 - [18] Leung, F. & Mak, K.L. (1996). Using delay-time analysis to study the maintenance problem of gearboxes, *International Journal of Operations and Production Management* **16**, 98–105.
 - [19] Leung, F.K.N. & Ma, T.W. (1997). A study on the inspection frequency of fresh water pumps, *International Journal of Industrial Engineering: Applications and Practice* **4**, 42–51.
 - [20] Redmond, D.F., Christer, A.H., Rigden, S.R., Burley, E., Tajelli, A. & Abutair, A. (1997). OR Modelling of the deterioration and maintenance of concrete structures, *European Journal of Operational Research* **99**, 619–631.
 - [21] Baker, R.D. (1996). Maintenance optimisation with the delay-time model, in *Reliability and maintenance of Complex Systems*, S. Özekici, ed, (proceedings of the 1995 NATO ASI, Antalya, Turkey), Springer-Verlag, Berlin, pp. 550–587.
 - [22] Wang, W. (1997). Subjective estimation of the delay time distribution in maintenance modelling, *European Journal of Operational Research* **99**, 516–529.
 - [23] Desa, M.I. & Christer, A.H. (2001). Modelling in the absence of data: a case study of fleet maintenance in a developing country, *Journal of the Operational Research Society* **52**, 247–260.
 - [24] Parmigiani, G. (1993). On optimal screening ages, *Journal of the American Statistical Association* **88**, 622–628.
 - [25] Baker, R.D. & Christer, A.H. (1994). Review of delay-time OR modelling of engineering aspects of maintenance, *European Journal of Operational Research* **73**, 407–422.
 - [26] Christer, A.H. (2002). A review of delay time analysis for modelling plant maintenance, in *Stochastic Models in Reliability and Maintenance*, S. Ozaki, ed, Springer, Berlin.

Related Articles

Damage Processes; Dynamic Programming Methods in Repair and Replacement; Early Detection Programs in Medical Care; Expert Opinion in Reliability; General Minimal Repair Models; Imperfect Repair; Maintenance Optimization; Poisson Processes; Replacement Strategies; System Availability.

ROSE D. BAKER

Total Time on Test Plots

Introduction

Suppose that we have a complete ordered sample $0 = t(0) \leq t(1) \leq \dots \leq t(n)$ of times to failure from n identical and independent nonrepairable units with life distribution $F(t)$ and survival function $R(t) = 1 - F(t)$. The *TTT* (total time on test) *plot* of these observations is then obtained in the following way

- Let $S_0 = 0$ and calculate the TTT values S_j as
- $S_j = S_{j-1} + (n - j + 1)(t(j) - t(j - 1))$ for $j = 1, 2, \dots, n$.
- Normalize these TTT values by calculating $u_j = S_j/S_n$ for $j = 0, 1, \dots, n$.
- Plot $(j/n, u_j)$ for $j = 0, 1, \dots, n$.
- Join the plotted points by line segments.

Accordingly, the TTT plot consists of line segments starting in the lower left corner and ending in the upper right corner of the unit square. When the sample size n increases to infinity the TTT plot converges to the *scaled TTT transform* of the life distribution $F(t)$ from which the sample has come

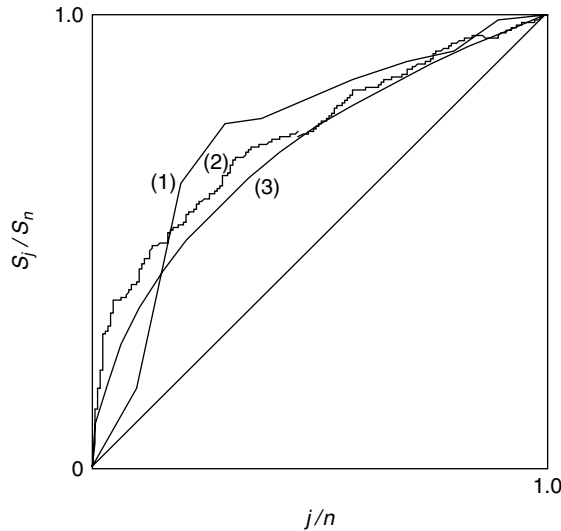


Figure 1 TTT plots based on simulated data from a Weibull distribution with $\beta = 2.0$, $n = 10$ in (1) and $\beta = 2.0$, $n = 100$ in (2) and the scaled TTT transform of a Weibull distribution with shape parameter $\beta = 2.0$ in (3). (From Bergman and Klefsjö [2])

[1]. An illustration of different TTT plots and the convergence is found in Figure 1. The scaled TTT transform, defined as

$$\varphi(u) = \frac{1}{\mu} \int_0^{F^{-1}(u)} R(t) dt \text{ for } 0 \leq u \leq 1 \quad (1)$$

where μ is the mean, is, as the name indicates, independent of scale. For instance, for a Weibull distribution with survival function $R(t) = \exp(-(t/\theta)^\beta)$, $t \geq 0$, the transform depends only on β and is independent of the value of θ . Furthermore, every exponential distribution $F(t) = 1 - \exp(-\lambda t)$, $t \geq 0$, is transformed to the diagonal in the unit square, independently of the value of λ ; see curve (3) in Figure 2. Accordingly, different deviations from the diagonal in the unit square of the scaled TTT transform of a life distribution $F(t)$ can be interpreted as different deviations of $F(t)$ from the exponential distribution.

The TTT plot and the scaled TTT transform were introduced by Barlow and Campo [1] primarily for model identification purposes. By comparing the TTT plot based on the obtained times to failure

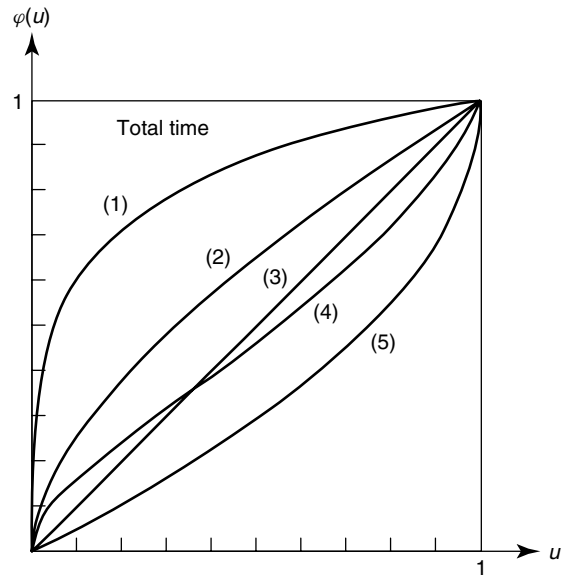


Figure 2 Scaled TTT transforms from five different life distributions: (1) normal with $\mu = 1$, $\sigma = 0.3$; (2) gamma with shape parameter 2.0; (3) exponential distribution; (4) lognormal with $\mu = 0$, $\sigma = 1$; (5) Pareto distribution with survival function $R(t) = (1 + t)^{-2}$, $t \geq 0$ [© John Wiley & Sons Ltd, New York.]

2 Total Time on Test Plots

with different scaled TTT transforms it is possible to choose a suitable life distribution model. Since then several other applications and generalisations of TTT plotting have appeared, both theoretical and practical. The aim of this presentation is to give an overview of a number of these practical and theoretical applications and support the reader with suitable references for further information.

TTT Plotting for Aging Classification

Different aging properties of $F(t)$ can be translated to properties of the scaled TTT transform of $F(t)$; see Barlow and Campo [1], Barlow [4], Bergman [5], Klefsjö [6]. This gives a source of subjective criteria of different aging properties; see Figure 3. As one important example, a life distribution has *Increasing Failure Rate (IFR)* exactly when the scaled TTT transform is concave [1]. If a unit has IFR “the proneness” to fail increases with the age, which means that there is a certain age at which the unit should be replaced with a new one to minimize the “long-run costs”, e.g., measured as the average long-run cost per unit time. To estimate the optimal replacement age according to this criteria the TTT

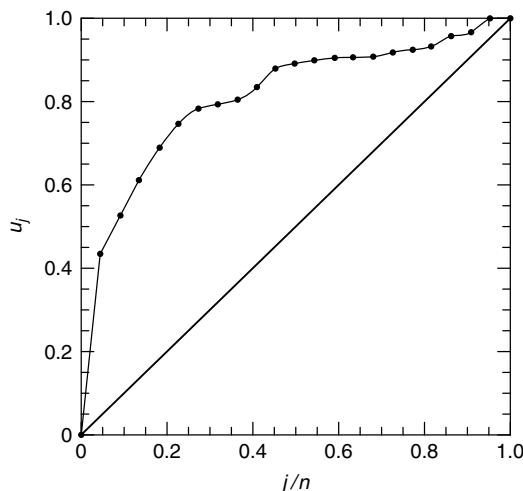


Figure 3 Since $F(t)$ has increasing failure rate (IFR) exactly when the scaled TTT transform $\varphi(u)$ is concave we expect a TTT plot based on failure times from an IFR life distribution to behave concavely. This is illustrated here by a TTT plot based on 22 failure times from a tractor engine [© John Wiley & Sons Ltd, WileyInterscience.]

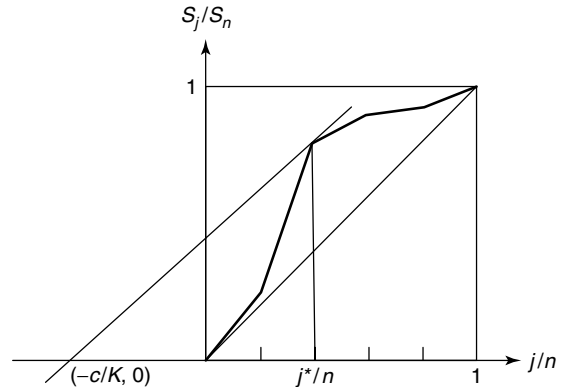


Figure 4 Illustration of estimation of the optimal replacement age based on the TTT plot when the replacement cost is c at age T and K at failure. The failure time $t(j^*)$ corresponding to the touching point at the TTT plot of the line through $(-c/K, 0)$ with the largest slope, is the estimator [© John Wiley & Sons Ltd, New York.]

plot is useful as well [3, 7, 8]; for an illustration see Figure 4.

More discussions related to the correspondence between TTT transforms and aging properties can be found in e.g., Klefsjö [6] and Bergman [5]. Illustrations of how to use the TTT plot and the scaled TTT transform for optimisation related to the age replacement problem with discounted costs are presented by Bergman and Klefsjö [9] and Dohi *et al.* [10], related to **burn-in** problems by Bergman and Klefsjö [11] and when discussing replacements to extend system life by Bergman and Klefsjö [12]. Some interesting papers dealing with problems from practice are Kumar *et al.* [13], Kumar and Westberg [14], Bergman and Klefsjö [15] and Al-Najjar [16].

TTT Plotting When Analyzing Data from Repairable Systems

Suppose that we are studying a single repairable unit. The data normally provided consists of epochs of failure events, measured on an operation timescale or a calendar timescale. In some cases the duration of repair is reported as well. However, in this paper repair times are excluded from the discussion. The sequence of event epochs, i.e., the times to failures, $T_1, T_2, \dots$ give rise to a sequence of interevent times, or interfailure times, $t_j = T_j - T_{j-1}$, also called *local times*. This means that the time to the j th failure T_j is equal to $T_j = t_1 + t_2 + \dots + t_j$.

When studying repairable units, the failure intensity at time T , which intuitively is a measure of the probability that a failure will occur during the next units of time, is of great interest [17]. The intensity of the failure process can be a function of both the running time T , also called *global time*, and the local time t .

When the intensity of the failure process is a function of the running time T only, the failures may occur according to a nonhomogeneous **Poisson process**. This situation is obtained if, for instance, we use a **minimal repair** policy, which means that the unit is repaired to the same state as it had just before the failure occurred. In this situation, the usual assumption of a renewal process to model such data is most often incorrect. For an excellent discussion on analysis of failure data from **repairable systems** and the abuse of the renewal process assumption, see Ascher and Feingold [17].

One of the most popular nonhomogeneous Poisson process models is the one where the **intensity function** $\rho(T)$ is a power of T , i.e., it can be written in the form

$$\rho(T) = \frac{\beta}{\theta} \left(\frac{T}{\theta} \right)^{\beta-1} \quad (2)$$

with constants $\theta > 0$ and $\beta > 0$. This model has been discussed under several different names, many of them including “Weibull” in the name, since the intensity function above has the same form as the failure rate of a Weibull life distribution. However, including “Weibull” in the name of the process has caused a lot of confusion. It gives the incorrect impression that the times between failures follow a Weibull distribution (for a discussion, see [18]). Therefore, we prefer to use the name *power law process*. For an excellent discussion of the power law process, see Rigdon and Basu [19, 20].

Suppose we observe a power law process, with intensity function $\rho(T) = (\beta/\theta)(T/\theta)^{\beta-1}$, until n events have occurred. Suppose further that the events occurred at $T_1, T_2, \dots, T_n$. Møller [21] has shown that if $w_j = -\ln(T_{n-j}/T_n)$, then $w_1, w_2, \dots, w_{n-1}$ is distributed as an ordered sample of size $n-1$ from an exponential distribution with mean $1/\beta$, i.e., from the distribution $F(t) = 1 - \exp(-\beta t)$, $t \geq 0$.

This means that we can use an ordinary TTT plot based on $w_1 \leq w_2 \leq \dots \leq w_{n-1}$ to check a power law process assumption. If the power law process is a suitable model we expect the TTT plot to

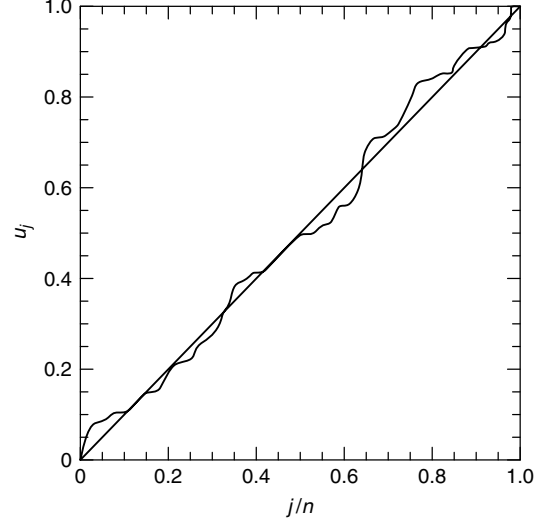


Figure 5 A TTT plot based on the w_j values from a simulated power law process with $\beta = 0.5$ and sample size $n = 52$ [Data from Crow and Basu [22], © IEEE, 1992.]

wiggle around the diagonal since the TTT plot is an estimate of the corresponding scaled TTT transform, see Figure 5.

Figure 6 illustrates the TTT plot based on the w_j values obtained from data simulated from a nonhomogeneous Poisson process with intensity function $\rho(T) = \exp(T-1)$, which is not a power law process. From Figure 6 it is seen that the data is not in agreement with a power law process assumption since the plot behaves convexly and deviates much from the diagonal.

If the TTT plot is based on the w_j values from a renewal process, it is hard to separate the plot from a plot based on data from a power law process [23]. However, this shortcoming can be eliminated by using the transform

$$d_j = (n-j)(w_j - w_{j-1}) \text{ for } j = 1, 2, \dots, n-1 \quad (3)$$

sometimes called the *Durbin transform*. Barlow and Proschan [24] have proven that if $w_1 \leq w_2 \leq \dots \leq w_{n-1}$ is an ordered sample from an exponential distribution then also $d_1, d_2, \dots, d_{n-1}$ is an independent sample from an exponential distribution. This means that if we have times to failures $T_1, T_2, \dots, T_n$ from a power law process then the d_j values are independent observations from an exponential distribution. Accordingly, the TTT plot based on the d_j values

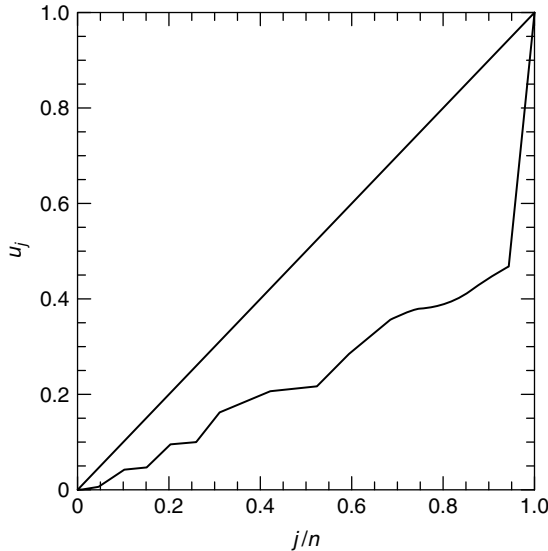


Figure 6 A TTT plot based on the w_j values from a simulated nonhomogeneous Poisson process with intensity function $\rho(T) = \exp(T - 1)$. This nonhomogeneous process is not a power law process and the plot is therefore expected to deviate from the diagonal [© IEEE, 1992.]

should wriggle around the diagonal in the unit square in the same manner as the plot based on the w_j values. See Klefsjö and Kumar [23] and Akersten *et al.* [25] for further discussion and illustration. Kvaløy and Lindqvist [26] also discuss trends in repairable systems data using the TTT concept.

Further TTT Plotting Applications

Extension of Power Law Process Testing

Another testing technique related to the power law process situation, presented by Akersten [27] is based on the fact that, with w_j as above, $w_1^\beta, w_2^\beta, \dots, w_n^\beta$ can be proven to be an ordered sample from a uniform distribution on $[0, 1]$. A consequence of this is that if the failures occur according to a power law process with power β , then the plot obtained by plotting and joining the points $(j/n, w_j^\beta)$, $j = 0, 1, \dots, n$ (where $w_0 = 0$), let us call it the w_j plot for short, is expected to twin around the diagonal in the unit square. By checking different values of β and choosing the one which gives the lowest value of a certain test statistic related to the cumulative TTT

statistic V [28], Akersten [27] obtains a possibility of both graphically and analytically testing a power law process assumption and also to estimate the value of β . Rigdon [29] also discusses test possibilities when studying power law processes.

In the case of a homogeneous Poisson process both the usual TTT plot and the w_j plot based on the points $(j/n, w_j)$ will tend to wriggle around the diagonal. If the failure history instead can be modeled by a nonhomogeneous Poisson process we expect the TTT plot based on the interevent times t_j to have a convex shape and the w_j plot to have a concave or convex shape depending on whether the intensity function is increasing or decreasing. If the TTT plot based on the interevent times has a concave shape, a nonhomogeneous Poisson process seems to be an incorrect assumption. See Akersten *et al.* [25] for more details.

Analysis of Test Statistics and Their Efficiency

The correspondence between the scaled TTT transform $\varphi(u)$ and different aging properties is a source for creating test statistics to analytically testing exponentiality against different aging properties and for studying the properties of these statistics. These issues are discussed e.g., by Klefsjö [30], Basu and Ebrahimi [31], Bergman and Klefsjö [32] and Bergman and Klefsjö [15]. For instance, Klefsjö [30] proved that the cumulative TTT statistic V introduced already by Barlow *et al.* [28] for testing exponentiality against IFR is consistent against all NBUE (New Better than Used in Expectation) life distributions, since V is a linear function of the area between the TTT plot and the diagonal. The TTT transform and TTT plot as an inspiration for efficiency calculations of test statistics is discussed in Bergman and Klefsjö [15]. In Klefsjö and Westberg [33] TTT transforms are used for a discussion of efficiency of a number of tests for exponentiality, among these Gnedenko's F-test [34].

TTT Plotting with Censored Data

The discussion above has focused on the situation when we have uncensored failure times. However, the TTT plotting positions j/n and u_j in the TTT plot for uncensored data can be expressed by using the empirical distribution function $F_n(t)$ and the empirical survival function $R_n(t) = 1 - F_n(t)$ as $(F_n(t_j), S_j/S_n)$;

Barlow and Campo [1]. This means that it is possible to generalize the TTT plotting technique to situations based on censored data by changing the empirical distribution function $F_n(t)$ to a suitable estimator of $F(t)$ based on the censored sample. One possibility then is to use the *Kaplan–Meier estimator (KME)*, whose properties were first studied by Kaplan and Meier [35]. Another one is to use the piecewise exponential estimator (PEXE) technique [36, 37]. A third idea is to use the *mean-order-number method* [38]. How to calculate and use TTT plotting with censored data is discussed e.g., in Westberg and Klefsjö [39], Westberg [40], Klefsjö and Westberg [33] and Sun and Kececioğlu [38].

Relation to the Lorenz Transform

Finally, we want to mention the close relation between the scaled TTT transform $\varphi(p)$ and the *Lorenz transform* $L(p)$, introduced by Lorenz [41] and frequently used in economics as a measure of concentration of wealth. The Lorenz transform is defined as

$$L(p) = \frac{1}{\mu} \int_0^p F^{-1}(s) ds \text{ for } 0 \leq p \leq 1 \quad (4)$$

which also is a curve within the unit square. Relations between the Lorenz curve and different aging properties of life distributions are discussed by e.g., Chandra and Singpurwalla [42], Klefsjö [43] and Perez-Ocun *et al.* [44]. Among other things, the *Gini Index (GI)*, often defined as twice the area between the egalitarian and the Lorenz transform, i.e.,

$$GI = 2 \int_0^1 (p - L(p)) dp \quad (5)$$

can be interpreted graphically as the area in the unit square above the scaled TTT transform. Applications of the Lorenz transform to repair problems are discussed by Dohi *et al.* [45].

TTT Plotting for Finding Active Effects in Design of Experiments

Bergman and Sandvik Wiklund [46] use the TTT plot to find active factors in **design of experiments**. A formal test is suggested based on the cumulative

TTT statistic V , see Barlow *et al.* [28], based on the squared contrasts.

Future Development

The TTT concept support some intuitive and insightful interpretations of life data and we foresee some future interesting applications.

Acknowledgments

The first authors would like to thank SKF AB and the Swedish Foundation for Strategic Research through Gothenburg Mathematical Modeling Center (GMMC) for financial support.

References

- [1] Barlow, R.E. & Campo, R. (1975). Total time on test processes and applications to failure data analysis, in *Reliability and Fault Tree Analysis*, R.E. Barlow, J. Fussell & N.D. Singpurwalla, eds, SIAM, Philadelphia, pp. 451–481.
- [2] Bergman, B. & Klefsjö, B. (1984). The total time on test concept and its use in reliability theory, *Operations Research* **32**, 596–606.
- [3] Bergman, B. & Klefsjö, B. (1988). Total time on test transforms, in *Encyclopedia of Statistical Sciences*, John Wiley & Sons, New York, Vol. 9, pp. 297–300.
- [4] Barlow, R.E. (1979). Geometry of the total time on test transform, *Naval Research Logistics Quarterly* **26**, 393–402.
- [5] Bergman, B. (1979). On age replacement and the total time on test concept, *Scandinavian Journal of Statistics* **6**, 161–168.
- [6] Klefsjö, B. (1982). On aging properties and total time on test transforms, *Scandinavian Journal of Statistics* **9**, 37–41.
- [7] Bergman, B. (1977). Some graphical methods for maintenance planning, in *Proceedings from 1977 Annual Reliability and Maintainability Symposium*, Philadelphia, pp. 467–471.
- [8] Bergman, B. & Klefsjö, B. (1982). A graphical method applicable to age replacement problems, *IEEE Transactions on Reliability* **31**(5), 478–481.
- [9] Bergman, B. & Klefsjö, B. (1983). TTT-transforms and age replacements with discounted costs. *Naval Research Logistics Quarterly*, **30**, 631–639.
- [10] Dohi, T., Kaio, N. & Osaki, S. (2003). A new graphical method to estimate the optimal repair-time limit with incomplete repair and discounting, *Computers and Mathematics with Applications* **46**(7), 999–1007.
- [11] Bergman, B. & Klefsjö, B. (1985). Burn-in models and TTT-transforms, *Quality and Reliability Engineering International* **1**, 125–130.

- [12] Bergman, B. & Klefsjö, B. (1987). The TTT-concept and replacements to extend system life, *European Journal of Operational Research* **28**, 302–307.
- [13] Kumar, U., Klefsjö, B. & Granholm, S. (1989). Reliability investigations for a fleet of load haul dump machines in a Swedish mine, *Reliability Engineering and System Safety* **26**, 341–361.
- [14] Kumar, D. & Westberg, U. (1997). Maintenance scheduling under age replacement policy using proportional hazards model and TTT-plotting, *European Journal of Operational Research* **99**(3), 507–513.
- [15] Bergman, B. & Klefsjö, B. (1998). Recent applications of TTT-plotting. Invited contribution to, in *Frontiers in Reliability*, A.P. Basu, ed, World Scientific, Singapore, pp. 47–61.
- [16] Al-Najjar, B. (2003). Total time on test, TTT-plots for condition monitoring of rolling elements in bearings in paper mills, *International Journal of COMADEM* **6**(2), 27–32.
- [17] Ascher, H. & Feingold, H. (1984). *Repairable Systems Reliability, Lecture Notes in Statistics*, Marcel Dekker, New York, Vol. 7.
- [18] Ascher, H. (1981). Weibull distribution versus Weibull process, in *Proceedings Annual Reliability and Maintainability Symposium 1981*, Philadelphia, pp. 426–429.
- [19] Rigdon, S.E. & Basu, A.P. (1989). The power law process – a model for the reliability of repairable systems, *Journal of Quality Technology* **21**, 251–260.
- [20] Rigdon, S.E. & Basu, A.P. (1990). The effect of assuming a homogeneous Poisson process when the true process is a power law process, *Journal of Quality Technology* **22**, 111–117.
- [21] Møller, S.K. (1976). The Rasch-Weibull process, *Scandinavian Journal of Statistics* **3**, 107–115.
- [22] Crow, L.H. & Basu, A.P. (1988). Reliability growth estimation with missing data, *II proceeding from Annual Reliability and Maintainability Symposium*, Los Angeles 248–253.
- [23] Klefsjö, B. & Kumar, U. (1992). Goodness-of-fit test for the power law process based on the TTT-plot, *IEEE Transactions on Reliability* **R-41**, 593–598.
- [24] Barlow, R.E. & Proschan, F. (1981). *Statistical Theory of Reliability and Life Testing*, To Begin With, Silver Spring.
- [25] Akersten, P.A., Klefsjö, B. & Bergman, B. (2001). Graphical techniques for analysis of data from repairable systems, in *Handbook of Statistics Volume 20*, N. Balakrishnan & C.R. Rao, eds, North-Holland, Elsevier, Amsterdam, Chapter 16, pp. 469–484.
- [26] Kvaløy, J. & Lindqvist, B.H. (1998). TTT-based tests for trend in repairable systems data, *Reliability Engineering and System Safety* **60**, 13–28.
- [27] Akersten, P.A. (1991). *Repairable systems reliability studied by TTT-plotting techniques*, Ph.D. dissertation, Division of Quality Technology, Linköping University of Technology, Sweden.
- [28] Barlow, R.E., Bartholomew, J.M., Bremner, J.M. & Brunk, H.D. (1972). *Statistical Inference under Order Restrictions*, John Wiley & Sons, New York.
- [29] Rigdon, S.E. (1989). Testing goodness-of-fit for the power law process, *Communications in Statistics Part A-Theory and Methods* **18**, 4665–4676.
- [30] Klefsjö, B. (1983). Some tests against aging based on the total time on test transform, *Communications in Statistics Part A-Theory And Methods* **12**, 907–927.
- [31] Basu, A.P. & Ebrahimi, N. (1985). Testing whether survival function is harmonic new better than used in expectation, *Annals of the Institute of Statistical Mathematics* **37**(1), 347–359.
- [32] Bergman, B. & Klefsjö, B. (1989). A family of test statistics for detecting monotone mean residual life, *Journal of Statistical Planning and Inference* **21**, 161–178.
- [33] Klefsjö, B. & Westerberg, U. (1996). Efficiency calculations of some tests for exponentiality using TTT-transforms, *Arab Journal of Mathematical Sciences* **2**, 129–149.
- [34] Gnedenko, B.V., Belyayev, Y.K. & Solov'yev, A.D. (1969). *Mathematical Methods of Reliability Theory*, Academic Press, New York.
- [35] Kapan, E.L. & Meier, P. (1958). Nonparametric estimation from incomplete observations, *Journal of the American Statistical Association* **53**, 457–481.
- [36] Kitchin, J. (1980). A new method for estimating life distributions from incomplete data, PhD dissertation, Florida State University.
- [37] Kitchin, J., Langberg, N. & Proschan, F. (1983). A new method for estimating life distributions from incomplete data, *Statistics* **1**, 241–255.
- [38] Sun, F.-B. & Kececioglu, D.B. (1999). A new method for obtaining the TTT-plot for a censored sample, in *Proceedings from Annual Reliability and Maintainability Symposium*, Washington DC, pp. 112–117.
- [39] Westberg, U. & Klefsjö, B. (1994). TTT-plotting for censored data based on the piecewise exponential estimator, *International Journal of Reliability, Quality and Safety Engineering* **1**, 1–13.
- [40] Westberg, U. (1994). Applications of the PEXE-concept for maintenance policies and proportional hazards model, Licentiate Thesis 1994:24L, Division of Quality Technology & Statistics, Luleå University, Sweden.
- [41] Lorenz, M.O. (1983). Methods of measuring the concentration of wealth, *Journal of the American Statistical Association* **9**, 209–219.
- [42] Chandra, M. & Singpurwalla, N.D. (1981). Relationships between some notions which are common to reliability and economics, *Mathematics of Operations Research* **6**, 113–121.
- [43] Klefsjö, B. (1984). Reliability interpretations of some concepts from economics, *Naval Research Logistics Quarterly* **31**, 301–308.
- [44] Perez-Ocon, R., Gamiz-Perez, M.L. & Ruin-Castro, J.E. (1997). Study of different ageing properties via total time on test transform and Lorenz curves, *Applied Stochastic Models and Data Analysis* **13**(3–4), 241–248.

-
- [45] Dohi, T., Othman, F.S., Kaio, N. & Osaki, S. (2001). The Lorenz transform approach to the optimal repair-cost limit replacement policy with imperfect repair, *RAIRO Operations Research* **35**, 21–36.
 - [46] Bergman, B. & Sandvik Wiklund, P. (1999). Finding active factors from un-replicated fractional factorials utilizing the total time on test (TTT) technique, *Quality and Reliability Engineering International*, **15**, 191–203.
 - Söderholm, P. (2004). Continuous improvements of complex technical systems: a theoretical quality management framework supported by requirements management and health management, *Total Quality Management and Business Excellence* **15**(4), 511–525.

Related Articles

Further Reading

- Kim, J.S. & Proschan, F. (1991). Piecewise exponential estimator of the survivor function, *IEEE Transactions on Reliability* **R-40**, 134–139.
- Klefsjö, B. (1986). TTT-transforms – a useful tool when analyzing different reliability problems, *Reliability Engineering* **15**, 231–241.

Burn-In and Maintenance Policies; General Minimal Repair Models; Intensity Functions for Non-homogeneous Poisson Processes; Maintenance Optimization; Poisson Processes; Replacement Strategies; Repairable Systems Reliability; Repairable Systems: Statistical Inference.

BO BERGMAN AND B. KLEFSJÖ

Stationary Replacement Strategies

Introduction

Most equipment used in industry and in daily life are deteriorating; to avoid costly failures and repairs it is sometimes wise to replace or maintain them before failures. However, maintenance incurs costs and a balance between different kinds of maintenance costs and different kinds of costs incurred by failures has to be achieved. See e.g. Barlow and Hunter [1] and Barlow and Proschan [2]. Some older reviews are given by Pierskalla and Voelker [3] and Valdez-Flores and Feldman [4]; a rather recent one by Wang [5]. In **Replacement Strategies** a general discussion is given. Here, we will restrict ourselves to the case of stationary replacement strategies and when such replacement rules are optimal. A replacement strategy is stationary, if the decision rule π to replace equipment is dependent only on data obtained since the last replacement; for the new unit used as a replacement the same decision rule is applied. It is assumed that the stochastic features in the different replacement cycles are stochastically independent. The conditions for stationary replacement rules will be scrutinized. Finally, arguments against stationary replacement rules will be given.

Some Stationary Decision Rules

Cost Structures

Two major types of cost structures and corresponding objective functions are used in the discussion of stationary replacement rules: long run average cost per time unit and expected discounted cost with a constant discount rate, see **Replacement Strategies** and e.g. Taylor [6], Zuckerman [7, 8], Bergman [9], Nummelin [10], and Aven [11]. Assuming independence between units of equipment, i.e. between replacement cycles, it is easily seen that the objective function for both of these cost structures often

may be expressed as

$$B_T = \frac{E \left[\int_0^T a(t)h(t) dt + c(0) \right]}{E \left[\int_0^T a(t) dt + p(0) \right]} \quad (1)$$

where T is a stopping time based on the information about the condition of the system (and possibly some **randomization** experiment), $\{a(t)\}$ is a nondecreasing stochastic process, $\{h(t)\}$ is a nonnegative stochastic process and $c(0)$ and $p(0)$ are nonnegative random variables; all variables are adapted to the information about the condition of the equipment under study (possibly formulated as a filtration $F_t, t \geq 0$); see for instance, Bergman [12] and Aven and Bergman [13]. The stopping time T is interpreted as the replacement time.

For the average long run cost per time unit the numerator in equation (1) represents the expected cost associated with one replacement cycle and the denominator, the expected length of that cycle; the objective function (equation 1) is derived using a renewal argument. For discounted expected cost, a similar renewal argument applies with the denominator $E[(1 - \exp\{-\alpha T\})]$, where α is a positive discount factor, and the numerator is the expected discounted cost during one replacement cycle. More generally, the objective function (equation 1) is obtained for a number of other regenerative stopping problems; see for instance, Ross [14]. In Aven and Bergman [13] the corresponding optimization problem, i.e. how to find an optimal stopping time T , is considered. In fact, the minimization problem (equation 1) is solved by minimizing a sequence of λ -functions

$$\begin{aligned} C_\lambda^T &= E \left[\int_0^T a(t)h(t) dt + c(0) \right] \\ &\quad - \lambda E \left[\int_0^T h(t) dt + p(0) \right] \\ &= E \left[\int_0^T (a(t) - \lambda)h(t) dt + c(0) - \lambda p(0) \right] \end{aligned} \quad (2)$$

These functions are minimized by $T_\lambda = \inf\{t \geq 0; a(t) \geq \lambda\}$ for $\lambda \in (-\infty, \infty)$ assuming

$$E \left[\int_0^{T_\lambda} h(t) dt + p(0) \right] < \infty \quad (3)$$

2 Stationary Replacement Strategies

A general algorithm for the solution is provided in an appendix of Aven and Bergman [13]. Essentially, this algorithm states that if $\lambda_{n+1} = B_{T_{\lambda_n}}$ (for some starting value λ_1), then $\lambda_n \rightarrow \lambda^*$ and T_{λ^*} is an optimal replacement time. A similar result, however, slightly less general, is given by Nummelin [10] using a different technique. In Aven and Bergman [13] also a discrete version of the replacement problem is discussed.

Age Replacement

For age replacement we refer **Multivariate Age and Multivariate Renewal Replacement**; here we only note that, for the **estimation** of an optimal replacement with objective function average long run cost per time unit, it is possible to utilize the total time on test (TTT) plotting concept, a solution is indicated in **Total Time on Test Plots**, see also Bergman [15] and Bergman and Klefsjö [16].

Replacement on Condition

Owing to a number of reasons the different equipment of the same kind are aging very differently. The most important reason is that of a variable environment (see **Maintenance Optimization in Random Environments**), but also variations in the production process of equipment are of importance. Anyhow, in many situations it would be advantageous to have some measurements on the equipment and its condition and base the replacement decision on these measurements. In fact, this idea is the basis for applied maintenance planning as suggested by Nolan and Heap [17]. They had found, studying replacements and **preventive maintenance** at United Airlines, that the generally applied **maintenance policies** were much too expensive, as they were not taking into account the information available on the aging of components. They suggested that condition monitoring should be utilized wherever economically feasible. Their general approach is called *reliability-centred maintenance* and in the aerospace area the corresponding maintenance planning scheme is called *maintenance steering group-3* (MSG-3) see for instance, Sandberg and Stromberg [18].

Replacement strategies based on condition monitoring of a state variable $X(t)$, assuming a known probability model, have been studied by e.g. Taylor [6] and Feldman [19] where X is a jump process

with failures occurring only at a jump and with a failure probability depending on the entered state. In Bergman [9], a general wear process is assumed together with a state-dependent failure rate $v(x)$. It is possible to use the objective function (equation 1) in these cases.

In Bergman [15], an empirically driven solution to the replacement problem based on condition monitoring of studied equipment is given. This solution is similar to the one suggested for the age replacement problem in **Total Time on Test Plots** and Bergman and Klefsjö [16]. At that time, however, sensors were not that refined and the solution was not widely utilized. Today, however, data retrieving mechanisms have become much better and replacement solutions based on condition monitoring are feasible. This is also seen from an increasing number of industry supported theses presented in this area; see e.g. Backlund [20], Söderholm [21] and Svensson [22].

Replacement of Repairable Systems

In some studies, (see e.g. Stadge and Zuckerman [23], Boshuizen [24], and Sandve and Aven [25]) the systems under study are assumed to be possible to repair minimally (see **General Minimal Repair Models**). Also in these cases it is possible to utilize the general framework introduced above.

Concluding Remarks

The idea of stationary replacement rules might be good as a basis for understanding the general replacement problem. However, it is assumed that the probability models are known from the beginning. In reality, that is almost never the case. Rather, it is important that we allow for learning in the strategies utilized. That leads to adaptive replacement strategies where successively obtained information about the failure model is obtained. The Bayesian approach is then a natural framework, see e.g. Mazzuchi and Soyer [26], Sheu *et al.* [27], and **Maintenance Optimization**. Stationary replacement strategies will not apply in this case.

Acknowledgment

This research has been partially supported by the Swedish foundation for strategic research *via* Gothenburg Mathematical Modeling Center (GMMC) and SKF AB.

References

- [1] Barlow, R.E. & Hunter, L.C. (1960). Optimum preventive maintenance policy, *Operations Research* **8**, 90–100.
- [2] Barlow, R.E. & Proschan, F. (1965). *Mathematical Theory of Reliability*, John Wiley & Sons, New York.
- [3] Pierskalla, W.P. & Voelker, J.A. (1976). A survey of maintenance models: the control and surveillance of deteriorating systems, *Naval Research Logistics Quarterly* **23**, 353–388.
- [4] Valdez-Flores, C. & Feldman, R.M. (1989). A survey of preventive maintenance models for stochastically deteriorating single-unit systems, *Naval Research Logistics Quarterly* **36**, 419–446.
- [5] Wang, H. (2002). A survey of maintenance policies of deteriorating systems, *European Journal of Operational Research* **139**, 469–489.
- [6] Taylor, H.M. (1975). Optimal replacement under additive damage and other failure models, *Naval Research Logistics Quarterly* **22**, 1–18.
- [7] Zuckerman, D. (1978). Optimal stopping in a semi-Markov shock model, *Journal of Applied Probability* **15**, 629–634.
- [8] Zuckerman, D. (1979). Optimal replacement rules-discounted cost criterion, *RAIRO Recherche Operationnelle/Operations Research* **13**, 67–74.
- [9] Bergman, B. (1978). Optimal replacement under a general failure model, *Advances in Applied Probability* **10**, 431–451.
- [10] Nummelin, E. (1980). A general failure model: optimal replacement with state dependent replacement and failure costs, *Mathematics of Operations Research* **5**(3), 381.
- [11] Aven, T. (1983). Optimal replacement under a minimal repair strategy – a general failure model, *Advances in Applied Probability* **15**, 198–211.
- [12] Bergman, B. (1980). On the optimality of stationary replacement strategies, *Journal of Applied Probability* **17**, 178–186.
- [13] Aven, T. & Bergman, B. (1986). Optimal replacement times – a general setup, *Journal of Applied Probability* **23**, 432–442.
- [14] Ross, S. (1971). Infinitesimal look-ahead stopping rules, *The Annals of Mathematical Statistics* **42**, 297–303.
- [15] Bergman, B. (1977). Some graphical methods for maintenance planning, in *Proceedings 1977 Annual Reliability and Maintainability Symposium*, Philadelphia.
- [16] Bergman, B. & Klefsjö, B. (1979). On age replacement and the total time on test concept, *Scandinavian Journal of Statistics* **6**, 161–168.
- [17] Nolan, F.S. & Heap, H.F. (1978). *Reliability-Centered Maintenance*, US Department of Commerce (NTIS), Springfield.
- [18] Sandberg, A. & Stromberg, U. (1999). G ripen: with focus on availability performance and life support cost over the product life cycle, *Journal of Quality in Maintenance Engineering* **5**(4), 325.
- [19] Feldman, R.M. (1976). Optimal replacement with semi-Markov shock models, *Journal of Applied Probability* **13**, 108–117.
- [20] Backlund, F. (2003). *Managing the introduction of reliability centred maintenance, RCM – RCM as a method of working within hydropower organizations*, Ph.D. thesis, Luleå University of Technology.
- [21] Söderholm, P. (2006). *Maintenance and continuous improvement of complex systems – linking stakeholder requirements to the use of built-in test systems*, Ph.D. thesis, Luleå University of Technology, Luleå.
- [22] Svensson, J. (2007). *Survival estimation of opportunistic maintenance*, Ph.D. thesis, Mathematical Statistics, Chalmers University of Technology, Gothenburg.
- [23] Stadjie, W. & Zuckerman, D. (1990). Optimal strategies for some repair replacement models, *Advances in Applied Probability* **22**, 641–656.
- [24] Boshuizen, F.A. (1997). Optimal replacement strategies for repairable systems, *Mathematical Methods of Operations Research* **45**, 133–144.
- [25] Sandve, K. & Aven, T. (1999). Cost optimal replacement of a monotone, repairable system, *European Journal of Operational Research* **116**(2), 235–249.
- [26] Mazzuchi, T.A. & Soyer, R.A. (1996). Bayesian perspective on some replacement strategies, *Reliability Engineering and System Safety* **51**, 295–303.
- [27] Sheu, S.-H., Yeh, R.H., Lin, Y.-B. & Juang, M.-G. (2001). A Bayesian approach to an adaptive preventive maintenance model, *Reliability Engineering and System Safety* **71**, 33–44.

Further Reading

Berg, M. (1976). A proof of optimality for age replacement policies, *Journal of Applied Probability* **13**, 751–759.

Related Articles

General Minimal Repair Models; Maintenance Optimization; Maintenance Optimization in Random Environments; Multivariate Age and Multivariate Renewal Replacement; Renewal Theory; Repairable Systems Reliability Replacement Strategies; Total Time on Test Plots.

BO BERGMAN

Maintenance Optimization in Random Environments

Introduction

In this article, we focus on a number of optimal management problems involving reliability and risk models operating in a random environment. The environment is described by a stochastic process, which represents important factors that affect the stochastic and deterministic parameters of the reliability or risk model. The term *environment* is used in the generic sense so that it represents any set of conditions that affect the stochastic structure of the model investigated. The concept of an “environmental” process, in one form or another, has been used in earlier literature for various purposes. Neveu [1] provides an early reference to paired stochastic processes where the first component is a Markov process while the second one has conditionally independent increments given the first. Ezhov and Skorohod [2] refer to this as a Markov process with a homogeneous second component. In a more modern setting, Çınlar [3, 4] introduced Markov additive processes and provided a detailed description on the structure of the additive component. The environment is modeled as a Markov process in all these cases and the additive process represents the stochastic evolution of a quantity of interest. Since the model parameters depend on the state of the environment at any time, the environmental process is also a process that modulates the model. As the state of the environment changes randomly, so do the model parameters and the environmental process indirectly generates the variation in the parameters.

The use of a modulating stochastic process as a source of variation in the model parameters and of dependence among the model components has proved to be quite useful in operations research and management science applications. In modeling hardware reliability for complex devices consisting of many interdependent components, the concept was introduced by Çınlar and Özekici [5] where the failure rate and hazard functions of the components all depend on the prevailing environmental conditions. For example, a complex device like a jet engine

consists of a large number of components where the failure structure of each component depends very much on the set of environmental conditions that it is subjected to during flight. The levels of vibration, atmospheric pressure, temperature, and so on, obviously change during take-off, cruising, and landing. Component lifetimes and reliabilities depend on these random environmental variations. Moreover, the components have dependent lifetimes since they operate in the same environment. The concept of random hazard functions is also used in Gaver [6] and Arjas [7]. The intrinsic aging model of Çınlar and Özekici [5] is studied further in Çınlar *et al.* [8] to determine the conditions that lead to associated component lifetimes, as well as multivariate increasing failure rate (IFR) and new better than used (NBU) life distribution characterizations. It was also extended in Shaked and Shanthikumar [9] by discussions on several different models with multicomponent replacement policies. Lindley and Singpurwalla [10] discuss the effects of the random environment on the reliability of a system consisting of components which share the same environment. Although the initial state of the environment is random, they assume that it remains constant in time and components have exponential life distributions in each possible environment. This model is also studied by Lefèvre and Malice [11] to determine partial orderings on the number of functioning components and the reliability of k -out-of- n systems, for different partial orderings of the probability distribution on the environmental state. The association of the lifetimes of components subjected to a randomly varying environment is discussed in Lefèvre and Milhaud [12]. Singpurwalla and Youngren [13] also discuss multivariate distributions that arise in models where a dynamic environment affects the failure rates of the components.

Although we discuss the use of environmental processes in managing hardware as well as **software reliability** and risk models, there is now considerable amount of literature on modulation in a variety of applications. An example in queuing is provided by Prabhu and Zhu [14] where customer arrival and service rates are modulated by a Markov process. Song and Zipkin [15] consider an inventory model with a demand process that fluctuates with respect to stochastically changing economic conditions. A general discussion on the idea can be found in Özekici [16]. The interested reader is referred to

Asmussen [17] and Rolski *et al.* [18] for further applications in queuing, insurance, and finance.

The stochastic structure of a rather general environmental process involving a **Markov renewal** process is described in the section titled “The Environmental Process”. Issues regarding the maintenance of complex systems under random environments are discussed in the section titled “Maintenance of Complex Systems”, and the last section presents some results on the structure of optimal **maintenance policies**.

The Environmental Process

We suppose that the environmental process is quite general as given in Çekyay and Özekici [19] (*see Markov Renewal Processes in Reliability Modeling*). Let T_n denote the time of the n th environment change and X_n denote the n th environmental state with $T_0 \equiv 0$. The main assumption is that the process $(X, T) = \{(X_n, T_n); n \geq 0\}$ is a Markov renewal process on the state space E with some semi-Markov kernel Q . Moreover, the environmental process $Y = \{Y_t; t \geq 0\}$ is the minimal semi-Markov process associated with (X, T) so that Y_t is the state of the environment at time t . More precisely, $Y_t = X_n$ whenever $T_n \leq t < T_{n+1}$. For any $i, j \in E$ and $t \geq 0$,

$$Q(i, j, t) = P\{X_{n+1} = j, T_{n+1} - T_n \leq t | X_n = i\} \quad (1)$$

and it is well-known that X is a Markov chain on E with transition matrix $P(i, j) = P\{X_{n+1} = j | X_n = i\} = Q(i, j, +\infty)$. We further assume that the Markov renewal process has infinite lifetime so that $\sup_n T_n = +\infty$. The probabilistic structure of T is given by the conditional distributions

$$\begin{aligned} G(i, j, t) &= P\{T_{n+1} - T_n \leq t | X_n = i, X_{n+1} = j\} \\ &= \frac{Q(i, j, t)}{P(i, j)} \end{aligned} \quad (2)$$

for $P(i, j) > 0$. Moreover, it is easy to see that

$$F_i(t) = P_i\{T_1 \leq t\} = \sum_{j \in E} Q(i, j, t) \quad (3)$$

denotes the distribution of the duration of state i .

Maintenance of Complex Systems

A complex system consists of many components that have stochastically dependent lifetimes. This dependence is induced by the common environment that they all operate in. Moreover, there is usually some economic dependence among the components due to the economies of scale obtained by maintaining many components at the same time. The most common type of maintenance involves the replacement of some of the functioning and nonfunctional components. The problem is to determine when replacement should be done and which components to replace. This should of course be done by observing the state of the reliability system. The intrinsic aging model introduced by Çınlar and Özekici [5] is a tractable one since it provides a realistic and simple stochastic structure of complex systems with many components. In Çekyay and Özekici [19], the authors consider various issues regarding reliability of such systems under the assumption that there is either no repair or maximal repair when all components are replaced at the beginning of each environmental state. In particular, the system reliability for the two cases are characterized through Markov renewal equations (15) and (37) in Çekyay and Özekici [19]. We will now extend their analysis by considering what happens when one incorporates a decision making process in their setting. We will follow their notation and terminology and the reader should first read that chapter. In particular, we set $R_+ = [0, \infty)$ and $F = R_+^m$ where m is the number of components. Moreover, we assume throughout that the components are serially connected so that the system fails as soon as any one of the components fails.

The intrinsic aging model is described by equations (25)–(33) in Çekyay and Özekici [19]. Recall that for any component k , $H_k(i, \cdot)$ is the cumulative hazard function and $r_k(i, \cdot)$ is the intrinsic aging rate function in environment i . The intrinsic age process is defined through the differential equation $dA_t(k)/dt = r_k(Y_t, A_t(k))$. The age of the system is $A_t = (A_t(1), A_t(2), \dots, A_t(m))$ at time t . It is clear that the process A takes values in F . Whenever the state of the environment changes at time T_n there is a decision made on maintenance based on the observation of the new environment X_n and intrinsic age $B_n = A_{T_n}$ of the system. The primary process is now given by the new Markov renewal process $((X, B), T)$.

We propose to use policies given by sets $M \subseteq E \times F$ so that **preventive maintenance** is done at the first-passage time

$$T_M = \inf\{T_n : (X_n, B_n) \in M\} \quad (4)$$

when the system is maintained if the observed environmental state i and the intrinsic ages $a = (a_1, a_2, \dots, a_m)$ of the components satisfy $(i, a) \in M$. The stopping time equation (4) identifies the time of maintenance. The type of maintenance to be done is another issue. The most common type involves replacements where one can replace either all components, or only a given a set of components, or only those that have failed. One can also use a repair type of policy where instead of replacing a component by a brand-new one, one repairs it to a lower intrinsic age level. A typical policy in this regard is the **minimal repair** one when a failed component is given a jump start and brought back to the working condition or age just before failure.

For any such policy, Çinlar and Özekici [5] make a Markov renewal argument to obtain the characterization

$$\begin{aligned} P_{ia}\{L > T_M\} &= \sum_{j \in E} \int_F \bar{R}(i, a; j, db) \\ &\quad \times \sum_{k \in E} \int_{R_+} Q(j, k, ds) e^{-\sigma(h(k, b, s) - b)} \\ &\quad \times 1_M(k, h(k, b, s)) \end{aligned} \quad (5)$$

for $(i, a) \notin M$ where $1_M(y)$ is the indicator function that is equal to 1 if and only if $y \in M$, and $\bar{R} = \sum_n \bar{P}^n$ is the potential kernel corresponding to the defective transition kernel

$$\begin{aligned} \bar{P}(i, a; j, db) &= \int_{R_+} Q(i, j, ds) H(i, a, s, db) \\ &\quad \times e^{-\sigma(b-a)} 1_{M^c}(j, b) \end{aligned} \quad (6)$$

where M^c denotes the complement of M . Note that equation (5) gives the probability that there is no system failure before the time of maintenance. The probability $P_{ia}\{L > t\}$ that the system does not fail until time t is already determined by equations (15) and (37) in Çekyay and Özekici [19] for the two extreme cases of no repair ($M = \phi$ and no preventive

replacement is done) and maximal repair ($M = E \times R_+^m$ and all components are replaced), respectively.

If $L(k)$ is the lifetime of the k th component, and L is the lifetime of the system, then $L = \min\{L(k); k = 1, 2, \dots, m\}$ and

$$S = \min\{T_M, L\} \quad (7)$$

is the first time either maintenance is done or the system fails, whichever occurs sooner. Then, using Markov renewal theory one obtains the joint distribution

$$\begin{aligned} P_{ia}\{X_S = j, A_S \in db, S \leq t, S = L(k)\} \\ = \sum_{l \in E} \int_{F \times [0, t]} \tilde{R}(i, a; l, dc; ds) \int_{[0, t-s]} du \bar{F}_l(u) \\ \times I(l, j) H(l, c, u; db) e^{-\sigma(b-c)} r_k(l, b) \end{aligned} \quad (8)$$

for $(i, a) \notin M$ where $\tilde{R} = \sum_n \tilde{Q}^n$ is the Markov renewal kernel corresponding to the semi-Markov kernel

$$\begin{aligned} \tilde{Q}(i, a; j, db; ds) &= Q(i, j, ds) H(i, a, s; db) \\ &\quad \times e^{-\sigma(b-a)} 1_{M^c}(j, b) \end{aligned} \quad (9)$$

and r_k is the intrinsic aging rate function given by equation (25) in Çekyay and Özekici [19]. Using a similar analysis, one also obtains

$$\begin{aligned} P_{ia}\{X_S = j, A_S \in db, S \leq t, S = T_M\} \\ = \sum_{l \in E} \int_{F \times [0, t]} \tilde{R}(i, a; l, dc; ds) \int_{[0, t-s]} Q(l, j, du) \\ \times H(l, c, u; db) e^{-\sigma(b-c)} 1_M(j, b) \end{aligned} \quad (10)$$

for $(i, a) \notin M$. In these expressions, equation (8) gives the probability that the system fails before time t due to the failure of component k in environmental state j at age db , while equation (10) gives the probability that the system is maintained before it fails until time t in environmental state j at age db . Both solutions are of the form $\tilde{R} * g$ for an appropriate function g . The sum of equation (8) for all k and equation (10) provides another Markov renewal characterization for the probability that preventive or failure maintenance is done before time t in environmental state j at age db . The probabilities (5), (8),

and (10) provide important performance measures to the decision maker to choose maintenance policies.

We will not attempt to formulate an optimization problem in our setting to determine optimal maintenance policies. It is indeed an interesting problem which has not been analyzed before for models with intrinsic aging in random environments depicted by a Markov renewal process. This will undoubtedly require various cost figures on system and component failures as well as costs of maintenance. We should point out, however, that the structure of optimal policies is by no means trivial and intuitive. The formulation of the optimization problem, the characterization of optimal policies, and the solution procedure is quite complicated. It is well known that the structure of optimal policies may be quite complex in multicomponent systems even when there are no environmental fluctuations. Preventive replacement is perhaps the most widely used maintenance policy to prevent the device from failure during operation, thus incurring excessive failure costs. The fixed environment case with several cost structures and objectives is discussed extensively in the literature by many authors and, in most cases, an age-replacement policy is optimal if the life distribution is IFR. Özekici [20] provides an example along this direction and a discussion on the optimality conditions for control-limit policies can be found in So [21].

Structure of Optimal Policies

In the optimal replacement problem, for example, perhaps the most widely used policy is the age-replacement policy where a component is replaced when its age is observed to be over a critical level. The problem then is to determine these ages optimally in order to minimize the average cost or total cost. A simple age replacement policy is one in which, in any environment i , there is critical age vector b_i such that component k is replaced if and only if the intrinsic age of the component exceeds $b_i(k)$. This policy is specified by m critical thresholds for each environment. However, such a policy is not necessarily optimal. It is clear that the component lifetimes are stochastically dependent due to their common environment. This implies that a component may be more prone to failure if the age of another component is higher and the critical threshold may depend on the age of other components. Another

complication is due to economic dependence among the components. This is the case when there is a fixed cost of replacement such that one is tempted to replace a functioning component when a failed component must be replaced since the fixed cost is already paid. In such cases, so-called opportunistic replacement policies with rather complex structures may be optimal. Özekici [22, 23] demonstrates this through several examples. For simplicity, suppose that there are two components only so that the optimal policy can be identified by the ages at which only component 1 is replaced (1), or only component 2 is replaced (2), or both components are replaced (1,2), or no component is replaced (0). Figure 1 depicts a typical optimal policy.

It is clear that the structure of the optimal policy, in general, is nontrivial and it cannot be identified by a few critical numbers. This is actually due to two reasons: economic and stochastic dependence among the components. Some rather counterintuitive observations can be made on Figure 1. For example, at point B, no component should be replaced but at point A, where there is substantial aging on component 2, both 1 and 2 should be replaced even though component 1 is at the same age as point B. This is called *opportunistic replacement* since 1 is replaced by making use of the opportunity in replacing component 2. Since an “old” component 2 is being replaced optimally, it may be best to replace component 1 at the same time. Opportunistic replacement is due to the fact that $b_1 \leq c_1$ and $b_2 \leq c_2$. This could be further clarified by considering the case where the components are stochastically independent while 1 has an exponentially distributed lifetime and

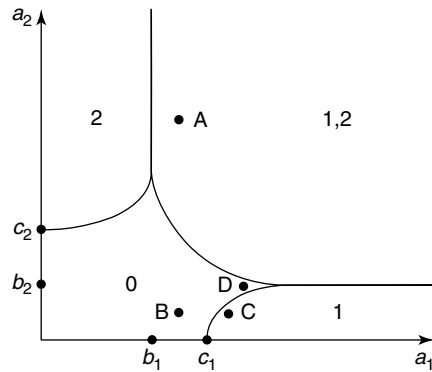


Figure 1 Optimal policy for 2 components

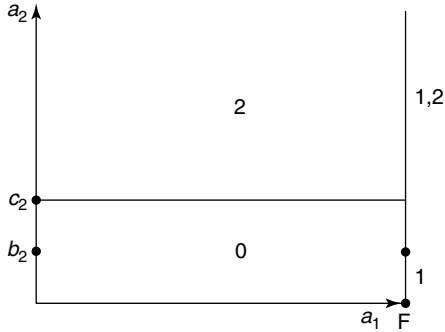


Figure 2 Optimal policy when one component has exponential lifetime

2 has an IFR life distribution. If the components are economically dependent, Radner and Jorgenson [24] showed that the optimal policy has the simple form given in Figure 2 for a fixed environment model. Note that component 1 is replaced only at failure (F) since it has an exponential lifetime. Component 2 is optimally replaced at the critical age c_2 , but if component 1 has failed and must be replaced, then 2 can be replaced opportunistically as early as age $b_2 \leq c_2$.

Another interesting observation follows from the comparison of points C and D on Figure 1. Note that component 1 is replaced at C but not at the “older” system age D. The system is not interfered with at the “worse” state D while a replacement decision is given at C. This is due to “opportunistic nonreplacement” since it may be best not to do any replacement at D and wait a little longer to replace both 1 and 2 at the same time.

The complexity in the structure of the optimal replacement policy creates computational as well as practical difficulties since it is not identified by a few critical numbers. Even if the optimal policy is determined by any one of the procedures, its implementation requires extreme precision. For these reasons, one may have to approximate the optimal policy by one that has a much simpler structure. An example is given by Van der Duyn Schouten, and Vanneste [25] analyzed the (n, N) policy given in Figure 3.

Note that the (n, N) policy provides substantial simplification when compared with the optimal policy in Figure 1. The fact that $n_1 \leq N_1$ and $n_2 \leq N_2$ show that, in essence, these policies reflect the existence of opportunistic replacement. However, this does not allow for “opportunistic nonreplacement” since the

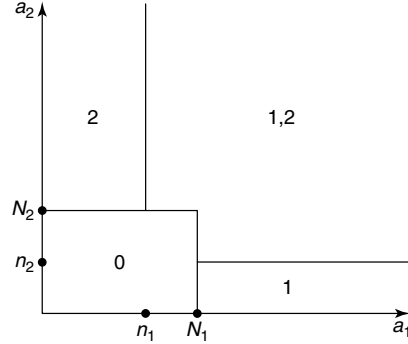


Figure 3 (n, N) policy

number of components replaced in an “older” system is at least as many as that of a “younger” system.

Acknowledgment

This research is supported by the Scientific and Technological Research Council of Turkey through grant 106M044.

References

- [1] Neveu, J. (1961). Une généralisation des processus à accroissements positifs indépendants, *Abhandlungen aus den Mathematischen Seminar der Universität Hamburg* **25**, 36–61.
- [2] Ezhov, I.I. & Skorohod, A.V. (1969). Markov processes with homogeneous second component: I, *Theory of Probability and its Application* **14**, 3–14.
- [3] Çinlar, E. (1972). Markov additive processes: I, *Zeitschrift für Wahrscheinlichkeitstheorie und verwandte Gebiete* **24**, 85–93.
- [4] Çinlar, E. (1972). Markov additive processes: II, *Zeitschrift für Wahrscheinlichkeitstheorie und verwandte Gebiete* **24**, 95–121.
- [5] Çinlar, E. & Özekici, S. (1987). Reliability of complex devices in random environments, *Probability in the Engineering and Informational Sciences* **1**, 97–115.
- [6] Gaver, D.P. (1963). Random hazard in reliability problems, *Technometrics* **5**, 211–226.
- [7] Arjas, E. (1981). The failure and hazard process in multivariate reliability systems, *Mathematics of Operations Research* **6**, 551–562.
- [8] Çinlar, E., Shaked, M. & Shanthikumar, J.G. (1989). On lifetimes influenced by a common environment, *Stochastic Processes and their Applications* **33**, 347–359.
- [9] Shaked, M. & Shanthikumar, J.G. (1989). Some replacement policies in a random environment, *Probability in the Engineering and Informational Sciences* **3**, 117–134.
- [10] Lindley, D.V. & Singpurwalla, N.D. (1986). Multivariate distributions for the lifetimes of components of

- a system sharing a common environment, *Journal of Applied Probability* **23**, 418–431.
- [11] Lefèvre, C. & Malice, M.-P. (1989). On a system of components with joint lifetimes distributed as a mixture of exponential laws, *Journal of Applied Probability* **26**, 202–208.
 - [12] Lefèvre, C. & Milhaud, X. (1990). On the association of the lifelenghts of components subjected to a stochastic environment, *Advances in Applied Probability* **22**, 961–964.
 - [13] Singpurwalla, N.D. & Youngren, M.A. (1993). Multivariate distributions induced by dynamic environments, *Scandinavian Journal of Statistics* **20**, 251–261.
 - [14] Prabhu, N.U. & Zhu, Y. (1989). Markov-modulated queueing systems, *Queueing Systems* **5**, 215–246.
 - [15] Song, J.S. & Zipkin, P. (1993). Inventory control in fluctuating demand environment, *Operations Research* **41**, 351–370.
 - [16] Özekici, S. (1996). Complex systems in random environments, in *Reliability and Maintenance of Complex Systems*, NATO ASI Series, S. Ozekici, ed, Springer-Verlag, Berlin, 137–157 Vol. F154.
 - [17] Asmussen, S. (2000). *Ruin Probabilities*, World Scientific, Singapore.
 - [18] Rolski, T., Schmidli, H., Schmidt, V. & Teugels, J. (1999). *Stochastic Processes for Insurance and Finance*, John Wiley & Sons, Chichester.
 - [19] Çekyay, B. & Özekici, S. (2007). Markov renewal processes in reliability modelling, in *Encyclopedia of Statistics in Quality and Reliability*, F. Ruggeri, F. Faltin & R. Kenett, eds, John Wiley & Sons, Chichester.
 - [20] Özekici, S. (1985). Optimal replacement of one-unit systems under periodic inspection, *SIAM Journal on Control and Optimization* **23**, 122–128.
 - [21] So, K.C. (1992). Optimality of control limit policies in replacement models, *Naval Research Logistics* **39**, 685–697.
 - [22] Özekici, S. (1988). Optimal periodic replacement of multicomponent reliability systems, *Operations Research* **36**, 542–552.
 - [23] Özekici, S. (1995). Optimal maintenance policies in random environments, *European Journal of Operational Research* **82**, 283–294.
 - [24] Radner, R. & Jorgenson, D.W. (1963). Opportunistic replacement of a single part in the presence of several monitored parts, *Management Science* **10**, 70–84.
 - [25] Van der Duyn Schouten, F.A. & Vanneste, S.G. (1990). Analysis and computation of (n, n) strategies for maintenance of a two-component system, *European Journal of Operational Research* **48**, 260–274.

Related Articles

Block Replacement; Group Maintenance Policies; Maintenance and Markov Decision Models; Maintenance Optimization; Markov Processes; Markov Renewal Processes in Reliability Modeling; Multicomponent Maintenance; Multivariate Age and Multivariate Renewal Replacement; Multivariate Imperfect Repair Models; Repairable Systems Reliability; Replacement Strategies; Total Productive Maintenance.

SÜLEYMAN ÖZEKICI

Group Maintenance Policies

Maintenance Policies

Suppose that n machines are operating in parallel. A loss of c_d per unit time is incurred for each failed machine. The cost of replacing a broken machine with a new machine is denoted by c_b , while the cost of servicing a working machine to be “as good as new” is given by c_r . A fixed cost of c_0 is incurred each time replacement or servicing of machines occurs. In this paper group replacement and servicing strategies are considered. The failure times of the machines are assumed to be independent identically distributed random variables with increasing hazard rate and density function $f(\bullet)$. The random variable representing the life length of a machine will be denoted by X .

A number of group replacement policies have appeared in the literature. An m -failure policy calls for group replacement at the m th failure (see Assaf and Shantikumar [1], Nakagawa [2], Park [3], Wilson and Benmerzouga [4]). T -age group replacement policies call for group replacement whenever the system reaches age T (see, e.g., Okumoto and Elsayed [5]). Ritchken and Wilson [6] introduce a class of policies, called (m, T) **group maintenance policies**, that call for group replacement at the m th failure or time T , whichever occurs first; Benmerzouga and Harous [7], a class of policies called $(m\text{-failure}, P\text{-repairmen})$ **group maintenance policies**, that call for group replacement at the m th failure with P -repairmen in the crew. Wilson and Benmerzouga [8] consider the case where the underlying failure distribution is exponential with an unknown parameter that must be estimated from the operation of the machines. A more general form of this policy for the case of three machines operating in parallel is considered by Wilson and Popova [9]. The case of phase distributed failure times has been analyzed by Wilson and Popova [9]. Mazzuchi and Soyer consider the case of a single machine with a Weibull failure time [11].

The m -failure, T -age, and (m, T) group replacement policies have one managerially unpalatable drawback: they ignore what happens during the cycle until the m th failure or time T has occurred. For

instance, consider a problem where 100 machines are operating in parallel and the optimal m -failure policy is to replace whenever the 51st machine breaks down. Consider the following two cases: the first 51 machines fail 1 year after being put into operation; 50 machines fail on the second day of operation and the 51st fails 1 year after the start of the cycle. The 51-failure policy treats the above two cases in exactly the same manner – group replacement is performed after 1 year of operation. However, most decision makers would want to treat the above two situations in a completely different manner. Of course, the analysis that would have provided the 51-failure policy as the optimal m -failure policy is based on average cost per unit time over an infinite horizon. Even though the cost of carrying 50 nonfunctioning machines for 1 year might be very large, the probability of such an event occurring would be so low that this case would be essentially disregarded in the average per unit time analysis. However, events of low probability can occur. Although it might be reasonable to discount them *a priori*, it would be foolish to ignore them if they actually occur. Assuming continuous inspection, this paper provides an easily implemented replacement policy that takes account of the actual state of the system at each stage.

Define a decision rule δ to be a function from $\{(i, x): 0 \leq i \leq n, 0 \leq x \leq \infty\}$, into the closed interval $[0, 1]$, where $\delta(i, x)$ denotes the probability with which group replacement will be made, if exactly i machines are broken and it has been x time units since the last group replacement. This class of strategies includes the m -failure, T -age and (m, T) policies as special cases. An m -failure policy is obtained by requiring $\delta(i, x)$ to equal 1 if $i \geq m$ and to equal 0, otherwise. A T -age replacement policy is obtained by setting $\delta(i, x)$ to equal 1 for $x \geq T$ and equal to 0, otherwise. An (m, T) policy is obtained by defining δ to satisfy

$$\delta(i, x) = \begin{cases} 1 & \text{if } i \geq m \\ 1 & \text{if } x \geq T \\ 0 & \text{otherwise} \end{cases} \quad (1)$$

Implementation of a policy δ as defined above only requires knowledge of the current state (i, x) . Thus no extra information, other than normally available when using an m -failure policy, is required. However, finding the optimal m -failure policy is relatively straightforward. Finding a policy δ that is optimal or ε -optimal (a discretized solution within ε of the true

2 Group Maintenance Policies

optimal) is not quite so straightforward. However, once such a policy has been found, implementation is not much harder than it would be for m -failure, T -age, or (m, T) policies.

The section titled “Optimal Policies When System is Observed at a Fixed Number of Times” provides a procedure for finding an optimal replacement policy when decisions can only be made at times $x_1 < x_2 < \dots < x_m$. As $x_m \rightarrow \infty$ and $\max_j |x_{j+1} - x_j| \rightarrow 0$, the cost associated with the policy computed using this assumption converges to the cost of the optimal policy for the case where decisions are allowed at anytime. The section titled “Continuous Observation” contains a procedure that will produce lower bounds on the optimal cost.

Optimal Policies When System is Observed at a Fixed Number of Times

In this section, the case where the state of the system is only observed at the fixed times $0 = x_0 < x_1 < x_2 < \dots < x_m$ is considered. It will be assumed that all machines will be replaced no later than time x_m and that none are replaced before time x_1 . The process is defined to be in state (i, x_j) if the current time is given by x_j and exactly i machines are broken.

In the section titled “Transition Probabilities and Costs”, the transition probabilities and costs of moving between states will be derived while the section titled “Policies” contains a description of the space of policies that will be considered. The section titled “An LP Formulation When the Expected Cycle Time is Provided” contains a linear programming (LP) formulation that will provide the best policy with a specified expected cycle time. (The LP approach adopted is similar to that described in Heyman and Sobel [12].) A search method over the expected cycle time then yields policies that are ε -optimal.

Transition Probabilities and Costs

At a given state, one of two actions must be taken: replace and/or repair all machines to as good as new ($a = 1$); wait until the next state ($a = 0$). Define $F_j(x)$ to be the conditional distribution function $P[X \leq x | X > x_j]$. For $j \geq 1$, define $P_{k,j-1}^{i,j}$ to be the probability that exactly i machines are broken at time x_j given that $k \leq i$ were broken at time x_{j-1} , i.e.,

$$P_{k,j-1}^{i,j} = \binom{n-k}{i-k} \{F(x_{j+1}|x_j)\}^{i-k} \times \{1 - F(x_{j+1}|x_j)\}^{n-i} \quad (2)$$

Let c_{ija} , $a \in \{0, 1\}$ denote the expected cost incurred while the process moves from state (i, x_j) to the next state if action a is taken. If the action is to perform group replacement then the cost consists of two machines: the cost of purchasing replacement machines for those that have failed and the cost of performing maintenance on the machines that are still working. Thus, $c_{ij1} = c_0 + ic_b + (n-i)c_r$. Now suppose that replacement is not performed at state (i, x_j) . The cost incurred between time x_j and x_{j+1} will consist of two parts: the downtime for the i machines that have failed by time x_j and the downtime for any machine that fails between x_j and x_{j+1} . The expected downtime over the period x_j to x_{j+1} for any one of the machines working at time x_j is given by $x_{j+1} - E[\min(x_{j+1}, X) | X > x_j] = \int_{x_j}^{x_{j+1}} F_j(x) dx$. Thus, the immediate expected cost of continuing at state (i, x_j) is given by

$$c_{ij0} = ic_d(x_{j+1} - x_j) + (n-i)c_d \int_{x_j}^{x_{j+1}} F_j(x) dx \quad (3)$$

where the first term is the downtime cost of the i machines that are broken at time x_j and the second term is the expected downtime cost for machines that fail between x_j and x_{j+1} . Thus, the costs c_{ija} can be written as follows:

$$c_{ija} = \begin{cases} c_0 + ic_b + (n-i)c_r & \text{for } a = 1 \\ ic_d(x_{j+1} - x_j) + (n-i)c_d \int_{x_j}^{x_{j+1}} F_j(x) dx & \text{for } a = 0 \end{cases} \quad (4)$$

Policies

A randomized policy δ is a function from $\{(i, x_j) : 0 \leq i \leq n, 1 \leq j \leq m\}$ to the interval $[0, 1]$. Here, $\delta(i, x_j) \in [0, 1]$ is to be interpreted as the probability with which the decision maker will perform group replacement. Set $\delta(i, x_m) \equiv 1$ for all i (i.e., no matter how many machines are working, replacement is required at time x_m).

Assume that the randomized policy δ is to be used. A cycle is defined to be the period between group replacements. Let $K(\delta)$, $C(\delta)$, and $T(\delta)$, respectively, denote the random variables for the cost per unit time, the cycle cost and the cycle length for a policy δ . Then, from the Renewal Reward theorem in terms of

expected values,

$$E[K(\delta)] = \frac{E[C(\delta)]}{E[T(\delta)]} \quad (5)$$

For purposes of mathematical convenience, the analysis that follows will be for this general class of randomized strategies. However, the solution procedure will always produce a deterministic policy. (This can be seen by applying the following argument from decision theory. Suppose $\delta^*(.,.)$ is an optimal policy with $0 < \delta^*(i_0, x_{j_0}) < 1$ for some i_0 and j_0 . Define two new policies, $\delta^1(.,.)$ and $\delta^2(.,.)$ to be identical to $\delta^*(.,.)$ except that at the state (i_0, x_{j_0}) , $\delta^1(i_0, x_{j_0})$ equals 0 while $\delta^2(i_0, x_{j_0})$ equals 1, i.e., $\delta^1(.,.)$ and $\delta^2(.,.)$ are nonrandomized at (i_0, x_{j_0}) . Now suppose that, at the beginning of each cycle, the decision maker chooses to implement either policy $\delta^1(.,.)$ with probability $1 - \delta^*(i_0, x_{j_0})$ or policy $\delta^2(.,.)$ with probability $\delta^*(i_0, x_{j_0})$. The expected cost per unit time of this procedure equals

$$\frac{(1 - \delta^*(i_0, x_{j_0}))E[C(\delta^1)] + \delta^*(i_0, x_{j_0})E[C(\delta^2)]}{(1 - \delta^*(i_0, x_{j_0}))E[T(\delta^1)] + \delta^*(i_0, x_{j_0})E[T(\delta^2)]} \quad (6)$$

which, by construction, equals $E[(\delta^*)]$. Note that

$$\frac{(1 - z)E[C(\delta^1)] + zE[C(\delta^2)]}{(1 - z)E[T(\delta^1)] + zE[T(\delta^2)]} \quad (7)$$

is a monotonic function of z for $z \in [0, 1]$ that equals $E[K(\delta^1)]$ at $z = 0$ and $E[K(\delta^2)]$ at $z = 1$.

Consequently,

$$E[K(\delta^*)] \geq \min\{E[K(\delta^1)], E[K(\delta^2)]\} \quad (8)$$

So, an optimal policy is given by either $\delta^1(.,.)$ or $\delta^2(.,.)$, neither of which is randomized at (i_0, x_{j_0}) . Now repeat this argument for any other state for which $0 < \delta^*(i, x_j) < 1$.

An LP Formulation When the Expected Cycle Time is Provided

Let $p_{ij0}(\delta)$ denote the probability that the state (i, x_j) will be visited during the cycle and the decision not to replace the system at that time is made. Let $p_{ij1}(\delta)$ denote the probability that the process will visit the state (i, x_j) and the decision to replace all machines is made.

The expected cost per cycle can be written as $c_{000} + \sum c_{ija} p_{ija}(\delta)$. For $j \geq 1$, let ℓ_{ija} be the transition time from state (i, x_j) to the next state when action a is taken, i.e.,

$$\ell_{ija} = \begin{cases} x_{j+1} - x_j & \text{for } a = 0 \\ 0 & \text{for } a = 1 \end{cases} \quad (9)$$

Then the expected cost per unit time from following the policy δ is given by

$$E[K(\delta)] = \frac{c_{000} + \sum_{i,j \geq 1,a} c_{ija} p_{ija}(\delta)}{x_1} + \sum_{i,j \geq 1,a} \ell_{ija} p_{ija}(\delta) \quad (10)$$

Theorem 1 shows that, given a policy with a specific value for the expected cycle length, the expected cost per unit time can be written in a linear form with linear constraints with appropriately chosen variables. In Theorem 2, it is shown that any such formulation leads to a policy with the specified expected cycle time.

Theorem 1 *Let $\delta(.,.)$ be any given policy that calls for replacement no sooner than x_1 . Denote the expected length of the cycle using this policy by B . Then the expected cost per unit time associated with this policy can be written in the form*

$$\frac{c_{000}}{B} + \sum_{i,j \geq 1,a} c_{ija} y_{ija} \quad (11)$$

where $i \in \{0, \dots, n\}$, $j \in \{1, 2, \dots, m\}$, $a \in \{0, 1\}$ and y_{ija} satisfy the following constraints:

$$y_{ija} \geq 0 \forall i, j, a \quad (12)$$

$$\frac{x_1}{B} + \sum_{i,j,a} \ell_{ija} y_{ija} = 1 \quad (13)$$

$$y_{ij0} + y_{ij1} = \sum_{k \leq i} P_{k,j-1}^{i-k,j} y_{i(j-1)0}$$

$$\text{for } j \geq 2 \text{ and } i \in \{1, \dots, n\} \quad (14)$$

$$y_{i10} + y_{i11} = B^{-1} P_{0,0}^{i,1} \text{ for } i = \{1, \dots, n\} \quad (15)$$

Proof. The expected length, $E[T(\delta)]$, of the renewal cycle when the policy δ is used is given by

$$B = E[T(\delta)] = x_1 + \sum_{i,j \geq 1,a} \ell_{ija} p_{ija}(\delta) \quad (16)$$

4 Group Maintenance Policies

For $j \geq 1$, define the quantities y_{ija} as follows:

$$y_{ija} \equiv \frac{p_{ija}}{E[T(\delta)]} \quad (17)$$

Then the expected cost per unit time can be written as

$$\begin{aligned} \frac{c_{000}}{E[T(\delta)]} + \sum_{i,j \geq 1,a} \frac{c_{ija} p_{ija}}{E[T(\delta)]} \\ = \frac{c_{000}}{B} + \sum_{i,j \geq 1,a} c_{ija} y_{ija} \end{aligned} \quad (18)$$

The quantities y_{ija} must satisfy a number of constraints. Since $p_{ija}(\delta) \geq 0$ and $x_1 + \sum_{i,j \geq 1,a} \ell_{ija} p_{ija}(\delta) \equiv E[T(\delta)] = B$, the quantities y_{ija} must satisfy equations (12) and (13).

The probability that the process visits state (i, x_j) during a cycle is given by $p_{ij0}(\delta) + p_{ij1}(\delta)$. The process can only be in this state if $k \leq i$ machines were broken at time x_{j-1} , action $a = 0$ was taken, and exactly $i - k$ machines broke down during the period from time x_{j-1} to x_j . Thus, for $j \geq 2$, the quantities $p_{ija}(\delta)$ must satisfy

$$p_{ij0}(\delta) + p_{ij1}(\delta) = \sum_{k \leq i} P_{k,j-1}^{i-k,j} p_{k(j-1)0}(\delta) \quad (19)$$

On dividing the above expression by $E[T(\delta)]$, the quantities y_{ija} are seen to satisfy equation (14).

Between times 0 and x_1 , i machines fail with probability equal $P_{0,0}^{i,1}$. Equation (15) follows on noting that

$$p_{i10}(\delta) + p_{i11}(\delta) = P_{0,0}^{i,1} \quad (20)$$

and dividing across by B .

Theorem 2 For a given B with $B \geq x_1$, suppose that y_{ija} for $i \in \{0, 1, 2, \dots, n\}$, $j \in \{1, \dots, m\}$, and $a \in \{0, 1\}$ satisfy the following:

$$y_{ija} \geq 0 \forall i, j, a \quad (21)$$

$$\frac{x_1}{B} + \sum_{i,j \geq 1,a} \ell_{ija} y_{ija} = 1 \quad (22)$$

$$\begin{aligned} y_{ij0} + y_{ij1} &= \sum_{k \leq i} P_{k,j-1}^{i-k,j} y_{i(j-1)0} \\ \text{for } j \geq 2 \text{ and } i \in \{0, \dots, n\} \end{aligned} \quad (23)$$

$$y_{i10} + y_{i11} = B^{-1} P_{0,0}^{i,1} \text{ for } i \in \{0, \dots, n\} \quad (24)$$

Define a group replacement policy, $\delta(.,.)$ as follows:

$$\delta(i, x_j) = \frac{y_{ij1}}{y_{ij0} + y_{ij1}}, \forall i, j \quad (25)$$

i.e., if i machines are down at time x_j , replace the system with probability equal to $y_{ij1}/(y_{ij0} + y_{ij1})$.

Then,

$$E[T(\delta)] = B \quad (26)$$

and

$$E[K(\delta)] = \frac{c_{000}}{B} + \sum_{i,j \geq 1,a} c_{ija} y_{ija} \quad (27)$$

Proof. First for $j \geq 1$, it will be shown that for a given cycle, the fraction of time i machines are down and decision a is taken at time x_j , is given by

$$p_{ija} = B y_{ija} \quad (28)$$

The proof will be by induction on j . First consider the case $j = 1$. At time x_1 , the probability that exactly i machines are down is given by $P_{0,0}^{i,1}$. Replacement will occur with probability equal to $\delta(i, x_1)$. Thus, using equation (25),

$$p_{i1a}(\delta) = \begin{cases} P_{0,0}^{i,1} \delta(i, 1) = \frac{P_{0,0}^{i,1} y_{i11}}{(y_{i10} + y_{i11})} \\ \text{for } a = 1 \\ P_{0,0}^{i,1} \{1 - \delta(i, 1)\} = \frac{P_{0,0}^{i,1} y_{i10}}{(y_{i10} + y_{i11})} \\ \text{for } a = 0 \end{cases} \quad (29)$$

Now use equation (24) in equation (29) to see that

$$p_{i1a} = B y_{i1a} \quad (30)$$

as required.

Now assume that equation (28) is true for $j - 1$. The goal is to show that this implies that equation (28) is true for j .

For i machines to be down at time x_j , $k \leq i$ machines must have been down at time x_{j-1} , the decision to continue was made at time x_{j-1} and $i - k$ must have failed between times x_{j-1} and x_j . A decision to replace will be made with probability $\delta(i, j)$. Thus,

$$p_{ija}(\delta) = \begin{cases} \delta(i, j) \sum_{k \leq i} P_{k,j-1}^{i-k,j} p_{k(j-1)0} \\ \text{for } a = 1 \\ (1 - \delta(i, j)) \sum_{k \leq i} P_{k,j-1}^{i-k,j} p_{k(j-1)0} \\ \text{for } a = 0 \end{cases} \quad (31)$$

Now use the induction hypothesis to replace $p_{k(j-1)0}$ with $y_{k(j-1)0}/B$ in equation (31) to see that

$$p_{ija}(\delta) = \begin{cases} \delta(i, j) B \sum_{k \leq i} P_{k, j-1}^{i, -k, j} y_{k(j-1)0} & \text{for } a = 1 \\ (1 - \delta(i, j)) B \sum_{k \leq i} P_{k, j-1}^{i, -k, j} y_{k(j-1)0} & \text{for } a = 0 \end{cases} \quad (32)$$

Use equation (23) and the definition of $\delta(i, j)$ to see that $p_{ija} = B y_{ija}$ as required.

From equations (22) and (28) the definition of expected cycle time:

$$\begin{aligned} E[T(\delta)] &= x_1 + \sum_{i, j \geq 1, a} \ell_{ija} p_{ija}(\delta) \\ &= x_1 + B \sum_{i, j \geq 1, a} \ell_{ija} y_{ija} \\ &= B \end{aligned} \quad (33)$$

as required.

From the above, equations (10) and (28), the expected cost associated with $\delta(., .)$ is given by

$$\begin{aligned} E[K(\delta)] &= \frac{c_{000} + \sum_{i, j \geq 1, a} c_{ija}(\delta) p_{ija}(\delta)}{E[T(\delta)]} \\ &= \frac{c_{000} + \sum_{i, j \geq 1, a} c_{ija}(\delta) p_{ija}}{B} \\ &= \frac{c_{000}}{B} + \sum_{i, j \geq 1, a} c_{ija}(\delta) y_{ija} \end{aligned} \quad (34)$$

as required.

Example Suppose there are six machines with i.i.d. $\gamma(\alpha = 5, \beta = 5)$ distributed failure times with costs as shown in Table 1.

Figure 1 displays optimal policy cost as a function of cycle time. The determination of the optimal

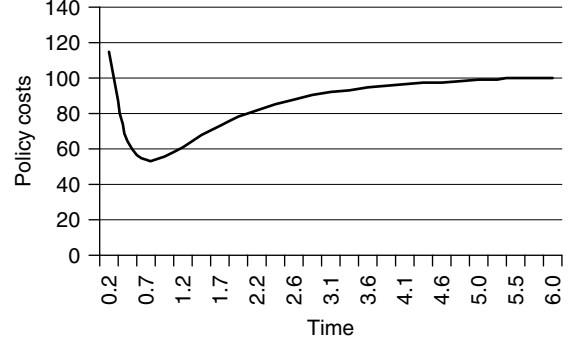


Figure 1 Best expected cost per unit time for a specified cycle time

policy required the linear program to be solved at each given expected cycle time followed a one-dimensional search to find the minimal cost. We speed up this one-dimensional search by using a simple golden-section search. (The results shown are for a grid with $x_j - x_{j-1} = 0.024$ for all j .) As can be seen from Figure 1, the optimal policy has an expected cycle length of 0.72. The actual policy is displayed graphically in Figure 2.

In Figure 2, the vertical axis contains the number of machines. The horizontal axis is the elapsed time since the beginning of a cycle and the shaded region is the region where group replacement should be made. From Figure 2, the optimal policy can be described as follows: at any time prior to 0.48, replace only if three or more machines are down; between times 0.48 and 0.9, replace if two or more machines are down; between times 0.9 and 1.2, replace if any machines are down; after 1.9, replace the system no matter how many machines have failed. This policy has an expected cost per unit time of 53.9. (The best (m, t) policy is $(3, 0.75)$: repair once three machines are down or 0.75 days have elapsed, with a cost of 56.04 (3.97% higher than optimal).) The best m policy is $m = 2$: repair once two machines are down with a cost of 58.71 (8.92% above optimal).

Table 1 System costs

Cost	Value
c_0	15
c_b	7
c_r	2
c_d	20

Continuous Observation

A more general problem is to allow the state space to be $\{(i, x): 0 \leq i \leq n, 0 \leq x < x_m\}$, where the assumption is that x_m is chosen so large that replacement must occur before x_m (the assumption of an increasing hazard rate implies that such an x_m exists).

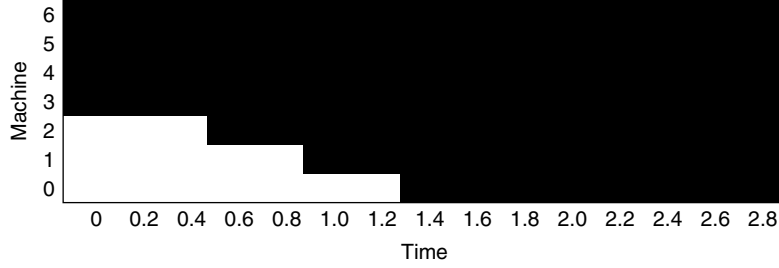


Figure 2 Optimal maintenance policy

As before, the analysis will be for randomized strategies. An optimal solution, of course, will always turn out to be a nonrandomized strategy. In this context, a randomized strategy s is a function from $\{(i, x): 0 \leq i \leq n, 0 \leq x < x_m\}$ to the closed interval $[0, 1]$, with $s(i, x)$ being the probability that group replacement will be performed when the system is in state (i, x) . Let K^* denote $\inf_s \{K(s)\}$, where, as before, $K(s)$ denotes the expected cost per unit time associated with the policy s .

As in the section titled “Optimal Policies When System is Observed at a Fixed Number of Times”, let $0 = x_0 < x_1 < x_2 < \dots < x_m$ define a given set of times. The objective of this section is to demonstrate that the times x_j can be chosen so that a policy defined on this restricted space is arbitrarily close to K^* . Any policy defined on this restricted space is also a policy that can be applied in the continuous observation case and consequently, its expected cost per unit time is an upper bound on K^* . Theorem 3, which follows, shows how to obtain a lower bound for K^* .

Theorem 3 Let $0 = x_0 < x_1 < x_2 < \dots < x_m$ be a given set of times at which the system is observed. (So the state space is $\{(i, x_j): 0 \leq i \leq n, 1 \leq j \leq m\}$.) Suppose that δ^* is the policy on this space that minimizes the expected cost per unit time where transition costs are given by

$$c'_{ija} \equiv \begin{cases} c_{ij1} - ic_d(x_j - x_{j-1}) + (n-i)c_d & \\ \int_{x_{j-1}}^{x_j} F_{j-1}(x) dx - (c_b - c_r)(n-i) & \\ \{F_{j-1}(x_j) - F_{j-1}(x_{j-1})\} \text{ for } a = 1 & \\ c_{ij0} \text{ for } a = 0 & \end{cases} \quad (35)$$

and the c_{ija} are provided in equation (4). Then the expected cost per unit time for δ^* using the costs c'_{ija} is a lower bound for K^* .

Proof. Let s be any given strategy defined on the space $\{(i, x): 0 \leq i \leq n, 0 \leq x < x_m\}$. For the given partition $\{x_j\}$, let δ^s denote the strategy on the discrete space $\{(i, x_j): 0 \leq i \leq n, 0 \leq j \leq m\}$ be defined as follows: group replacement using strategy δ^s occurs at time x_j if and only if a decision maker who had been using s would have replaced the system at time $x \in (x_{j-1}, x_j]$; set $\delta^s(i, x_m) \equiv 1$ for all i .

The expected downtime cost over a cycle using δ^s is larger than that using s since replacement is being deferred. Suppose that, using δ^s , group replacement is required at (i, x_j) . Then, for this case, the extra downtime cost can be no more than

$$ic_d(x_j - x_{j-1}) + (n-i)c_d \int_{x_{j-1}}^{x_j} F_j(x) dx \quad (36)$$

where the first part of equation (36) is the downtime accrued by the i failed machines over the period from x_{j-1} to x_j and the second part of equation (36) is an upper bound for the expected downtime for extra machines that might have failed during $(x_{j-1}, x_j]$.

The expected replacement cost over a cycle using δ^s is larger than that using s since replacement is being deferred and extra replacement costs might be incurred – the extra cost increment incurred for each extra failed machine being $c_b - c_r$. Suppose that, using δ^s , group replacement is required at (i, x_j) , the conditional probability that a machine will fail between x_{j-1} and x_j equals $F_{j-1}(x_j) - F_{j-1}(x_{j-1})$. Thus, the expected extra replacement costs incurred by waiting to replace at x_j do not exceed

$$(c_b - c_r)(n-i)\{F_{j-1}(x_j) - F_{j-1}(x_{j-1})\} \quad (37)$$

The extra expected downtime cost incurred by using δ^s is no more than the amounts given in equations (36) and (37). So subtracting these amounts from the replacement cost on replacement, more than compensates for these amounts. So calculating the expected cost per unit time for δ^s using replacement costs of

$$\begin{aligned} & c_{ij1} - i c_d(x_j - x_{j-1}) + (n - i)(c_b - c_r) \\ & \times \int_{x_{j-1}}^{x_j} F_{j-1}(x) dx - (c_b - c_r)(n - i) \\ & \times \{F_{j-1}(x_j) - F_{j-1}(x_{j-1})\} \end{aligned} \quad (38)$$

instead of c_{ij1} will result in a lower value than the expected cost over the cycle using s .

The expected length of each renewal cycle is longer when using δ^s than s . This is because δ^s defers group replacement to the first time x_j that comes after the strategy s would have replaced the machines. If the policy s replaces at $x \in (x_{j-1}, x_j]$, then the length of the cycle using δ^s is no more than

$$x_j - x_{j-1} \quad (39)$$

units longer than if s were used.

From the above, the expected cost per unit time for δ^s , where the transition costs are given by the modified values c'_{ija} , is lower than the expected cost per unit time for s where this latter cost is computed using the (true) transition costs given by c_{ija} . This gives the required result since it holds true for any s .

From equations (36), (37) and (39), the difference between the expected cost per unit time for the case of a fixed partition and the value obtained from procedure in Theorem 3 goes to zero as $\max_j |x_j - x_{j-1}| \rightarrow 0$. In other words, the optimal policy for the state space $\{(i, x_j) : 0 \leq i \leq n, 1 \leq j \leq m\}$ can have an expected cost per unit time as close as desired to K^* by choosing a fine enough partition $\{x_j\}$.

Example Suppose there are six machines with i.i.d. $\gamma(\alpha = 5, \beta = 5)$ distributed failure times with costs as shown in Table 1. Table 2 contains the values of the objective function for different equally spaced partitions of the interval 0 to 3. As can be seen, even moderately sized partitions produced reasonable results.

Table 2 Results for varying partition step sizes $x_j - x_{j-1}$

Costs		
Partition step size	Partition size	Upper bound–lower bound
0.12	25	0.971
0.06	50	0.861
0.04	75	0.825
0.03	100	0.706
0.024	125	0.524

Table 3 Comparison of costs for optimal policy to m and (m, t) policies

Parameter		Policy costs		Cost increase over optimal		
α	β	Optimal	(m, t)	m	(m, t) (%)	m (%)
7	7	49.44	50.56	53.93	2.27	9.08
7	5	36.87	37.82	39.28	2.58	6.54
7	3	22.72	23.15	23.61	1.89	3.92
5	7	73.61	75.42	79.18	2.46	7.57
5	5	53.9	56.04	58.71	3.97	8.92
5	3	33.84	35.55	37.18	5.05	9.87
3	5	118.79	118.92	127.16	0.11	7.05
3	7	91.87	92.27	100.17	0.44	9.03
3	3	60.56	61.34	66.47	1.29	9.76

Example Again, suppose there are six machines with costs as shown in Table 1. A comparison of the optimal policy with the results of m -failure and (m, T) -policies for various gamma distributions is shown in Table 3.

Conclusion

This paper introduces a group replacement policy that takes account of the actual state of the system at each stage. This general approach cannot have a higher expected cost per unit time than the m -failure, T -age, or (m, T) policies and is a generalization of those policies. The general approach is no harder to implement than any of these policies and has the great advantage of being intuitively much more acceptable. The drawback of the procedure is that more computational effort is required to find good policies. The adaptive policy allows the production manager to efficiently use the information provided by failing machines. This will result in a reduction in the variability of costs among cycles.

References

- [1] Assaf, D. & Shantikumar, G. (1987). Optimal group maintenance policies with continuous and periodic inspection, *Management Science* **33**, 1440–1452.
- [2] Nakagawa, T. (1983). Optimal number of failures before replacement time, *IEEE Transactions on Reliability* **R-32**, 114–116.
- [3] Park, K.S. (1979). Optimal number of minimal repairs before replacement, *IEEE Transactions on Reliability* **R-28**, 137–140.
- [4] Wilson, J.G. & Benmerzouga, A. (1990). Optimal m-failure policies with nonzero repair time, *Operations Research Letters* **9**, 203–209.
- [5] Okumoto, K. & Elsayed, E.A. (1983). An optimum group maintenance policy, *Naval Research Logistics Quarterly* **30**, 667–674.
- [6] Ritchken, P.H. & Wilson, J.G. (1990). (m, T) group maintenance policies, *Management Science* **36**, 632–639.
- [7] Benmerzouga, A. & Harous, S. (1999). Optimal (m-Failure, P-Repairmen) policies with random repair time, *Stochastic Analysis and Applications* **17**, 327–338.
- [8] Wilson, J.G. & Benmerzouga, A. (1995). Bayesian group replacement policies, *Operations Research* **43**, 471–476.
- [9] Wilson, J.G. & Popova, E. (1996). Adaptive replacement policies for a system of parallel machines, in *Lifetime Data: Models in Reliability and Survival Analysis*, N.P. Jewell *et al.*, eds, Kluwer Academic Publishers, pp. 371–375.
- [10] Mazzuchi, T.A. & Soyer, R. (1996). A Bayesian perspective on some replacement strategies, *Reliability Engineering and Systems Safety* **61**, 295–303.
- [11] Heyman, D.P. & Sobel, M.J. (1984). *Stochastic Models in Operation Research*, McGraw-Hill Book Company, Vol. II.
- [12] Nakagawa, T. (1979). Further results of replacement problem of a parallel system in a random environment, *Journal of Applied Probability* **16**, 923–926.

Related Article

Maintenance Optimization in Random Environments; Replacement Strategies; Imperfect Repair; Maintenance Optimization; Multivariate Age and Multivariate Renewal Replacement; System Availability; General Minimal Repair Models; Multi-component Maintenance; Age-Dependent Minimal Repair and Maintenance; Block Replacement .

JOHN G. WILSON, ALI BENMERZOUGA AND
CHRIS K. ANDERSON

Replacement Strategies

Definitions and Main Characteristics

Our daily lives involve the observation, use, construction, and destruction of many systems. We replace the light bulbs when they burn out, take our cars to the repair shop when the brakes need to be replaced or the oil is due to be changed. The situations requiring a *replacement* are usually connected to wear out, aging, deterioration, or failure of the item/system involved. In this presentation we will assume that all of these processes are random.

Definition: A replacement policy π is a decision making rule that defines the time and type of replacement of an item or a system such that an objective defined by the decision maker is optimized. The main characteristics of any replacement policy are:

Objective function – A common criteria in deciding when to perform a replacement is economically justified, that is a cost function (called *objective function*), $G(x)$, is optimized with respect to a set of parameters. Examples are the smallest average replacement cost per unit time or the total discounted cost. Alternative objectives are the maximum reliability, minimum net present cost, internal rate of return (IRR), or any general utility function defined by the decision maker.

Time of replacement – A replacement can be performed: as soon as the item fails (also known as corrective maintenance), before failure (preventive maintenance), or in a certain amount of time after failure. This will be one of the objective function's parameters, call it T . We will try to find the value of T that optimizes the objective function.

Type of replacement – The item can be replaced with a new one or with an old (used) one. This is another one of the objective function parameters, denote it by a . As before, we would like to find the age of the system with which we will replace the “old” one that optimizes the objective function.

Failure time – The failure occurrences are assumed to be random following a general counting stochastic process $\{N(t), t \geq 0\}$, see Rausand and Høyland [1, p. 231]. The “failure time” will be the time between two consecutive failures.

We can formulate the general stochastic optimization problem of finding the optimal replacement policy $\pi = (T^*, a^*)$ as

$$\min_{T, a \in R^+} E[G(T, a, t)] \quad (1)$$

where E stands for the expectation operator taken with respect to a filtration up to time t defined by the counting process $N(t)$. The decision variables are T – time of replacement and a – age of the replacing item. Note that if we want to find the maximum of the objective function, it will be equivalent to finding the minimum of its negative value.

Here are some common replacement policies one can obtain for specific values of the time of replacement, T :

- Age replacement policy: the item/system is replaced upon failure or at age Y , whichever comes first (see **Multivariate Age and Multivariate Renewal Replacement**).
- Block replacement: the item/system is replaced at regular time intervals $X, 2X, \dots$ regardless of age (see **Block Replacement**).
- Group replacement: a group of items is replaced at the same time to take advantage of economies of scale (see **Group Maintenance Policies**).
- Condition-based replacement: a collection of variables measuring the state of system's **degradation** is monitored and replacement decision is made based on their values (see **Degradation Models**).
- Opportunity-based replacement: the replacement is performed at a time when *opportunity* arrives. Examples of opportunities are: scheduled downtime, lunch breaks, failure of a system in close proximity to the item of interest. The arrival process of the opportunities is assumed to be random.

There are three main replacement types based on the values of a :

- The failed item/system is replaced with a brand new one with age 0. This type of replacement is called *as good as new* or also known as **preventive maintenance**.
- $a = a_t$, where t is the time since last replacement and a_t is the age of the failed item. This replacement is called *as good as old*, or **minimal repair** (see **General Minimal Repair Models**).

- The age of the replacing item is different from 0 or the age of the failed item. This is called *imperfect repair* (see **Imperfect Repair; Multivariate Imperfect Repair Models**).

The literature on replacement policies is enormous. The earlier work was done in the early 1960s by Barlow and Proschan [2]. An excellent survey of the literature is presented by Florez and Feldman [3] and Dekker [4]. Rausand and Høyland [1] is one of the contemporary textbooks on reliability. Marquez and Heguedas [5] present a review of the recent research on **maintenance policies** and solve the problem of periodic replacement in the context of a semi-Markov decision processes methodology. We will review some of the most recent work on replacement policies and refer the readers to the above references for earlier research.

The reliability literature on replacement models and policies that allow for a change of the future failure behavior of the system is limited. There are several papers that could be classified as either models where replacement actions reduce the rate of failures or models where the replacement action reduces the (virtual) age of the system, see Rausand and Høyland [1, p. 287], for details.

Such problems where the replacement decisions may influence the future stochastic nature of the system are referred to as *decision-dependent-randomness* problems. For a general overview of the existing literature that relates to this class of problems see Morton and Popova [6]. Models with decision-dependent uncertainty are discussed by Jonsbråten [7] and Jonsbråten and Woodruff [8].

Lai *et al.* [9] analyze a single-unit system subjected to external shocks. The unit can fail due to aging or shock. This is an example where the failure rate of the system increases after a lethal shock or with aging. The policy is to replace the system after the n th shock or failure, whichever occurs first. They minimize the long-run expected cost per unit time to obtain the optimal value of n . Lai and Tang [10] consider a two-unit system where the failure of each unit either increases the failure rate of the other or brings it to an instantaneous failure. The system is replaced at age T or at failure, whichever occurs first. The value of T is obtained by minimizing the long-run expected cost per unit time.

Time Horizon and Objective Functions

An important characteristic of the replacement decisions is the time horizon over which we want to solve the optimization problem – it can be either finite or infinite.

Infinite Time Horizon

In the infinite case one can use existing limit theorems (for example from renewal theory and Markov decision process) to obtain the structure of the optimal replacement policy, see Ross [11] (see **Maintenance and Markov Decision Models; Markov Renewal Processes in Reliability Modeling**). The long-run expected cost per unit time and the total discounted cost are the two common objective functions used in this case.

Juang and Sheu [12] analyze a k -out-of- n system and find the optimal age replacement policy (see **Total Time on Test Plots**). They consider two approaches – one that minimizes the long-run cost per unit time, and a graphical approach based on the **total time on test** concept. Jiang and Ji [13] consider the age replacement policy by introducing a different objective – multiple attribute utility. In addition to cost they include the availability, reliability, and lifetime as attributes of the objective function.

The optimal replacement of an item under warranty is an important problem both for the manufacturer and the user. Iskandar and Sandoh [14] consider a system with a warranty period $[0, S]$ where opportunities for replacement occur according to a **Poisson process**. The policy is as follows: when the system fails at age $x \leq S$, a minimal repair is performed, if an opportunity occurs at age $S < x < T$, the system is replaced with probability p , and it is also replaced at age T . The optimal values of the parameters are obtained by minimizing the long-run expected cost. Iskandar and Murthy [15] study a combination of repair and replace strategies for a system under warranty. They obtain the optimal policy by minimizing the expected cost of servicing the warranty (see **Warranty Servicing**).

Finite Time Horizon

The finite time horizon is a bit more challenging since there are no existing general results from the stochastic/reliability literature to indicate which replacement

policy is optimal. Su and Chang [16] find the periodic maintenance policies that minimize the life-cycle cost over a predefined finite horizon. Galenko *et al.* [17] solve the combination of preventive replacement and minimal repair policies over a finite horizon by minimizing the total replacement cost where the time of replacement is an integer variable.

There are other objective functions used in the finite time horizon case. Net present value (NPV) is the difference between the sum of present values of the project's (replacement policy) future cash flows (computed as the difference between inflows and costs) and the initial cost of the project. The NPV approach is the most common method used in capital budgeting (since the decision when and how to replace fits within the capital budgeting framework).

When making the replacement decision, one chooses the option with the highest NPV. This simple rule does not work very well when the replacement decision involves choosing between two machines with unequal lives. Suppose that both machines can do the same job, but they have different operating costs and will last for different time periods. A simple application of the NPV rule suggests taking the machine whose costs have the lower present value. This choice might be suboptimal because the lower-cost machine may need to be replaced before the other one. The easiest way to compare the two machines involves calculating something called the *equivalent annual cost* of each machine. In other words, the NPV of the costs computed over the finite time horizon is converted into an annuity cost assuming that the machines will have to be part of the production process forever, that is into infinite horizon case. Examples of such necessities can be found in the health care system where a new admitting system may be needed every 5 years; making a decision about copy machines, replacing computers, and so on.

Another frequently used approach is the Internal Rate of Return, IRR. The basic rationale behind the IRR method is that it provides a single number summarizing the merits of a project. That number does not depend on the interest rate prevailing in the capital markets. The IRR is the discount rate that equates the NPV of the project to zero. The general decision rule is very simple: accept the project if the IRR is greater than the discount rate; reject the project if IRR is less than the discount rate. One very common cited difficulty is when this approach is

used to choose between mutually exclusive projects, which might be the case in the replacement policy framework. The problem arises since the IRR method ignores issues of scale. In other words, one project may have a higher IRR but its NPV may be low compared to the other project that has lower IRR but higher NPV. This is a very important issue in capital budgeting since the main objective of every company is to increase the shareholder's value (i.e., the value of the firm). Making a decision to choose between mutually exclusive projects by accepting the one with the highest IRR may not be the **optimal decision** from the point of view of maximizing the firm value. There are several ways in which one can solve this problem. One is to use the NPV method and make a decision based on highest NPV. Another way is to compute the incremental NPV and accept if it is greater than zero. A third one is to compare the incremental IRR to the discount rate and accept if the IRR is greater than the discount rate.

Some companies use the profitability index (PI) method to evaluate projects. It is the ratio of the present value of the future expected cash flows after initial investment divided by the amount of the initial investment. This approach shares the same advantages and disadvantages with the IRR rule. For mutually exclusive projects it ignores the issue of scale. Incremental cash flows have to be constructed and then the PI rule can be used. The project will be accepted if $PI > 1$. This rule is also useful in the capital rationing context. Suppose that the firm does not have enough capital to fund all positive NPV projects. In the case of limited funds, we cannot rank projects according to their NPVs. Instead, we should rank them according to PI. The PI measures the dollar return for the dollar invested, that is "the bang for the buck", and is useful for capital rationing. We should note that the PI does not work if funds are also limited beyond the initial time period and it is not useful over multiple time periods.

Both objectives, IRR and PI, can be applied to the replacement problem in the following situation. Suppose we have to compare an existing replacement policy to a "proposed" one (claimed to be economically better). Then we can compute the cost savings when the new policy is used, and construct both the IRR and PI objectives.

The above methods briefly describe some of the most commonly used approaches by companies. For additional examples and more detailed explanation

4 Replacement Strategies

see Ross *et al.* [18]. In reality, in addition to these methods companies use methods like payback, discounted payback, and accounting rate of return. The use of the methods varies with the industry. For example, firms that are better able to estimate cash flows are more likely to use NPV. Companies in the oil business or in energy-related businesses can do such **estimation**. But companies in the entertainment business, like motion-picture production may not be able to produce reliable cash flow estimates.

Failure Time

An important input to the optimization problem given in equation (1) is the failure time distribution, which measures the time between two consecutive failures. The exponential distribution is a common assumption, which implies that the system's failure rate function is constant over time. In this situation, preventive replacement will not be optimal and one needs to consider replacement only at failure.

It is also important to recognize the implications of the type of replacement to the future failure time distribution. The replacement with a new item corresponds to having independent and identically distributed (i.i.d.) times between failures. This scenario is the best one if we have to estimate the failure time distribution's parameters from a data set since standard statistical techniques will be applicable.

The replacement with an old item has a variety of options: the age of the "new" item is the same as the age of the item that it is replacing, or it can be different. The first case is referred to as "replacement with as good as old" (or minimal repair) and the resulting stochastic failure process is the nonhomogeneous Poisson process, see Brown and Proschan [19].

Recently, several papers analyzed a system in which failure and repair times follow a geometric process. Zhang *et al.* [20] consider a deteriorating system and find the optimal replacement time. They investigate two replacement policies: one based on the accumulated working time T of the system, and the other based on the accumulated number N of failures. The optimal values of T and N are obtained by minimizing the long-run expected cost per unit time. Zhang [21] studies a deteriorating system and obtains the optimal replacement policy based on the number of failures by minimizing the long run expected cost per unit time.

Hu and Yue [22] model a deteriorating system (according to semi-Markov process) that operates in semi-Markov environment. They show the existence of an optimal control limit policy that minimizes the discounted total cost over finite and infinite time horizon. Moustafa *et al.* [23] analyze a system that deteriorates according to a multistate semi-Markov process. They obtain a control limit policy using the policy-iteration algorithm for the expected long-run cost rate of the system. Chen and Feldman [24] formulate a Markov decision process model for a deteriorating system and show that the control limit policy for a modified repair/replacement is optimal over the space of all possible policies under the discounted cost criterion.

A different failure time model is presented by Castro and Alfa [25], who model a single-unit system with operational time according to a phase-type distribution. In addition the system is subjected to external failures that arrive as a Bernoulli process. They present two versions of the age replacement policy. Archibald *et al.* [26] use Cox regression model to estimate the underlying maintenance behavior. They develop a stochastic dynamic programming approach to obtain the optimal replacement policy for a system subject to a decay. In addition they investigate the stability of the results to changes in the cost parameters.

Castanier *et al.* [27] propose a combined replacement policy (based on periodic and aperiodic inspections) for a stochastically deteriorating system. They model the deterioration process as a gamma process (see **Warranty: Usage and Wear Process for**) and derive the long-run expected cost per unit time for this policy. Barros *et al.* [28] focus on a parallel system of two units with dependent failures and subject to external shocks that arrive according to a Poisson process (see **Poisson Processes**). The failures, however, are detected with a given probability. The replacement policy is a group replacement, that is both items are replaced regardless of which one has failed. The objective function is the long-run expected cost per unit time.

A different approach to failure times modeling is via the Bayesian perspective. The classical failure time modeling assumes a parametric distribution (like exponential or Weibull) with parameter values that are, for instance, the maximum-likelihood estimates. Standard Bayesian modeling assumes a parametric lifetime distribution and makes inferences about its parameters via the posterior distribution

derived using Bayes' theorem, see Bernardo and Smith [29]. A review of the Bayesian approaches to maintenance intervention is presented in Wilson and Popova [30]. Chen and Popova [31] propose two types of Bayesian policies that learn from the failure history and adapt the next maintenance point accordingly. They find that the optimal time to observe the system depends on the underlying failure distribution. A combination of **Monte Carlo simulation** and optimization methodologies is used to obtain the problem's solution. Another related reference is Mazzuchi and Soyer [32]. In Popova [33], the optimal structure of Bayesian group-replacement policies for a parallel system of n items with exponential failure times and random failure parameter is presented. The paper shows that it is optimal to observe the system only at failure times. For the case of two items operating in parallel the exact form of the optimal policy is derived.

A more general approach is to use nonparametric Bayesian approach which, for instance, takes the hazard rate of the failure time to be a random "parameter", and puts the continuum of all distributions as its prior. This notion of placing a prior on function spaces is termed *nonparametrics*. Ferguson [34] was the first to consider this idea from a Bayesian perspective. Since then, there have been hundreds of papers in a variety of contexts that use Bayesian nonparametric models, see Dey and Rao [35]. Once a prior has been put on the space of hazard rates, the posterior distribution of the hazard rate process is derived. Typically, this posterior distribution will not have a closed form solution. **Markov chain Monte Carlo** (MCMC) methods are used to obtain inferences from the posterior distributions of interest; see, for example, Krishnan *et al.* [36]. Merrick *et al.* [37] discuss optimal **replacement strategies** for machine tools when their failure behavior is modeled via a Bayesian semiparametric proportional hazards model. The problem of a finite horizon, single-item maintenance optimization structured as a combination of preventive and corrective maintenance in a nuclear power plant environment is analyzed by Damien *et al.* [38]. In addition they present Bayesian semiparametric models to estimate the failure time distribution and costs involved.

A real-world application of replacement strategies is presented by Brezavscek and Hudoklin [39], who develop a block-replacement and spare-provision policies for the electric locomotives in Slovenian

Railways. Aka *et al.* [40] investigate the relationship between a replacement and inventory policy for a manufacturing system.

Acknowledgments

This research has been partially supported by NSF grant #DMI-0457558 and South Texas Project Nuclear Operating Company grant #B02857.

References

- [1] Rausand, M. & Høyland, A. (2004). *System Reliability Theory: Models, Statistical Methods, and Applications*, 2nd Edition, John Wiley & Sons.
- [2] Barlow, R.E. & Proschan, F. (1965). *Mathematical Theory of Reliability*, John Wiley & Sons.
- [3] Valdez-Florez, C. & Feldman, R. (1989). A survey of preventive maintenance models for stochastically deteriorating single-unit systems, *Naval Research Logistics* **36**, 419–446.
- [4] Dekker, R. (1996). Applications of maintenance optimization models: a review and analysis, *Reliability Engineering and Systems Safety* **51**, 229–240.
- [5] Marquez, A. & Heguedas, A. (2002). Models for maintenance optimization: a study for repairable systems and finite time periods, *Reliability Engineering and System Safety* **75**, 367–377.
- [6] Morton, D. & Popova, E. (2001). Monte Carlo simulations for stochastic optimization, in *Encyclopedia of Optimization*, C. Floudas & P. Pardalos, eds, Kluwer Academic Publishers.
- [7] Jonsbråten, T. (1998). Oil field optimization under price uncertainty, *Journal of the Operational Research Society* **49**, 811–818.
- [8] Jonsbråten, T.W., Woodruff, D.L. & Wets, R.J.B. (1998). A class of stochastic programs with decision dependent random elements, *Annals of Operations Research* **82**, 83–106.
- [9] Lai, M.-T., Shih, W. & Tang, K.-Y. (2006). Economic discrete replacement policy subject to increasing failure rate shock model, *The International Journal of Advanced Manufacturing Technology* **27**, 1242–1247.
- [10] Lai, M.-T. & Tang, K.-Y. (2006). Optimal periodic replacement policy for a two-unit system with failure rate interaction, *The International Journal of Advanced Manufacturing Technology* **29**, 367–371.
- [11] Ross, S.M. (1996). *Stochastic Processes*, John Wiley & Sons.
- [12] Juang, M.-G. & Sheu, S.-H. (2003). Graphical approach to replacement policy of a k-out-of-n system subject to shocks, *International Journal of Reliability, Quality and Safety Engineering* **10**(1), 55–68.
- [13] Jiang, R. & Ji, P. (2002). Age replacement policy: a multi-attribute value model, *Reliability Engineering and Systems Safety* **76**, 311–318.

- [14] Iskandar, B.P. & Sandoh, H. (1999). An opportunity-based age replacement policy considering warranty, *International Journal of Reliability, Quality and Safety Engineering* **6**(3), 229–236.
- [15] Iskandar, B.P. & Murthy, D.N.P. (2003). Repair-replace strategies for two-dimensional warranty policies, *Mathematical and Computer Modelling* **38**, 1233–1241.
- [16] Su, C.-T. & Chang, C.-C. (2000). Minimization of the life cycle cost for a multistate system under periodic maintenance, *International Journal of Systems Science* **31**, 217–227.
- [17] Galenko, A., Morton, D., Popova, E., Kee, E., Grantom, R. & Sun, A. (2005). Operational level models and methods for risk informed nuclear asset management, in *Proceedings of the American Nuclear Society International Topical Meeting on Probabilistic Safety Assessment*, San Francisco.
- [18] Ross, S., Westerfield, R.W. & Jaffe, J. (2005). *Corporate Finance*, 7th Edition, McGraw-Hill.
- [19] Brown, M. & Proschan, F. (1983). Imperfect repair, *Journal of Applied Probability* **20**, 851–859.
- [20] Zhang, Y.L., Yam, R.C.M. & Zuo, M.J. (2001). Optimal replacement policy for a deteriorating production system with preventive maintenance, *International Journal of Systems Science* **32**(10), 1193–1198.
- [21] Zhang, Y.L. (2004). An optimal replacement policy for a three-state repairable system with a monotone process model, *IEEE Transactions on Reliability* **53**(4), 452–457.
- [22] Hu, Q. & Yue, W. (2003). Optimal replacement of a system according to a semi-Markov decision process in a semi-Markov environment, *Optimization Methods and Software* **18**(2), 181–196.
- [23] Moustafa, M.S., Maksoud, E.Y.A. & Sadek, S. (2004). Optimal major and minimal maintenance policies for deteriorating systems, *Reliability Engineering and Systems Safety* **83**, 363–368.
- [24] Chen, M. & Feldman, R.M. (1997). *European Journal of Operational Research* **98**, 75–84.
- [25] Castro, I.T. & Alfa, A.S. (2004). Lifetime replacement policy in discrete time for a single unit system, *Reliability Engineering and System Safety* **84**, 103–111.
- [26] Archibald, T.W., Ansell, J.I. & Thomas, L.C. (2004). The stability of an optimal maintenance strategy for repairable assets, *Journal of Process Mechanical Engineering* **218**, 77–82.
- [27] Castanier, B., Grall, A. & Bérenguer, C. (2001). A stochastic model for hybrid maintenance policies evaluation and optimization, *International Journal of Reliability, Quality and Safety Engineering* **8**(3), 233–248.
- [28] Barros, A., Bérenguer, C. & Grall, A. (2006). A maintenance policy for two unit parallel systems based on imperfect monitoring information, *Reliability Engineering and System Safety* **91**, 131–136.
- [29] Bernardo, J.M. & Smith, A. (1994). *Bayesian Theory*, John Wiley & Sons.
- [30] Wilson, J. & Popova, E. (1996). Bayesian approaches to maintenance intervention, in *Proceedings of the Section on Bayesian Science of the American Statistical Association*, pp. 244–248.
- [31] Chen, T. & Popova, E. (2000). Bayesian maintenance policies during a warranty period, *Communications in Statistics: Stochastic Models* **16**, 121–142.
- [32] Mazzuchi, T.A. & Soyer, R. (1996). Adaptive Bayesian replacement strategies, in *Bayesian Statistics 5*, J.O. Berger, J.M. Bernardo, A.P. Dawid & A.F.M. Smith eds, Oxford University Press, pp. 667–674.
- [33] Popova, E. (2004). Basic optimality results for Bayesian group replacement policies, *Operations Research Letters* **32**, 283–287.
- [34] Ferguson, T. (1973). A Bayesian analysis of some non-parametric problems, *Annals of Statistics* **1**, 209–230.
- [35] Dey, D. & Rao, C. (eds) (2006). *Handbook of Statistics, Bayesian Thinking: Modeling and Computation*, Elsevier, Vol. 25.
- [36] Krishnan, M., Ramaswamy, V., Meyer, M. & Damien, P. (1999). Customer satisfaction for financial services, *Management Science* **45**(9), 1192–1209.
- [37] Merrick, J.R.W., Soyer, R. & Mazzuchi, T.A. (2003). A Bayesian semiparametric analysis of the reliability and maintenance of machine tools, *Technometrics* **45**(1), 58–69.
- [38] Damien, P., Galenko, A., Popova, E. & Hanson, T. (2007). Bayesian semiparametric analysis for a single item maintenance optimization, *European Journal of Operational Research* **182**, 794–805.
- [39] Brezavscek, A. & Hudoklin, A. (2003). Joint optimization of block-replacement and periodic-review spare-provisioning policy, *IEEE Transactions on Reliability* **52**(1), 112–117.
- [40] Aka, M., Gilbert, S.P. & Ritchken, P. (1997). Joint inventory/replacement policies for parallel machines, *IIE Transactions* **29**, 441–449.

Related Articles

Block Replacement; Burn-In and Maintenance Policies; Degradation Models; General Minimal Repair Models; Group Maintenance Policies; Imperfect Repair; Maintenance and Markov Decision Models; Maintenance Optimization; Markov Renewal Processes in Reliability Modeling; Multicomponent Maintenance; Multivariate Age and Multivariate Renewal Replacement; Multivariate Imperfect Repair Models; Maintenance Optimization in Random Environments; Stationary Replacement Strategies; Warranty Servicing.

ELMIRA POPOVA AND IVILINA POPOVA

Imperfect Repair

Introduction

When a system fails, one has a choice of replacing it with a brand new system or repairing it. When a system is replaced with a brand new system, the lifetime of the replaced system has the same distribution F as the original system. What will be the distribution of the lifetime of a repaired system? The term *imperfect repair* is used both to describe the nature of the repair and the distribution of the lifetime of the repaired system and comes in many flavors. In fact, replacement by a brand new system itself can be called a *perfect repair*. Ascher [1] introduced the concept of *minimal repair* in which the repaired system has the same distribution as a working new system with the age of the recently failed system. In symbols, the survival function of the minimally repaired system is given by $\bar{F}_t$, where

$$\bar{F}_t(s) = \frac{1 - F(s+t)}{1 - F(t)}, \quad s > 0 \quad (1)$$

and t denotes the time of the previous failure.

Let $\{S_j\}$ denote the failure times of the system maintained by repairs and let $\{T_j\}$ the interfailure times, where $T_j = S_j - S_{j-1}$ and $S_0 \stackrel{\text{def}}{=} 0$. A stochastic model for repairs of such a system (also called a *repair model*) should specify the joint distribution of $\{T_1, T_2, \dots\}$.

Note that F is the distribution of T_1 , the time of first failure of the system. Under a *perfect repair model*, $\{T_1, T_2, \dots\}$ are independent and identically distributed according to F . Again, $P(T_2 > s | T_1 = t) = F_t(s)$ when a minimal repair is performed after the first failure. More generally, $P(T_{n+1} > s | A_n = t) = \bar{F}_t(s)$, where A_n is the effective age of the system following the first n failures and repairs. Brown and Proschan (BP) [2] describe another repair model by stipulating that, following any system failure, a perfect repair or a minimal repair is made with probability p and $1 - p$, respectively where p is in $[0, 1]$. This will allow one to write down the joint distribution of $\{T_1, T_2, \dots\}$ to completely define the *imperfect repair model*. This will be called the BP repair model. Another repair model was introduced by Block, Borges, and Savits [3] (BBS). In this model, the p of the BP model is generalized to

be a function of time t and upon the n th failure a perfect repair is made with probability $p(A_n)$ and a minimal repair otherwise, where A_n is the age of the system calculated from the beginning or from the last full repair, whichever came later. More properties of these repair models are described in the section titled “Repair Models”.

There have been several other repair models described in the literature including the Dorado, Hollander, and Sethuraman (DHS) [4] model, the Kijima [5] models, and the Last and Szekli [6] model. The DHS model is more fully described in **Nonparametric Methods for Analysis of Repair Data**.

A common inference problem in repair models is the **estimation** of F . This will allow the estimation of other features of the repair model. In the section entitled “Inference in Repair Models”, we describe the nonparametric maximum-likelihood estimator (NPMLE) of F derived by Whitaker and Samaniego (WS) [7]. Hollander *et al.* [8] (HPS) used counting processes to find more properties of the WS estimator. This is described in the section titled “Inference using Counting Processes”. The estimation of F in the DHS model is described in **Nonparametric Methods for Analysis of Repair Data**. The section titled “Other Inference Studies on Imperfect Repair” considers additional inferential methods for imperfect repair.

Repair Models

The BP and BBS models discussed in the section titled “Introduction” involve the basic notions of *perfect repair*, which restores a repairable system, upon failure, to its condition when new and *minimal repair*, which restores a failed system to its condition just prior to failure. When a *perfect repair* is performed, the repaired system has the lifetime distribution F of a new system. When a *minimal repair* is made, the repaired system has the distribution of the residual lifetime of a system of age t , where t is the failure time of the system that is, the survival function is given by $\bar{F}_t(s)$ defined in equation (1). BP and BBS introduce a stochastic mechanism for determining whether a given repair will be perfect or minimal.

The BP model for imperfect repair simply stipulated that, following any system failure, a perfect

2 Imperfect Repair

repair was made with probability $p \in [0, 1]$, with a minimal repair being made otherwise. Under this modeling assumption, BP showed that the distribution of the time Y until the first perfect repair is $\bar{F}_Y(t) = (\bar{F}(t))^p$, where $\bar{F} = 1 - F$ is the survival function of the system when new. BP noted the interesting fact that the model implied the proportional hazard property for Y , (i.e., $R_Y(t) = pR(t)$ for Y , where $R = -\ln(\bar{F})$ is the hazard function of the system when new), obtaining this property as a derived result rather than as an initial assumption. Proschan, well-known for his sense of humor, liked to say that this property was evidence that “among the various ways one might model imperfect repair, God favored this particular formulation”. A number of preservation properties for the model were established in BP, including results such as: if F belongs to the IFR, DFR, IFRA, DFRA, NBU, or NWU classes, the distribution F_Y belongs to the same class. The paper also proved several monotonicity results, including (a) $E(Y)$ is an increasing (decreasing) function of the model parameter p when F belongs to the NBUE (NWUE) class and (b) $N(t)$, the number of failures in the interval $[0, t]$ is stochastically decreasing (increasing) in p when F is NBU (NWU). The paper is generally considered the seminal contribution on “imperfect repair”.

The BBS model, which generalizes the BP model, allows the probability of a perfect repair following any given failure of the system to depend on the age t at which the system failed; this probability is thus represented by a function $p(t)$ taking values in $[0, 1]$ for $t > 0$. While the model is well defined for any choice of such a function $p(\cdot)$, the paper shows that the time Y until a perfect repair is finite with probability 1 if and only if

$$\int_0^\infty p(t) dR(t) = \infty \quad (2)$$

In this case, assuming for simplicity that F is absolutely continuous, the distribution F_Y of Y is given by

$$F_Y(t) = 1 - \exp \left\{ - \int_0^t \frac{p(s)}{(1 - F(s))} dF(s) \right\} \quad (3)$$

and has failure rate $r_Y(t) = p(t)r(t)$, where $r(t)$ is the failure rate of F . The intuition that $p(t)$

might well be a nondecreasing function in any real application, making a perfect repair more likely as the system grows older, is well accommodated by the BBS model. While monotonicity of p is not required in general, several results are derived for this important special case. BBS also identify the successive times until complete repair as a specific renewal process. When $p(t)$ is increasing (decreasing), BBS prove that F_Y is IFR, IFRA, NBU or DMRL (DFR, DFRA, NWU or IMRL) when F belongs to the same class. The paper also contains some monotonicity properties of the model and provides an example showing that F_Y need not be NBUE when F belongs to the NBUE class.

A number of subsequent papers have generalized the BP or the BBS model in a variety of ways. Notable among these are the multivariate generalizations studied by Shaked and Shantikumar [9] Savits [10] Sheu and Griffith [11, 12] and Hu [13]. Several papers have shed light on various forms of stochastic ordering between or among random variables obtained *via* the BP or the BBS model. Contributions in this direction may be found in Hu [13], Last and Szekli [6] and Belzunce and Shaked [14]. Papers that study imperfect repair models in special contexts include Finkelstein [15], where these models are examined for systems subject to shocks, Biswas and Sarkar [16], where imperfect repair is studied with maintainability and availability issues in mind, and Li and Shaked [17], where **preventive maintenance** strategies are studied within an imperfect repair framework.

Inference in Repair Models

The first published treatment of problems of nonparametric statistical inference for imperfect repair models is that of WS. These authors addressed the problem of estimating the unknown parameters (p, F) of the BP model and $(p(\cdot), F)$ of the BBS model. HPS, using different techniques, provide a broader asymptotic analysis of the WS estimator, while Presnell *et al.* [18] (PHS) examine the question of testing the “minimal repair” assumption of the BP and BBS models. These papers have been followed by a number of inference papers on imperfect repair. The sections titled “Inferences using Counting Processes” and “Other Inference Studies on Imperfect Repair” discuss several such papers with a view toward

providing the reader with a sense of the scope of generalizations and alternative approaches contained in the statistical literature on imperfect repair models.

WS took a NPMLE approach to estimation in the BP and BBS models. They demonstrate that the BP and BBS models are not identifiable on the basis of observed interfailure times alone and that data on the mode of repair after each observed failure would need to be observed (or imputed in some justifiable fashion) in order to estimate p and F using classical statistical methods. WS assume that, under inverse sampling until the occurrence of the m th perfect repair, data of the form $\{(t_i, z_i), i = 1, 2, \dots, n\}$ is available, where t_i is the age of the system at the time of the i th failure and $z_i = 1$ if the i th repair is perfect and $z_i = 0$ otherwise. Under this assumption, WS show that the NPMLE is well defined only when $z_{(n-1)} = 1$, where $z_{(n-1)}$ is the value of z corresponding to the second oldest age of failure. Under the BP model, WS show that the pair of estimators $\hat{F}(t) = 1 - \prod_{j=1}^i k_j / (k_j + 1)$ for $x_{(i)} \leq t < x_{(i+1)}$, where $k_i = \sum_{j=1}^i z(j)$, and $\hat{p} = m/n$ to be the NPMLE when it exists and to be a “neighborhood NPMLE” with respect to the sup norm in general. WS then prove the weak convergence of the process $\sqrt{m}(\hat{F}(t) - F(t))$ to a zero-mean Gaussian process on any fixed interval $[0, T]$ with $F(T) < 1$, and show that $\hat{F}$ has a smaller asymptotic variance than two natural alternative estimators of F . Extensions to estimation for the BBS model and to inference in renewal testing are also discussed.

Inference using Counting Processes

Using counting process techniques, HPS extended the large sample theorems of WS to the whole line. HPS also derived a nonparametric asymptotic simultaneous confidence band for F . HPS assumed

$$F(\tau_F) = 1, \int_{(0, \tau_F)} p(t) \frac{dF(t)}{\bar{F}(t-)} = \infty \quad (4)$$

where $\tau_F = \inf\{t \in [0, \infty] : F(t) = 1\}$. They considered the situation where one observes n independent copies of the BBS process, each until the time of its first perfect repair. Let T be the first failure age at which only one item is at risk. Let $X_{(k)}$ denote the ordered values of the observed failure ages of

the n BBS processes. The WS estimator can be written as

$$\hat{\bar{F}} = \prod_{X_{(k)} \leq t \wedge T} \left(1 - \frac{1}{Y(X_{(k)})}\right) \quad (5)$$

where $Y(X_{(k)})$ is the number of processes still being observed just prior to $X_{(k)}$.

Because the BP and BBS models are based on an assumption that imperfect repairs are indeed minimal repairs, it is useful to test that assumption. Presnell *et al.* [18] constructed two asymptotically distribution-free tests of the minimal repair assumption in the BBS model. Let $\hat{\bar{F}}_e$ denote the empirical survival function based on the initial failure times of the systems of the systems being observed. PHS developed a Kolmogorov–Smirnov test based on the maximum absolute difference between $\hat{\bar{F}}_e$ and $\hat{\bar{F}}$, the WS estimator. PHS also developed a Wilcoxon-type test based on $\int_0^\infty \hat{\bar{F}}_e d\hat{\bar{F}}$ (see **Nonparametric Methods for Analysis of Repair Data**).

Other Inference Studies on Imperfect Repair

Among the subsequent inferential studies involving imperfect repair, several papers provide analytical treatments in substantially new settings – either by considering different **sampling schemes**, by placing additional constraints on the underlying model or by taking a different inferential approach. Examples of such papers include the following. Lim [19] considers the estimation of the underlying distribution F in the BP model using masked data; the case when F is Weibull is given special attention. Sethuraman and Hollander [20, 21] introduced a new class of Partition Based priors and showed that it is a conjugate class for a large class of general repair models that includes the BBS model. They used this class and a subclass called Partition Based Dirichlet priors, which also forms a conjugate family, to develop Bayesian nonparametric estimators for repair models, including a Bayesian nonparametric competitor of the WS estimator. Baker [22] proposed data-based modeling of the failure rate in an imperfect repair setting. Kvam *et al.* [23] derived the NPMLE $\hat{r}(t)$ of the failure rate of a system subject to imperfect repair and the additional condition that the system’s lifetime distribution belongs to the IFR class. They established

the strong consistency of $\hat{r}(t)$ in an interval $[a, b]$ under the assumption that the true failure rate $r(t)$ is continuous and increasing in that interval.

References

- [1] Ascher, H. (1968). Evaluation of repairable systems using the "Bad as Old" concept, *IEEE Transactions on Reliability* **17**, 105–110.
- [2] Brown, M. & Proschan, F. (1983). Imperfect repair, *Journal of Applied Probability* **20**, 851–859.
- [3] Block, H., Borges, W. & Savits, T. (1985). Age dependent minimal repair, *Journal of Applied Probability* **22**, 370–385.
- [4] Dorado, C., Hollander, M. & Sethuraman, J. (1997). Nonparametric estimation for a general repair model, *Annals of Statistics* **25**, 1140–1160.
- [5] Kijima, M. (1989). Some results for repairable systems with general repair, *Journal of Applied Probability* **26**, 89–102.
- [6] Last, G. & Szekli, R. (1998). Stochastic comparison of repairable systems by coupling, *Journal of Applied Probability* **35**, 348–370.
- [7] Whitaker, L.R. & Samaniego, F.J. (1989). Estimating the reliability of systems subject to imperfect repair, *Journal of the American Statistical Association* **84**, 301–309.
- [8] Hollander, M., Presnell, B. & Sethuraman, J. (1992). Nonparametric methods for imperfect repair models, *Annals of Statistics* **20**, 879–896.
- [9] Shaked, M. & Shanthikumar, J.G. (1987). Multivariate hazard rates and stochastic ordering, *Advances in Applied Probability* **19**, 123–137.
- [10] Savits, T.H. (1988). Some multivariate distributions derived from a non-fatal shock model, *Journal of Applied Probability* **25**, 383–390.
- [11] Sheu, Sh-H. & Griffith, W.S. (1991). Multivariate age-dependent imperfect repair, *Naval Research Logistics* **38**, 839–850.
- [12] Sheu, Sh-H. & Griffith, W.S. (1992). Multivariate imperfect repair, *Journal of Applied Probability* **29**, 947–956.
- [13] Hu, T. (1996). Stochastic comparisons of order statistics under multivariate imperfect repair, *Journal of Applied Probability* **33**, 156–163.
- [14] Belzunce, F. & Shaked, M. (2001). Stochastic comparisons of mixtures of convexly ordered distributions with applications in reliability theory, *Statistics and Probability Letters* **53**, 363–372.
- [15] Finkelstein, M.S. (1997). Imperfect repair models for systems subject to shocks, *Applied Stochastic Models and Data Analysis* **13**, 385–390.
- [16] Biswas, A. & Sarkar, J. (2000). Availability of a system maintained through several imperfect repairs before a replacement or a perfect repair, *Statistics and Probability Letters* **50**, 105–114.
- [17] Li, H. & Shaked, M. (2003). Imperfect repair models with preventive maintenance, *Journal of Applied Probability* **40**, 1043–1059.
- [18] Presnell, B., Hollander, M. & Sethuraman, J. (1994). Testing the minimal repair assumption in an imperfect repair model, *Journal of the American Statistical Association* **89**, 289–297.
- [19] Lim, T.J. (1998). Estimating system reliability with fully masked data under Brown-Proschan imperfect repair model, *Reliability Engineering and System Safety* **59**, 277–289.
- [20] Sethuraman, J. & Hollander, M. (2006a). Bayesian methods in repair models, in *Proceedings of the International Conference on Degradation, Damage, Fatigue and Accelerated Life Models in Reliability Testing*, Angers, 217–221.
- [21] Sethuraman, J. & Hollander, M. (2006b). Nonparametric Bayes Estimation in Repair Models, submitted.
- [22] Baker, D. (2001). Data-based modeling of the failure rate of repairable equipment, *Lifetime Data Analysis* **7**, 65–83.
- [23] Kvam, P.H., Singh, H. & Whitaker, L.R. (2002). Estimating distributions with increasing failure rate in an imperfect repair model, *Lifetime Data Analysis* **8**, 53–67.

Related Articles

Accelerated Life Tests: Analysis with Competing Failure Modes; Multivariate Imperfect Repair Models; Repair Data, Sets of: How to Graph, Analyze, and Compare.

MYLES HOLLANDER, FRANCISCO J. SAMANIEGO AND JAYARAM SETHURAMAN

Load-Sharing Models

Introduction

Consider a system of components whose lifetimes are governed by a probability distribution. *Load sharing* refers to a model of stochastic interdependency between components that operate within a system. If components are set up in a parallel system (*see Parallel, Series, and Series-Parallel Systems*) for example, the system survives as long as at least one component is operating. In a typical load-sharing system, once a component fails, the remaining components suffer an increase in failure rate due to the extra “load” they must encumber due to the failed component.

Component interdependencies are discussed in **Dependence; Aging and Positive Dependence; Measures of Association**. Another model for component interdependency is called a *shock model*, such as Marshall and Olkin [1] bivariate exponential model, which enables the user to model stochastic dependencies between component lifetimes by incorporating latent variables to allow simultaneous component failures. For example, suppose two components in a system are characterized by failure processes Λ_1 and Λ_2 . In a shock model, we can add a new process Λ_3 that characterizes events causing both components to fail, hence the overall failure process for the components become $\Lambda_1 + \Lambda_3$ and $\Lambda_2 + \Lambda_3$. This is the basis for most dependent failure models in probabilistic risk assessment, including *Common Cause Failure* models used in the nuclear industry; see Chapter 8 of Bedford and Cooke [2].

In contrast to the shock model, the load sharing creates a *dynamic* reliability model in which component lifetime distributions may or may not change throughout the course of a system lifetime. Daniels [3] originally adopted the load-share model to describe how the strain on yarn fibers increases as individual fibers within a bundle break. In this case, reliability is measured from strength instead of time until failure. Freund [4] formalized the **probability theory** for a bivariate exponential load-share model.

In general, the shock model provides an easier avenue for multivariate modeling of system component lifetimes. However, dynamic models such as the load-share model are deemed more realistic in environments where a component’s performance can

change once another component in the system fails or degrades. In typical applications of the load-share model, the group of operating components share a load (e.g., current, stress, weight) and once a component within the system fails, the failure rate of the surviving components increase because there are fewer components to share the constant load.

The load-sharing framework can apply to general problems of detecting members of a finite population. If resources allocated for finding a finite set of items are defined globally, rather than assigned individually, then once items are detected, resources can be redistributed for the problem of detecting the remaining items. This action gives rise to a load-sharing model. For bridge construction, some beams connected to the girder are further supported by a set of welded joints. The failure of one or two welded joints can cause the increase of stress on the remaining joints, inducing a load-share model.

Load-Share Rules

Perhaps the most important element of the load-share model is the rule that governs how failure rates change after some components in the system fail. This rule depends on the reliability application and how the components within the system interact that is, through the structure function. For example, an *equal load-sharing* rule, illustrated in Figure 1, implies that the extra load caused by the failed component is shared equally among the surviving ones. A *local*

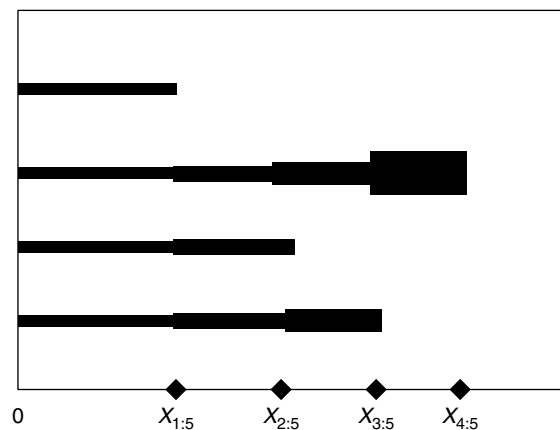


Figure 1 Illustrative load transfer with an equal load-sharing rule

2 Load-Sharing Models

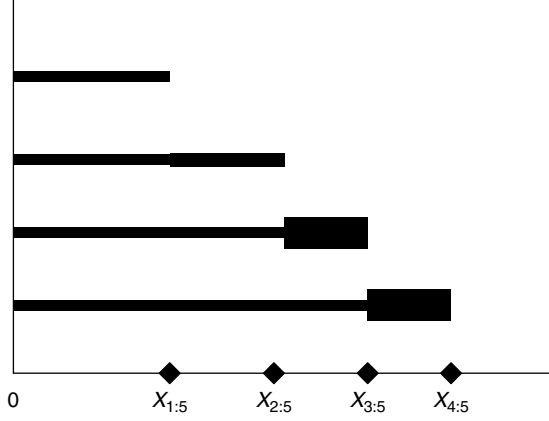


Figure 2 Illustrative load transfer with an local load-sharing rule, where the load from a failed component is transferred to its closest neighbors

load-sharing rule, illustrated in Figure 2, dictates that a failed component's load is transferred to adjacent components; the proportion of the load the surviving components inherit depends on their distance to the failed component.

More generally, a *monotone load-sharing* rule states that the load on the surviving components is nondecreasing with respect to the failure of other components in the system. Not all load-share rules need to be monotone. In finite population sampling, the event of finding one of the n observations might decrease the rate for finding the others. To model the detection rate in software debugging for example, the detection time for existing faults can depend on the number of other faults in the software that have already been found (*see Software Reliability Modeling and Analysis*). The discovery of a critical fault in the software might help reveal or conceal other yet undetected bugs. If other faults are more concealed, they have a decreased rate of discovery in the debugging process.

For assessing an unknown load-sharing rule, Kim and Kvam [5] model the dynamic nature of **component reliability** through proportional hazards. The baseline component failure time distribution is denoted by F . The hazard function (or cumulative hazard rate) corresponding to F is $R(x) = -\log(1 - F(x))$, and the hazard rate is $r(x) = f(x)[1 - F(x)]^{-1}$, where $f(x)$ is the density of F . Until the first component failure, the failure rate of each of k components in the system equals the baseline rate

$r(x)$. Upon the first failure within a system, the failure rates of the $k - 1$ remaining components jump to $\gamma_1 r(x)$, and remain at that rate until the next component failure. After this failure, the failure rates of the $k - 2$ surviving components jump to $\gamma_2 r(x)$, and so on. The failure rate of the last remaining component is $\gamma_{k-1} r(x)$. An equal load-share rule is implied by the $k - 1$ parameters $\gamma_1, \gamma_2, \dots, \gamma_{k-1}$.

To illustrate the relationship of the load-share rule to the system failure process, we consider Daniels' original application for modeling fiber strength.

Stress-Strength Models

For researchers in the textile industry who investigate the reliability of composite materials, a bundle of fibers can be considered as a parallel system subject to a steady tensile load. The rate of failure for individual fibers depends on how the unbroken fibers within the bundle share the load of this overall stress. The load-share rule of such a system depends on the physical properties of the fiber composite. Yarn bundles or untwisted cables tend to spread the stress load uniformly after individual failures. This leads to an equal load-share rule, which implies the existence of a constant system load that is distributed equally among the working components.

In more complex settings, a bonding matrix joins the individual fibers as a composite material, and an individual fiber failure affects the load of certain surviving fibers (e.g., neighbors) more than others, dictated by a local load-sharing rule. If individual fiber strengths are independently and identically distributed with distribution function F , the (parallel) system strength is

$$Y_n = \max_n \left\{ \frac{n - k + 1}{n} X_{k:n}, k = 1, \dots, n \right\} \quad (1)$$

Unfortunately, the distribution F_{Y_n} for Y_n becomes intractable quickly as n increases. Suh *et al.* [6] developed a recursive formula for F_n

$$F_{Y_n} = \sum_{i=1}^n \binom{n}{i} (-1)^{i+1} (F(x))^i F_{n-i} \left(\frac{nx}{n-i} \right) \quad (2)$$

that allows algorithmic computation, but even this solution becomes computationally infeasible with large n .

As an alternative approach, consider a system with each fiber undergoing a positive stress level x . As

x increases, more and more individual fibers break. In this scenario, however, the system load actually decreases at the instance of failure because here the system load is measured as the total stress divided by the number of surviving fibers. Now the actual load in the system (per fiber) at stress level x is

$$L(x) = xn^{-1} \sum_{i=1}^n I(X_i > x) \quad (3)$$

Because $x^{-1}nL(x)$ is binomially distributed, $L(x)$ has an asymptotic normal distribution with mean $x(1 - F(x))$ and variance $x^2F(x)(1 - F(x))$. The mean is unimodal in x , so there is a stress value x_0 such that $L(x_0)$ represents the maximum load the system can endure. For example, if $F(x) = x$, $0 \leq x \leq 1$ (strength is uniformly distributed between 0 and 1) then $x_0 = 0.5$, the maximum value of $x(1 - x)$.

Various techniques have been used to approximate $P_n(x) = n^{-1} \sum I(X_i \leq x)$ in $L(x) = x(1 - P_n(x))$. These are summarized in Crowder *et al.* [7]; Barbour [8] approximation shows a slight improvement over the others based on $P_n(x) \approx \Phi(z_x)$, where

$$z_x = \frac{\sqrt{n}(x - L_0 - \lambda \Delta^{1/3} n^{-2/3})}{\sigma^2 + \gamma n^{-1/2} \Delta^{2/3}} \quad (4)$$

$L_0 = L(x_0)$, $\sigma^2 = x_0^2 F(x_0)(1 - F(x_0))$, and

$$\Delta = \frac{x_0^4 F'(x_0)^2}{2F'(x_0) + x_0 F''(x_0)} \quad (5)$$

(λ, γ) are constant with $\gamma \approx -0.317$ and $\lambda \approx 0.996$. McCartney and Smith [9] carefully evaluated several approximation techniques, noting that the Barbour technique outperform other methods in the lower tail of the distribution.

Time-to-Failure Models

The load-share model has found broad application in life testing and dependent systems analysis since first being developed for material strength models. In most examples, load-sharing models serve as a detriment to systems, especially **parallel systems**. Most research has emphasized parametric lifetime distributions for the dependent components (e.g., Exponential, Weibull) with an assigned load-sharing rule to characterize the dependence. Rydeñ [10] extended the load-sharing framework to nonparametric **estimation** of component and

system lifetime using the same load-sharing rules of Daniels [3].

As a simple example, if the components of a parallel system have independent exponential distributions with failure rate λ , it is easy to show through spacings (the time between successive failures in the system) that the expected system lifetime is $\binom{n+1}{2} \lambda^{-1}$. However, if there is a load-sharing system in which the load of the failed component is added to the other components, the expected system lifetime is reduced to λ^{-1} . This is another consequence of the memoryless property that is, for $0 < a < t$, if $X \sim \text{Exp}(\lambda)$, $P(X > t | X > a) = P(X > t - a)$. In this case, it does not even matter if the load is shared equally by all the surviving components or whether the load is put only on one or two survivors.

Figure 3 shows the expected lifetime for parallel system of Weibull(a, b)-lifetime components. The scale parameter (a) equals one and the shape parameter (b) changes from one to five, so the left-most case refers to the exponential model, where the expected lifetime is 1.0. For the components, the failure rate is $r(t) = t^b$. Here, the load from a failed component is assumed to be distributed equally among the surviving components *via* the scale parameter as described by Kvam and Peña [11] and Liu [12].

With complex load-sharing rules, the number of parameters can escalate dramatically. Lynch [13] showed that complex systems of Weibull-distributed components, in a load-sharing model, can be modeled using a mixture of gamma-type distributions.

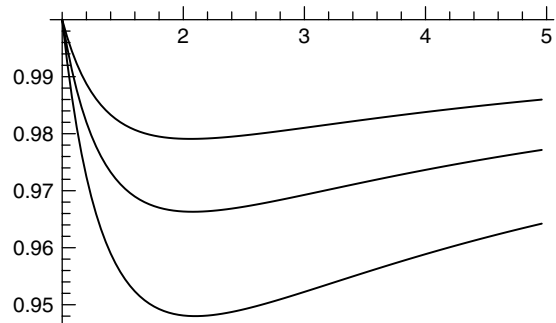


Figure 3 Expected parallel system lifetime versus shape parameter of Weibull-distributed components. The shape parameter changes from 1 (Exponential case) to 5. Curves are for cases $n = 20$ (highest), $n = 10$ (middle) and $n = 5$ (lowest)

Estimating the Load-Sharing Rule

Using nonparametric **maximum likelihood estimation** (MLE), Kvam and Peña [11] estimated a simple load-sharing rule based on independent identically distributed (i.i.d). parallel systems for which the failure rates of all components were equal, but the change in rate after a component failure depends on the set of functioning components in the system. The proportional hazard model by Kim and Kvam [5] is applied to estimate the load-sharing rule; that is, failure rate changes are characterized by $k - 1$ unknown parameters $\gamma_1, \gamma_2, \dots, \gamma_{k-1}$ and the unknown baseline distribution or hazard function. In this case, estimating the underlying baseline functions F was of primary interest.

If X_{ij} represents the lifetime of the j th component in the i th parallel system, then the random spacing T_{ij} is the time between j th failure and $(j - 1)$ th failure for the i th system. Kim and Kvam [5] considered the exponential model. For this special case where $i = 1, \dots, n$ and $j = 1, \dots, k$, the likelihood function for the i th system is

$$L_i(\theta, \boldsymbol{\gamma}; t_{i1}, t_{i2}, \dots, t_{ik}) = k! \theta^k \prod_{j=1}^{k-1} \gamma_j \exp\left(-\theta \sum_{j=1}^k (k - j + 1) \gamma_{j-1} t_{ij}\right) \quad (6)$$

Although the likelihood provides no closed form solution to obtaining MLEs, a Gauss-Seidel algorithmic method can be used to solve them. Kim and Kvam [5] also considered order restricted inference caused by monotone load sharing where $1 \leq \gamma_1 \leq \gamma_2 \leq \dots \leq \gamma_{k-1}$. In this case, a component failure causes the increase in the work-load of the other components, which can equate to an increase of failure rate.

Kvam and Peña [11] generalized the special case of Kim and Kvam [5] in which the components are distributed exponential. $S_{i,1} < S_{i,2} < \dots$ represents the successive component failure times for the i th system. With the counting processes

$$N_i(t) = \sum_{j=1}^k I(S_{i,j} \leq t), \quad i = 1, 2, \dots, n \quad (7)$$

$N_i(t)$ represents the number of component failures for the i th system that occurred by time t . We can express γ in terms of the counting processes

as $\gamma[N_i(w)] = \sum_{j=0}^{k-1} \gamma_j I(N_i(w) = j)$. To formulate the nonparametric likelihood in terms of stochastic processes, let

$$Y_i(w) = (k - N_i(w-)) I(\tau \geq w) \quad (8)$$

Define $\mathcal{F}_{it} = \sigma\{(N_i(w), Y_i(w+)); w \leq t\}$ to be the filtration generated by the i th system up to time t , and let $\mathcal{F}_t = \bigvee_{i=1}^n \mathcal{F}_{it}$. The load-share model can be described by specifying the failure intensities as

$$\begin{aligned} \Pr\{dN_i(t) = 1 | \mathcal{F}_{it-}\} \\ = r(t) Y_i(t) \gamma[N_i(t-)] dt, \quad i = 1, \dots, n \end{aligned} \quad (9)$$

Note that

$$M_i(t) = N_i(t) - \int_0^t \gamma[N_i(u-)] r(u) Y_i(u) du \quad (10)$$

is a vector of orthogonal square-integrable zero-mean martingales. See, for example, Andersen *et al.* [14]. In terms of $\{(N_i(w), Y_i(w)), 0 \leq w \leq \tau\}$: $i = 1, 2, \dots, n$, the likelihood is

$$\begin{aligned} L(R(\cdot), \boldsymbol{\gamma}) \\ = \left\{ \prod_{i=1}^n \prod_{0 \leq w \leq \tau} \pi [Y_i(w) \gamma[N_i(w-)] dR(w)]^{dN_i(w)} \right\} \\ \times \exp \left\{ - \int_0^\tau Y_i(w) \gamma[N_i(w-)] dR(w) \right\} \end{aligned} \quad (11)$$

We can estimate $R(\cdot)$ from the equation (10) by fixing $\boldsymbol{\gamma}$, in $\hat{R}(\cdot; \boldsymbol{\gamma})$. By plugging $\hat{R}(\cdot; \boldsymbol{\gamma})$ into the equation (10), we have the profile likelihood $L_p(\boldsymbol{\gamma})$ for $\boldsymbol{\gamma}$, which is then maximized in $\boldsymbol{\gamma}$ to obtain the estimator $\hat{\boldsymbol{\gamma}}$. The semiparametric estimator of $R(\cdot)$ is $\hat{R}(\cdot) = \hat{R}(\cdot; \hat{\boldsymbol{\gamma}})$. To estimate R this way, we define $J(w) = I(\sum_{i=1}^n Y_i(w) > 0)$, where $J(w) = 0$ indicates all nk components have already failed at time w . If $\boldsymbol{\gamma}$ is known, by using the zero-mean property of the martingale $\sum_{i=1}^n M_i(\cdot)$, we have

$$\hat{R}(s; \boldsymbol{\gamma}) = \int_0^s \frac{J(w) dN(w)}{\sum_{i=1}^n Y_i(w) \gamma[N_i(w-)]} \quad (12)$$

This is analogous to the derivation of the Nelson-Aalen estimator in Aalen [15].

The estimator in the equation (12) is a generalized Nelson-Aalen estimator, and is similar in structure to the hazard function estimator for tensile strengths derived by Rydeń [10]. To obtain the estimator of $R(\cdot)$ for the more general case where γ is unknown, we first obtain the profile likelihood for γ by plugging in $\hat{R}(\cdot; \gamma)$ given in equation (12) into the likelihood function in equation (10). From equation (11) and equation (12) we obtain the *profile* likelihood to be

$$L_p(s; \gamma) = \prod_{i=1}^n \pi_{0 \leq w \leq s} \left[\frac{Y_i(w) \gamma [N_i(w-)]}{\sum_{l=1}^n Y_l(w) \gamma [N_l(w-)]} \right]^{dN_i(w)} \quad (13)$$

This profile likelihood is maximized with respect to γ to obtain $\hat{\gamma}$, which is then plugged in into $\hat{R}(\cdot; \gamma)$ to obtain the semiparametric estimator of R given by

$$\hat{R}(s) = \hat{R}(s; \hat{\gamma}) \quad (14)$$

By virtue of the product representation of $\bar{F} = 1 - F$ given by $\bar{F}(s) = \prod_{0 \leq w \leq s} [1 - R(dw)]$, we can estimate $\bar{F}$ with

$$\hat{\bar{F}}(s) = \prod_{0 \leq w \leq s} [1 - \hat{R}(dw)] \quad (15)$$

The solution to the nonparametric MLE is detailed in Kvam and Peña [11], and asymptotic properties for $\hat{R}$ and γ are discussed in the section titled “Time-to-Failure Models” of that paper.

References

- [1] Marshall, A.W. & Olkin, I. (1967). A multivariate exponential distribution, *Journal of the American Statistical Association* **62**, 30–44.
- [2] Bedford, T. & Cooke, R. (2001). *Probabilistic Risk Analysis: Foundations and Methods*, Cambridge University Press.
- [3] Daniels, H.E. (1945). The statistical theory of the strength of bundles of threads, *Proceedings of the Royal Society of London. Series A* **183**, 405–435.
- [4] Freund, J. (1961). A bivariate extension of the exponential distribution, *Journal of the American Statistical Association* **56**, 971–977.
- [5] Kim, H.T. & Kvam, P.H. (2004). Reliability estimation based on system data with an unknown load share rule, *Lifetime Data Analysis* **10**, 83–94.
- [6] Suh, M.W., Bhattacharyya, B.B. & Grandage, A. (1970). On the distribution and moments of the strength of a bundle of fibers, *Journal of Applied Probability* **7**, 712–720.
- [7] Crowder, M.J., Kimber, A.C., Smith, R.L. & Sweeting, T.J. (1991). *Statistical Analysis of Reliability Data*, Chapman & Hall.
- [8] Barbour, A.D. (1981). Brownian motion and a sharply curved boundary, *Advances in Applied Probability* **13**, 736–750.
- [9] McCartney, L.N. & Smith, R.L. (1983). Statistical theory of the strength of fiber bundles, *ASME Journal of Applied Mechanics* **105**, 601–608.
- [10] Rydeń, P. (1999). *Estimation of the reliability of systems described by the Daniels load-sharing model*, Licentiate Thesis, Department of Mathematical Statistics, Umeå University, Sweden.
- [11] Kvam, P.H. & Peña, E.A. (2005). Estimating load-sharing properties in a dynamic reliability system, *Journal of the American Statistical Association* **100**, 262–272.
- [12] Liu, H. (1999). Reliability of a load-sharing k-out-of-n:G system: non-iid components with arbitrary distributions, *IEEE Transactions on Reliability* **47**, 279–284.
- [13] Lynch, J.D. (1999). On the joint distribution of component failures for monotone load sharing systems, *Journal of Statistical Planning and Inference* **78**, 13–21.
- [14] Andersen, P.K., Borgan, O., Gill, R.D. & Keiding, N. (1993). *Statistical Models Based on Counting Processes*, Springer-Verlag, New York.
- [15] Aalen, O.O. (1978). Nonparametric estimation for a family of counting processes, *Annals of Statistics* **6**, 701–726.

Further Reading

- Gücer, D.E. & Gurland, J. (1962). Comparison of statistics of two fracture models, *Journal of the Mechanics and Physics of Solids* **10**, 365–373.
- Smith, R.L. (1982). The asymptotic distribution of the strength of a series-parallel system with equal load-sharing, *Annals of Probability* **10**, 137–171.

Related Articles

Competing Risks; Cumulative Distribution Function (CDF); Dependence; Markov Processes; Measures of Association; Parallel, Series, and Series-Parallel Systems; Renewal Theory; Software Failure Data Analysis; Software Reliability Modeling and Analysis; Stochastic Deterioration; Stochastic Orders and Aging; Stress-Strength Model.

PAUL H. KVAM AND JYE-CHYI LU

Maintenance Optimization

Introduction

As the cost of maintenance is continuously increasing [1], a scientific approach to maintenance optimization is of considerable interest. Maintenance optimization is the problem of determining cost-optimal maintenance decisions for an object (system or structure or one of its components) to ensure a safe and economic operation. An important concept in maintenance optimization is that of life-cycle costing (*see Life Cycle Costs and Reliability Engineering*), where the total cost of design, building, maintenance, and demolition is considered over the entire life span of the object in question. In optimizing maintenance, the uncertainties in the time to failure and/or the deteriorating condition should be taken into account (*see Stochastic Deterioration*). Many maintenance models, under uncertain deterioration, have been developed in various fields of engineering and a large number of articles, mainly focusing on the mathematical aspects, have been published [1–12]. Herein, a number of probabilistic models for optimizing the life-cycle cost of deteriorating objects are reviewed.

Maintenance optimization techniques are applicable in the design phase and the application phase of the life span of an object. In the design phase, the initial cost of investment has to be balanced against the future cost of maintenance by optimizing the design: the larger the initial resistance, the higher the cost of investment, but the lower the cost of maintenance. In the application phase, the costs of inspection and preventive maintenance (PM) have to be balanced against the costs of corrective maintenance (CM) and failure, by optimizing intervals of inspection and PM: the more PM is carried out, the higher the costs of inspection and PM, but the smaller the costs of CM and failure.

Maintenance

Maintenance Strategies

Usually, maintenance is defined as a combination of actions carried out to restore a component to, or to “renew” it to, a specified condition in which

the component can perform its required functions. Inspections, replacements, perfect repairs, partial repairs (lifetime-extending maintenance), and minimal repairs (restoring the component to its prefailure state) are possible maintenance actions. Roughly, there are two types of maintenance: corrective maintenance (after failure) and preventive maintenance (mainly before failure). The decision diagram for the choice between CM and PM is given in Figure 1.

PM can be further subdivided into: time-based maintenance carried out at regular intervals of time, use-based maintenance carried out after a fixed cumulative use, operation or load, and condition-based maintenance (CBM) carried out at times determined by (non)periodic inspection or continuous monitoring of a component's condition. According to Moubray [13], industrial maintenance changed from corrective to preventive in 1950s; initially time-based maintenance in 1960s (such as equipment overhauls at fixed intervals) and from 1970s on more CBM (such as condition monitoring). The above maintenance strategies can be applied to a single component or groups of components.

To determine optimal maintenance decisions under stochastic deterioration (*see Stochastic Deterioration*), we distinguish two approaches [14, Chapter 1]: the actuarial approach based on the notion of lifetime (time to failure) and the physical approach based on the notion of failure due the stress exceeding the resistance (*see Stress–Strength Model*). The former is used to model time-based maintenance and the latter to model CBM.

Assumptions, Notations, and Abbreviations

In the material reviewed below, the following assumptions are made: inspection is perfect in the sense that deterioration will be observed with certainty; inspection, maintenance, and replacement take negligible time and do not degrade the component at hand.

In addition, the following notation will be used in the description of the models: c_P is the cost of preventive replacement or maintenance, c_F is the cost of corrective replacement or maintenance and failure, c_I is the cost of a single inspection, and c_U is the cost of unavailability per unit time in the event of an unrevealed failure.

Finally, the following standard abbreviations will be used for brevity of presentation: PM (preventive maintenance), CM (corrective maintenance), CBM

2 Maintenance Optimization

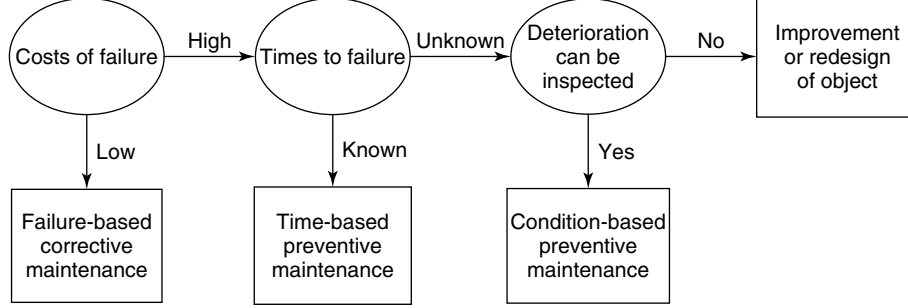


Figure 1 Decision diagram for corrective and preventive maintenance

(condition-based maintenance), and EEAC (expected equivalent average cost).

Renewal Reward Processes

Maintenance can often be modeled as a renewal process, whereby the renewals are the maintenance actions that bring a component back into its original condition or “good as new” state (*see Renewal Theory*). A renewal process $\{N(t), t \geq 0\}$ is a nonnegative integer-valued stochastic process that registers the successive renewals in the time interval $(0, t]$. Let $F(t)$ be the cumulative distribution function of the renewal time $T \geq 0$ and let $c(t)$ be the cost associated with a renewal at time t .

There are three cost-based criteria that can be used to compare maintenance decisions [15, Chapter 11]: (a) the expected average cost per unit time, (b) the expected discounted cost over an unbounded time horizon, and (c) the EEAC per unit time. Because the planned lifetime of most systems and structures is very long, maintenance decisions may be compared over an unbounded time horizon. For an overview on maintenance optimization over a bounded horizon, see [16].

Using renewal reward theory [17, Chapter 3], the expected average cost per unit time is

$$\lim_{t \rightarrow \infty} \frac{E(K(t))}{t} = \frac{\int_0^\infty c(t) dF(t)}{\int_0^\infty t dF(t)} = \frac{E(\text{cycle cost})}{E(\text{cycle length})} \quad (1)$$

where $E(K(t))$ represents the expected nondiscounted cost in the bounded time interval $(0, t]$. Let a renewal cycle be the time period between two

renewals, and recognize the numerator of equation (1) as the expected cycle cost and the denominator as the expected cycle length (mean lifetime).

To account for the time value of money, Samuelson [18] proposed to use the exponential discount function e^{-rt} with constant discount rate $r > 0$ for time $t \geq 0$. Using renewal reward theory with discounting [19, 20], the expected discounted cost over an unbounded horizon can be written as

$$c_0 + \lim_{t \rightarrow \infty} E(K(t, r)) = c_0 + \frac{\int_0^\infty e^{-rt} c(t) dF(t)}{1 - \int_0^\infty e^{-rt} dF(t)} = L(r) \quad (2)$$

where c_0 is the investment cost and $E(K(t, r))$ is the expected discounted cost in the bounded time interval $(0, t]$, $t > 0$. Similar results can be obtained for discrete-time **renewal processes** [21]. For renewal theory with other types of discounting (such as generalized hyperbolic or gamma discounting), see van der Weide *et al.* [20].

For the purpose of reserving budget for performing future maintenance actions, it is important to determine the amount of money these actions cost per unit time while taking the discounting into account. This cost is known as the *equivalent average cost per unit time* [15, Chapter 11]. The EEAC per unit time computed over an unbounded time horizon is defined as

$$EEAC = \lim_{t \rightarrow \infty} \frac{c_0 + E(K(t, r))}{\int_0^t e^{-rx} dx} = rL(r) \quad (3)$$

The EEAC per unit time can also be interpreted as a stream of fixed identical cost per unit time sufficient

to recover all the necessary discounted cost. As r tends to zero from above, the EEAC approaches the expected average cost per unit time [21]; that is,

$$\lim_{r \downarrow 0} \frac{E(K(t, r))}{\int_0^t e^{-rx} dx} = \frac{E(K(t, 0))}{t} \quad (4)$$

The cost-based criteria of discounted cost and EEAC are most suitable for balancing the initial building cost optimally against the future maintenance cost, because the contribution of the initial investment cost is not ignored (as opposed to the expected average cost per unit time). The criterion of the expected average cost per unit time can be used in situations in which no large investments are made (like inspections) and in which the time value of money is of no consequence. Often, it is preferable to spread the cost of maintenance over time and to use discounting. The area of optimizing maintenance through mathematical models based on lifetime distributions was founded in the early 1960s [2, 3]. Well-known decision models of this period are the age-replacement model and the block-replacement model.

Age Replacement

Under an age-replacement policy [2, Chapters 3–4], a replacement is carried out at age $k > 0$ (preventive replacement) *or* at failure (corrective replacement), whichever occurs first. Let the time to failure T have a cumulative distribution function $F(t)$. According to equation (1), the expected average cost of age replacement per unit time is

$$\lim_{t \rightarrow \infty} \frac{E(K(t))}{t} = \frac{c_F F(k) + c_P \bar{F}(k)}{\int_0^k t dF(t) + k \bar{F}(k)} \quad (5)$$

where k is the age-replacement interval and $0 < c_P \leq c_F$. The optimal age-replacement interval k^* is an interval for which the expected average cost per unit time is minimal. Note that PM only makes sense for deteriorating components (i.e., having increasing failure rates). According to equation (2) and Fox [22], the expected discounted cost of age replacement over

an unbounded horizon is

$$\begin{aligned} \lim_{t \rightarrow \infty} E(K(t, r)) &= \frac{c_F \int_0^k e^{-rt} dF(t) + c_P e^{-rk} \bar{F}(k)}{1 - \left[\int_0^k e^{-rt} dF(t) + e^{-rk} \bar{F}(k) \right]} \\ &= L(r) \end{aligned} \quad (6)$$

Note that the replacement model can also be applied for determining an optimal component in the design phase, which balances the initial cost of investment c_P optimally against the future cost of maintenance, by adding c_P to equation (6). The age-replacement model is one of the maintenance optimization models that has been applied most [1, 9]. For example, it has been extended with the possibility of lifetime-extending maintenance in [23].

Block Replacement

Under a block-replacement policy (*see Block Replacement*), a replacement is carried out at failure (corrective replacement) *and* periodically at the times $k, 2k, 3k, \dots$ (preventive replacement). Let the failure times $T_1, T_2, T_3, \dots$ be nonnegative, independent, identically distributed, random quantities having the cumulative distribution function $F(t)$. The expected average cost of block replacement per unit time when the decision maker chooses block-replacement interval k is

$$\lim_{t \rightarrow \infty} \frac{E(K(t))}{t} = \frac{c_F E(N(k)) + c_P}{k} \quad (7)$$

where $E(N(t))$ is the expected number of failures in $(0, t]$:

$$N(t) = \max\{j \mid S_j \leq t\} = \sum_{i=1}^{\infty} 1_{\{S_i \leq t\}}, \quad t \geq 0 \quad (8)$$

where $S_j = T_1 + \dots + T_j$, $j = 1, 2, \dots$, and 1_A denotes the indicator function of the set A . The expected discounted cost of block replacement over an unbounded horizon when the decision maker chooses block-replacement interval k is

$$\lim_{t \rightarrow \infty} E(K(t, r)) = \frac{c_F E(N(k, r)) + c_P e^{-rk}}{1 - e^{-rk}} \quad (9)$$

4 Maintenance Optimization

where $E(N(t, r))$ is the expected number of “discounted failures” in $(0, t]$; that is,

$$E(N(t, r)) = E\left(\sum_{i=1}^{\infty} e^{-rS_i} 1_{\{S_i \leq t\}}\right), \quad t \geq 0 \quad (10)$$

The block-replacement policy may be modified by performing a **minimal repair** when a component fails in the block interval [24]. Assuming the failure rate h before and after the minimal repair to be identical, then $E(N(k)) = \int_0^k h(t) dt$ and $E(N(k, r)) = \int_0^k e^{-rt} h(t) dt$.

Inspection to Detect Failures

The age- and block-replacement models assume failures to be detected immediately (revealed failures). When failures are detected only by inspection (unrevealed failures), then a different maintenance model should be used. For this purpose, we assume that the inspection times are $0 = t_0 < t_1 < t_2 < \dots$. A renewal is assumed when inspection reveals failure. According to Barlow and Proschan [2, Chapter 4], the expected costs of inspection and unavailability can be written as

$$E(\text{cycle cost}) = \sum_{j=1}^{\infty} \int_{t_{j-1}}^{t_j} [c_I j + c_U(t_j - t)] dF(t) + c_F \quad (11)$$

and the expected renewal time as

$$E(\text{cycle length}) = \sum_{j=1}^{\infty} \int_{t_{j-1}}^{t_j} t_j dF(t) \quad (12)$$

Another inspection model is the delay-time model [25, 26], which assumes the failure process to have two stages: the first stage at which the component is functioning well and the second stage at which the component is still functioning but defective, in the sense that failure can be expected soon. Periodic inspection is performed to reveal whether a component is defective or not with inspection interval τ . The time interval between the first moment at which a defect can be noticed until the time of failure is called the *delay time* and is assumed to have cumulative distribution function G . The occurrence process of defects is assumed to be a Poisson with a

rate λ . Using equation (1), the expected average cost per unit time can be written as

$$\frac{E(\text{cycle cost})}{E(\text{cycle length})} = \frac{c_I + c_F \int_0^{\tau} \lambda G(\tau - x) dx}{\tau} \quad (13)$$

Multicomponent Maintenance

When production downtime occurs or when other system component failures occur requiring downtime, maintenance opportunities may be presented for all components. When such is the case, opportunity-based maintenance models are used. These approaches use either age [27–31] or block [32, 33] **replacement strategies** with opportunity times described by a renewal, Markov or **Poisson process**.

Dependence among components can arise due to economies of scale for maintenance cost, structural relationships, or stochastic dependence of failure times [34]. Accounting for these types of dependencies has also been an important research area in maintenance optimization. Reviews of models for multicomponent maintenance can be found in [8, 35] (see **Multicomponent Maintenance**).

Condition-based Maintenance

Due to the usual lack of failure data, a reliability approach solely based on lifetime distributions is unsatisfactory. If possible, it is recommended to model deterioration in terms of a time-dependent stochastic process $\{X(t), t \geq 0\}$ where $X(t)$ is a random quantity for all $t \geq 0$. The deterioration as a function of time t , $X(t)$, is usually assumed to be a Markov process [2, Chapter 5] (see **Markov Processes**). Classes of Markov processes which are useful for modeling stochastic deterioration are discrete-time Markov processes having a finite or countable state space called *Markov chains* and continuous-time nondecreasing Markov processes with independent increments such as the compound Poisson process and the gamma process. Compound Poisson processes (gamma processes) are jump processes with a finite (infinite) number of jumps in a bounded time interval (for a sample path of the gamma process, see Figure 2 taken from [36]). Compound Poisson processes are suitable for modeling

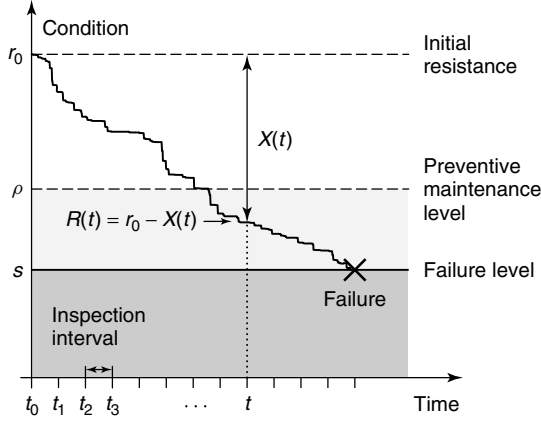


Figure 2 Condition-based maintenance model under gamma process deterioration

damage due to sporadic shocks and gamma processes for describing gradual damage by continuous use [36].

Monotone Jump Process. Modeling the deterioration as a monotone stochastic jump process is especially suited when inspections are involved [36]. In optimizing periodic inspection, the two decision variables are the inspection interval τ and the PM level ρ (Figure 2). Such a policy is called a *control-limit policy* with the PM level called the *control limit* [11]. The deterioration is inspected periodically (i.e., $t_j = j\tau$, $j = 1, 2, \dots$) and is regarded as a gamma process. At an inspection, the object can be in a functional (good), marginal, or failed state. If the object is found to be in a functional state, no maintenance is required and only the cost of the inspection is incurred. If it is found to be in a marginal state, the cost of PM is added to the cost of the inspection. For an object in failed state, there are two scenarios: either the failure is immediately noticed at the time of occurrence [37–41], without the necessity of an inspection, or failure is only detected at the next planned inspection [42–44]. In the first case, only the cost of CM is incurred and in the second case, the cost of the inspection and unavailability has to be included as well. A renewal can be PM or CM and it brings the object back to its “good as new” condition. The optimal maintenance decision is determined by minimizing the long-term expected average cost per unit time in equation (1).

When a failure is detected immediately, the expected renewal-cycle cost is the sum of the costs of all inspections during the cycle and either PM or CM:

$$E(\text{cycle cost}) = \sum_{j=1}^{\infty} [(jc_I + c_P) \Pr \{\text{PM in } (t_{j-1}, t_j]\} + ((j-1)c_I + c_F) \times \Pr \{\text{CM in } (t_{j-1}, t_j]\}] \quad (14)$$

In terms of the deterioration process $X(t)$, the event $\{\text{PM in } (t_{j-1}, t_j]\}$ means PM at the j th inspection and is equivalent to $\{r_0 - X(t_{j-1}) \geq \rho, s \leq r_0 - X(t_j) < \rho\}$. The event $\{\text{CM in } (t_{j-1}, t_j]\}$ means CM during the j th inspection interval, which is equivalent to $\{r_0 - X(t_{j-1}) \geq \rho, r_0 - X(t_j) < s\}$. In a similar way, the expected renewal-cycle length is

$$E(\text{cycle length}) = \sum_{j=1}^{\infty} [t_j \Pr \{\text{PM in } (t_{j-1}, t_j]\} + E(T_F; \text{CM in } (t_{j-1}, t_j])] \quad (15)$$

where T_F is the time to failure. Extensions and variations of the CBM model with failures detected immediately include discounting [39, 40], non-periodic inspection [40, 45], and partial or **imperfect repair** [46].

When a failure is detected only by inspection, the object is renewed when an inspection reveals either that the PM level ρ is crossed while no failure has occurred (PM) or that the failure level s is crossed (CM). The expected cycle cost in equation (14) should now be extended with the cost of the inspection at which failure is detected and a penalization for the unavailability of the object due to the unrevealed failure being

$$\sum_{j=1}^{\infty} [c_I \Pr \{\text{CM in } (t_{j-1}, t_j]\} + c_U E(t_j - T_F; \text{CM in } (t_{j-1}, t_j])] \quad (16)$$

where c_U is the cost of unavailability per unit time. The expected cycle length is

$$E(\text{cycle length}) = \sum_{j=1}^{\infty} t_j \Pr \{\text{PM or CM in } (t_{j-1}, t_j]\} \quad (17)$$

Extensions and variations of the CBM model with failures detected only by inspection include nonperiodic inspection [47, 48], **optimal design** [49], and damage initiation [50]. A CBM model with deterioration described as a compound Poisson process can be found in Zuckerman [44].

Markov Decision Process. When a finite or countable state space is assumed, implying that the condition of a component can be in any one of $n \geq 0$ discrete states, a Markov-chain model can be used (see **Maintenance and Markov Decision Models**). The condition-based inspection model described in the previous section can also be applied when deterioration is modeled by a finite-state Markov process. See [51] for an application to bridge inspections. However, a more common optimization framework for these processes are the so-called Markov decision processes.

Decisions about an optimal policy for maintenance actions are made on a finite set of actions A and costs $C(i, a)$, which are incurred when the process is in state i and action $a \in A$ is taken. The costs are assumed to be bounded and a policy is defined to be any rule for choosing actions. When the process is in state i at time $t = 0, 1, 2, \dots$ and an action a is taken, the process moves into state j after one unit of time with probability

$$P_{ij}(a) = \Pr\{X_{t+1} = j | X_t = i, a_t = a\} \quad (18)$$

Like for regular Markov chains, this transition probability does not depend on the state history. If a stationary policy is selected, then this process is a Markov decision process. A stationary policy arises when the decision for an action only depends on the current state of the process and not on the time at which the action is performed.

Given that the state of the process at times $0, 1, \dots$ is modeled by a Markov decision process $X_0, X_1, \dots$ governed by the transition probabilities $P_{ij}(a)$, the optimization of inspection and/or maintenance policies using this process can be performed. For example, when the object is in state i the expected discounted costs over an unbounded horizon are given by the recurrent relation

$$V_\alpha(i) = C(i, a) + \alpha \sum_{j=1}^n P_{ij}(a) V_\alpha(j) \quad (19)$$

where $\alpha = 1/(1+r)$ is the discount factor for one unit of time, r the discount rate, and V_α the value function using α . Starting from state i , $V_\alpha(i)$ gives us the cost of performing an action a given by $C(i, a)$ and adds the expected discounted costs of moving into another state one unit of time later with probability $P_{ij}(a)$. The discounted costs over an unbounded horizon associated with a start in state j are given by $V_\alpha(j)$, therefore, equation (19) is a recursive equation. The choice for the action a is determined by the maintenance policy and also includes no repair.

A cost-optimal decision can now be found by minimizing equation (19) with respect to the action a . There are a number of ways to find this optimal solution. One of these is the so-called policy improvement algorithm, where equation (19) is calculated for increasingly better policies until no more improvement can be made. Also, it is possible to formulate the minimization problem as a linear programming problem. This is used in the Arizona pavement management system [52] and the PONTIS bridge management system [53]. As Golabi [54] illustrates, we can choose to maximize the condition of the road system under a budget constraint or we can minimize the maintenance cost under a minimum safety constraint. This can be achieved by using the original linear programming formulation or with its dual formulation.

The difference between the use of the condition-based inspection model from the previous section and Markov decision processes lies primarily in the fact that the latter approach is used to make a decision about what to do at fixed times $t = 0, 1, 2, \dots$. The condition-based inspection model has a fixed policy, namely, do nothing if the condition is functional, do a preventative repair if the condition is marginal, and do a corrective repair if the object is in a failed condition. Given this policy, the inspection interval and the thresholds for marginal and failed conditions may be optimized. The decision is therefore not what to do, but when to do it.

Adaptive Maintenance Policies

In most cases, the maintenance optimization approaches are developed under the assumption that the parameters of the process are known (or estimated) within a degree of certainty. Adaptive methods for

maintenance optimization using a Bayes approach where **prior distributions** are formulated for distribution parameters and updated as data becomes available have been proposed for only a few basic models. Wilson and Benmerzouga [55] have used this approach for analysis of group replacement strategies for parallel components with exponential failure times. Mazzuchi and Soyer [56] propose a Bayes approach for the standard age- and block-replacement model where the underlying failure time distribution is Weibull. There are several extensions in the literature. In Markov decision processes, Bayes theorem can be applied to update prior distributions on transition probabilities elicited from engineers with inspections [57]. Finally, Kallen and van Noortwijk [41] provide a Bayes analysis for a periodic (imperfect) inspection and maintenance plan where the deterioration process is modeled by a gamma process.

References

- [1] Dekker, R. & Scarf, P.A. (1998). On the impact of optimisation models in maintenance decision making: the state of the art, *Reliability Engineering and System Safety* **60**(2), 111–119.
- [2] Barlow, R.E. & Proschan, F. (1965). *Mathematical Theory of Reliability*, John Wiley & Sons, New York.
- [3] McCall, J.J. (1965). Maintenance policies for stochastically failing equipment: a survey, *Management Science* **11**, 493–524.
- [4] Pierskalla, W.P. & Voelker, J.A. (1976). A survey of maintenance models: the control and surveillance of deteriorating systems, *Naval Research Logistics Quarterly* **23**, 353–388.
- [5] Sherif, Y.S. & Smith, M.L. (1981). Optimal maintenance models for systems subject to failure - a review, *Naval Research Logistics Quarterly* **28**, 47–74.
- [6] Sherif, Y.S. (1982). Reliability analysis: optimal inspection and maintenance schedules of failing systems, *Microelectronics Reliability* **22**(1), 59–115.
- [7] Valdez-Flores, C. & Feldman, R.M. (1989). A survey of preventive maintenance models for stochastically deteriorating single-unit systems, *Naval Research Logistics* **36**, 419–446.
- [8] Cho, D.I. & Parlar, M. (1991). A survey of maintenance models for multi-unit systems, *European Journal of Operational Research* **51**(1), 1–23.
- [9] Dekker, R. (1996). Applications of maintenance optimization models: a review and analysis, *Reliability Engineering and System Safety* **51**(3), 229–240.
- [10] Pintelon, L.M. & Gelders, L.F. (1992). Maintenance management decision making, *European Journal of Operational Research* **58**(3), 301–317.
- [11] Wang, H. (2002). A survey of maintenance policies of deteriorating systems, *European Journal of Operational Research* **139**(3), 469–489.
- [12] Frangopol, D.M., Kallen, M.J. & van Noortwijk, J.M. (2004). Probabilistic models for life-cycle performance of deteriorating structures: review and future directions, *Progress in Structural Engineering and Materials* **6**(4), 197–212.
- [13] Moubray, J. (1997). *Reliability-Centered Maintenance*, 2nd Edition, Industrial Press, New York.
- [14] Rausand, R. & Høyland, A. (2004). *System Reliability Theory; Models, Statistical Methods, and Applications*, 2nd edition, John Wiley & Sons, Hoboken.
- [15] Wagner, H.M. (1975). *Principles of Operations Research*, 2nd Edition, Prentice-Hall, Englewood Cliffs.
- [16] Nakagawa, T. & Mizutani, S. (2007). A summary of maintenance policies for a finite interval, *Reliability Engineering and System Safety* In Press; doi: 10.1016/j.res.2007.04.004.
- [17] Ross, S.M. (1970). *Applied Probability Models with Optimization Applications*, Dover Publications, New York.
- [18] Samuelson, P.A. (1937). A note on measurement of utility, *The Review of Economic Studies* **4**(2), 155–161.
- [19] Rackwitz, R. (2001). Optimizing systematically renewed structures, *Reliability Engineering and System Safety* **73**(3), 269–279.
- [20] van der Weide, J.A.M., Suyono, & van Noortwijk, J.M. (2008). Renewal theory with exponential and hyperbolic discounting, *Probability in the Engineering and Informational Sciences* **22**(1), (In Press).
- [21] van Noortwijk, J.M. (2003). Explicit formulas for the variance of discounted life-cycle cost, *Reliability Engineering and System Safety* **80**(2), 185–195.
- [22] Fox, B. (1966). Age replacement with discounting, *Operations Research* **14**(3), 533–537.
- [23] van Noortwijk, J.M. & Frangopol, D.M. (2004). Two probabilistic life-cycle maintenance models for deteriorating civil infrastructures, *Probabilistic Engineering Mechanics* **19**(4), 345–359.
- [24] Barlow, R.E. & Hunter, L.C. (1960). Optimum preventive maintenance policies, *Operations Research* **8**(1), 90–100.
- [25] Christer, A.H. (1982). Modeling inspection policies for building maintenance, *Journal of the Operational Research Society* **33**(8), 723–732.
- [26] Baker, R.D. & Christer, A.H. (1994). Review of delay-time OR modelling of engineering aspects of maintenance, *European Journal of Operational Research* **73**(3), 407–422.
- [27] Dekker, R. & Dijkstra, M.C. (1992). Opportunity-based age replacement: exponentially distributed times between opportunities, *Naval Research Logistics* **39**(2), 175–190.
- [28] Coolen-Schrijner, P., Coolen, F.P.A. & Shaw, S.C. (2006). Nonparametric adaptive opportunity-based age replacement strategies, *Journal of the Operational Research Society* **57**(1), 63–81.

- [29] Satow, T. & Osaki, S. (2003). Opportunity-based age replacement with different intensity rates, *Mathematical and Computer Modelling* **38**(11–13), 1419–1426.
- [30] Jhang, J.P. & Sheu, S.H. (1999). Opportunity-based age replacement policy with minimal repair, *Reliability Engineering and System Safety* **64**(3), 339–344.
- [31] Iskandar, B.P. & Sandoh, H. (2000). An extended opportunity-based age replacement policy, *RAIRO Recherche Opérationnelle – Operations Research* **34**(2), 145–154.
- [32] Dekker, R. & Smeitink, E. (1991). Opportunity-based block replacement, *European Journal of Operational Research* **53**(1), 46–63.
- [33] Dekker, R. & Smeitink, E. (1994). Preventive maintenance at opportunities of restricted duration, *Naval Research Logistics* **41**(3), 335–353.
- [34] Thomas, L.C. (1986). A survey of maintenance and replacement models for maintainability and reliability of multi-item systems, *Reliability Engineering* **16**(4), 297–309.
- [35] Nicolai, R.P. & Dekker, R. (2007). Optimal maintenance of multi-component systems: a review, in *Complex System Maintenance Handbook*, K.A.H. Kobbacy & D.N.P. Murthy, eds, Springer Verlag.
- [36] van Noortwijk, J.M. (2007). A survey of the application of gamma processes in maintenance, *Reliability Engineering and System Safety* In Press; doi:10.1016/j.res.2007.03.019.
- [37] Kong, M.B. & Park, K.S. (1997). Optimal replacement of an item subject to cumulative damage under periodic inspections, *Microelectronics Reliability* **37**(3), 467–472.
- [38] Park, K.S. (1988). Optimal continuous-wear limit replacement under periodic inspections, *IEEE Transactions on Reliability* **37**(1), 97–102.
- [39] van Noortwijk, J.M., Kok, M. & Cooke, R.M. (1997). Optimal maintenance decisions for the sea-bed protection of the Eastern-Scheldt barrier, in *Engineering Probabilistic Design and Maintenance for Flood Protection*, R. Cooke, M., Mendel & H., Vrijling, eds, Kluwer Academic Publishers, Dordrecht, pp 25–56.
- [40] Newby, M. & Dagg, R. (2004). Optimal inspection and perfect repair, *IMA Journal of Management Mathematics* **15**(2), 175–192.
- [41] Kallen, M.J. & van Noortwijk, J.M. (2005). Optimal maintenance decisions under imperfect inspection, *Reliability Engineering and System Safety* **90**(2–3), 177–185.
- [42] Abdel-Hameed, M. (1987). Inspection and maintenance policies of devices subject to deterioration, *Advances in Applied Probability* **19**, 917–931.
- [43] Abdel-Hameed, M. (1995). Correction to: “Inspection and maintenance policies of devices subject to deterioration”, *Advances in Applied Probability* **27**, 584.
- [44] Zuckerman, D. (1980). Inspection and replacement policies, *Journal of Applied Probability* **17**, 168–177.
- [45] Grall, A., Bérenguer, C. & Dieulle, L. (2002). A condition-based maintenance policy for stochastically deteriorating systems, *Reliability Engineering and System Safety* **76**(2), 167–180.
- [46] Castanier, B., Bérenguer, C. & Grall, A. (2003). A sequential condition-based repair/replacement policy with non-periodic inspections for a system subject to continuous wear, *Applied Stochastic Models in Business and Industry* **19**(4), 327–347.
- [47] Dieulle, L., Bérenguer, C., Grall, A. & Roussignol, M. (2003). Sequential condition-based maintenance scheduling for a deteriorating system, *European Journal of Operational Research* **150**(2), 451–461.
- [48] Grall, A., Dieulle, L., Bérenguer, C. & Roussignol, M. (2002). Continuous-time predictive-maintenance scheduling for a deteriorating system, *IEEE Transactions on Reliability* **51**(2), 141–150.
- [49] Speijker, L.J.P., van Noortwijk, J.M., Kok, M. & Cooke, R.M. (2000). Optimal maintenance decisions for dikes, *Probability in the Engineering and Informational Sciences* **14**(1), 101–121.
- [50] van Noortwijk, J.M. & Klatte, H.E. (1999). Optimal inspection decisions for the block mats of the Eastern-Scheldt barrier, *Reliability Engineering and System Safety* **65**(3), 203–211.
- [51] Kallen, M.J. & van Noortwijk, J.M. (2006). Optimal periodic inspection of a deterioration process with sequential condition states, *International Journal of Pressure Vessels and Piping* **83**(4), 249–255.
- [52] Golabi, K., Kulkarni, R.B. & Way, G.B. (1982). A state-wise pavement management system, *Interfaces* **12**(6), 5–21.
- [53] Golabi, K. & Shepard, R. (1997). Pontis: a system for maintenance optimization and improvement of US bridge networks, *Interfaces* **27**(1), 71–88.
- [54] Golabi, K. (1983). A Markov decision modeling approach to a multi-objective maintenance problem, in *Essays and Surveys on Multiple Criteria Decision Making*, P. Hansen, ed Springer, Berlin, Vol. 209, pp. 115–125.
- [55] Wilson, J.G. & Benmerzoug, A. (1995). Bayesian group replacement policies, *Operations Research* **43**(3), 471–476.
- [56] Mazzuchi, T.A. & Soyer, R. (1996). A Bayesian perspective on some replacement strategies, *Reliability Engineering and System Safety* **51**(3), 295–303.
- [57] Lee, T.C., Judge, G.G. & Zellner, A. (1970). *Estimating the Parameters of the Markov Probability Model from Aggregate Time Series Data*, North-Holland, Amsterdam.

Related Articles

Age-Dependent Minimal Repair and Maintenance; Block Replacement; Analysis of Recurrent

Events from Repairable Systems; General Minimal Repair Models; Group Maintenance Policies; Imperfect Repair; Inspection Policies for Reliability; Life Cycle Costs and Reliability Engineering; Maintenance and Markov Decision Models; Maintenance Optimization in Random Environments; Markov Processes; Multicomponent Maintenance; Multivariate Age and Multivariate Renewal Replacement; Multivariate Imperfect Repair Models;

Nonparametric Methods for Analysis of Repair Data; Renewal Theory; Repairable Systems Reliability; Replacement Strategies; Stochastic Deterioration; Stress–Strength Model.

THOMAS A. MAZZUCHI, JAN M. VAN
NOORTWIJK AND MAARTEN-JAN KALLEN

Expert Opinion in Reliability

Introduction

Expert opinion can be a useful source of information for assessing reliability, provided it is obtained in a manner that minimizes common sources of **bias** and ambiguity. The term for obtaining expert opinion in this manner is *elicitation*. In reliability assessment, elicitation is most often directed at two distinct elements of expert opinion: expert *judgment* about the value of a specific quantity or probability of interest, or expert *knowledge* about the dependency structure of a system.

Expert opinion elicitation is typically used in one of two ways. First, it is used in situations where data are scarce, but a relevant body of expertise exists. This might be the case, for example, when assessing the reliability of a new type of system where test data are not yet available, but there is experience with similar systems, and the relevant physical mechanisms and design principles are known. Second, expert opinion can be a source of additional information even where a body of data already exists. For example, elicitation of expert knowledge about system dependencies can reveal interactions between parts that have not yet been quantified or formally documented, but are relevant to reliability. Similarly, experts may condition their judgment about part reliability on a wider range of contextual information than can be captured in a particular test. As a result, statistically combining expert judgment with test data may provide more information about reliability than test data alone. In Bayesian statistical analyses, expert judgment elicitation is often used to obtain **prior distributions** which are then combined with other data to generate posterior distributions for use in assessing system reliability [1].

The first step in successful elicitation of expert opinion is the identification of relevant experts. The term *expert* may be applied to a wide range of individuals, not all of whom are appropriate elicitation subjects. In general, for the purposes of expert opinion elicitation, the relevant experts are those with *contributory expertise* in the problem area [2]. These are individuals who have experience either with the specific system in question, or with similar

systems; who have experience directly related to the problem at hand; and who are (or were recently) active participants in technical work. Expertise in this sense may or may not be related to formal technical qualifications (such as an academic degree), although it is usually related to experience over some period of time in the relevant technical area. Experts are also usually chosen based on recognition of their expertise within the relevant technical community, as this can be crucial to the credibility of the elicitation results.

When quantitative results are desired, elicitation of expert opinion is usually carried out according to a formalized list of questions, which the experts may be asked to answer verbally or in writing. Formulating a useful set of questions is an important aspect of elicitation methodology. Overly complex or ambiguous questions, or questions that require the expert to make multiple assumptions, have been shown to yield less accurate responses. For this reason, general questions are often decomposed into a series of simpler questions that depend on only a few specified assumptions. To develop such a list of questions, the persons conducting the elicitation often work in collaboration with one or more *advisory experts*. These experts assist in defining the problem space and refining the elicitation questionnaire, providing a check on the quality of the questions before they are administered to a larger set of experts. Because many technical experts do not have expertise in statistics or reliability, one role of the advisory expert is to translate reliability terms into technical terms familiar to the expert community. The quality of the questions can also be assessed during the elicitation process by monitoring the experts' problem-solving processes, for example by asking them to explicitly report on the reasoning and assumptions behind their responses [3].

Expert opinion can be aggregated in a variety of ways. Elicitation can be conducted with each expert individually, and their responses aggregated mathematically [3, 4]. Expert opinions can also be aggregated behaviorally by conducting group elicitations or through the use of other methods in which experts try to reach consensus, such as the Delphi method [3, 5].

Expert Judgment Elicitation

Expert judgment elicitation is typically focused on obtaining the best judgment of an expert about the

2 Expert Opinion in Reliability

value of a quantity or probability of interest [6]. For example, an expert might be asked to estimate the likely failure temperature of a part, or could be asked to estimate the probability that a part will fail at a given temperature. The expert is most often asked to provide a point estimate, but may also be asked to estimate uncertainty bounds or provide a complete probability distribution.

The main driver for elicitation methodology is the need to identify and mitigate bias in experts' responses. While some biases are problem specific, research in the social and behavioral sciences has documented a number of common biases in the way both experts and laypeople answer elicitation questions. It is useful to distinguish three forms of bias in elicitation. *Cognitive bias* occurs when expert reasoning fails to adhere to some objective standard, such as the laws of probability, due to inherent limitations on the brain's ability to process information. *Motivational bias* occurs when experts alter their responses (often unconsciously) in response to the elicitation setting, so that their responses do not reflect their true subjective beliefs. For example, experts may downplay uncertainties in order to provide more authoritative answers if they believe that is what the interviewer is looking for. *Interviewer bias* may also be introduced into the elicitation process if the person conducting the elicitation or analyzing the results misinterprets expert responses or steers them toward a preferred answer [3].

Because motivational bias and interviewer bias can only be measured in relation to subjective standards, methods for mitigating these forms of bias are not highly formalized; a general principle is to ask nondirective questions that carefully avoid leading the expert toward a particular answer. Cognitive bias, however, has been the subject of extensive behavioral research, and a number of specific techniques have been developed to minimize it.

Some of the most significant types of cognitive bias are availability, anchoring, and inconsistency [4, 7–9]. *Availability bias* refers to experts focusing on what they have seen prior to the exclusion of the full range of what is possible. So an expert may report an inappropriately high probability of a part failing in a way previously observed, especially if it was a dramatic failure, while discounting the probability of more likely, more mundane failures. Group elicitation methods can help counter this bias by allowing experts to share information. Free

association methods, in which experts are asked to quickly generate a list of factors bearing on a question, are also used [3].

Anchoring bias is the tendency, when faced with new information, to update quantities conservatively from an initial estimated value, rather than comprehensively reassessing the problem. This can occur when experts are asked to estimate the probability of an event based on various preconditions; their answer to the first question is likely to influence their answers to subsequent questions. Similarly, if an expert is asked to estimate probabilities for a string of highly unlikely events, they may then have difficulty giving a good estimate of a more likely event. Providing feedback from other experts, such as through group elicitation methods or the Delphi method, can inhibit anchoring. Elicitation questions can also explicitly request that experts not base their answers on previous responses [3].

The inconsistency bias appears in complex elicitations where experts may not be able to keep track of all previous answers, assumed rules, and hypothetical frameworks in a way that allows them to provide consistent estimates. This bias is related to memory, so one way it can be addressed is by frequently reminding elicitation subjects of key definitions and assumptions during the elicitation process. Experts may also be provided with documentation of previous responses so they can monitor their own consistency [3]. Pairwise comparison methods such as the Saaty Analytical Hierarchy Process, can also be used to monitor consistency [10].

Elicitation Methods

Elicitation methods can be categorized by what part of the probability distribution they are designed to capture: center value, interval, or functional form. Examples of eliciting a center value include asking an expert for a single-value estimate of a physical quantity or probability of an event occurring. This can be done on a continuous or discrete scale. Asking the expert to give a probability between 0 and 1 or a temperature in degrees Celsius are examples of using a continuous scale. A discrete scale is one that breaks the range of possible values into a series of intervals. For example, the Sherman Kent scale breaks probabilities into five categories: nearly impossible, improbable, even chance, probable, and nearly certain, which are specified to correspond

to percent chances of approximately 10, 30, 50, 70, and 90%. Experts may be more comfortable providing probability estimates as percent chances or, particularly for small probabilities, as odds ratios (e.g., 1/100) [3]. An indirect method for eliciting center values, known as the *interval technique*, is to specify the range of possible values of a variable, and then ask experts to adjust a dividing point until the chances of the value being above or below that point are approximately equal, which will be the likely median value [7, 11]. Availability and anchoring are potential biases in eliciting center values.

Interval elicitation is generally used to elicit uncertainty in experts' estimates. This can be done in two ways. First, a probability interval can be specified and the expert can be asked which values fall within the interval. For example, in addition to the most likely temperature at which a part might fail, the interviewer might ask for upper and lower bounds on failure temperature. The interval may be specified in terms of absolute maximum and minimum values, or in terms of **confidence intervals** (e.g., specifying a 90% chance of failure within the interval). A second approach is to specify a range of values and ask the expert what the chances of failure are within that range [7, 11]. For example, if a part specification provides a nominal range of temperature tolerance, this could be verified by asking the expert what he or she thinks the odds are of the part failing if its temperature stays within that range [12]. Experts frequently underestimate uncertainty, a form of anchoring bias [3, 13]. Elicitation methods designed to counter availability and anchoring may be used to counteract this tendency.

Elicitation methods have also been developed to capture the overall shape of a probability distribution. Statistically, sophisticated experts can be asked to draw the **probability density function** or cumulative distribution function directly. Distributions can be approximated by aggregating point estimates – for example, the probability distribution on failure temperature for a part can be assembled by asking for the odds of failure at various points in the temperature range. Similarly, fractile methods ask the expert to estimate the value of a parameter that has a certain probability of occurring (e.g., 0.25) at several points between 0 and 1, and then aggregate these to approximate a probability distribution. The interval technique, mentioned above, may be extended to estimate **quartiles** or even smaller intervals. Methods

for eliciting the shape of probability functions are attractive from a statistical point of view because they provide the most information possible for further analysis. However, they are complex and time consuming, which may lead to problems with inconsistency, and ask for a level of discrimination about probabilities that many experts may not possess.

Elicitation of System Dependencies

Often what an expert knows in the context of reliability goes beyond what they can estimate about a single probability or range of values. Elicitation of expert knowledge about system dependencies focuses on determining the structure of variables that make up a reliability model, and the degree to which system outcomes can be explained in terms of combinations of component measurements.

Eliciting system knowledge is a challenge for a variety of reasons. Systems are often complex: frequently no single person truly understands how the entire system works and what all measures gathered about the system mean. In addition, for all but the most trivial systems, experts typically have incomplete or conflicting knowledge about these factors. Generally, system reliability depends on part characteristics and interactions that may not be documented in formal system representations.

If a system is made up of several components, it is generally necessary to elicit what experts know about dependencies within the system (i.e., model structure) in order to accurately assess reliability. Take the trivial case of a two-component system with the probability of each component functioning properly as $P1$ and $P2$. Unless it is known whether or not these components are in a serial relationship, the system reliability cannot be estimated as $P1 \times P2$ – under different system logical structures, with different assumptions about system control logic, the overall system reliability could alternatively be maximum ($P1, P2$), or $P1 + (1 - P1) \times P2$, etc.

In addition, once the assumption that the reliability of $P1$ and $P2$ are independent is violated, a host of additional unavoidable elicitation complications are introduced. For example, $P1$ and $P2$ may have more than one failure mode, or system components may have more than two states. In such a case, it may be necessary to model not only whether a component fails, but how the failure may cascade

along different pathways and with different overall system level effects through to other components. *P1* may fail gracefully, allowing *P2* to pick up the slack, or *P1* may explode, greatly increasing the odds that *P2* will also fail.

Accordingly, unless a set of fairly stringent assumptions about a system's internal logic can be sustained (i.e., serial relationships, binary states, stable and fixed relationships, and no changes with age), eliciting expert knowledge of dependencies is not a trivial task [14]. In the case of some systems, these assumptions are not impossibly stringent. Electrical circuits and some types of electronics, for example, lend themselves to models that incorporate these assumptions. Modeling a great many other types of systems, especially those involving human interactions or extended life cycles, greatly stresses all of these assumptions. Military systems and electrical power facilities are examples of systems that face varying conditions over their lifetimes.

These challenges divide the elicitation of expert knowledge about system dependencies into two discernable tasks: eliciting knowledge about system structure directly from experts, and harvesting information that experts and their communities have encoded in representations of the system (e.g., engineering block diagrams, event timelines, flowcharts, etc.).

Eliciting system structure depends on choosing a system representation method that will support a successful description of the system by (a) successfully capturing the complexities of component relationships, while (b) presenting information that is meaningful both to system operators (generally engineers) and reliability modelers (generally statisticians). If the system is fairly simple, it may be possible to sketch and annotate a representation of important components and relationships on a piece of paper in a single conversation. If the system is more complicated, elicitation may involve multiple experts with focused knowledge of different parts of the system. In such cases, a more formalized approach is helpful. For example, software has been developed that enables representation of the parts of a system and their interactions as a collection of boxes and connecting lines, and also keeps track of those parts and interactions in a database. A more complicated system may also require more time iterating representations with the experts over several conversations

as they think through and come to consensus on the functional dependencies [15].

There are a number of possible approaches to developing representations of system dependencies. A typical elicitation process might run as follows:

First, elicit a list (or ontology) of possible objects that will appear in the system representations. This ontology dictates what relationships are permissible and how objects are named. Among other things, this serves to prevent confusion about similar parts – for example, whether “Battery A” on diagram 4 is the same as “battery-a” on diagram 64. Next, create a high-level diagram that only shows the main components, and indicates how these components interact. Then “drill down” into each of those high-level boxes and create a diagram of the parts that comprise that component, also noting interactions. Parts may appear on more than one diagram (e.g., a battery may provide electricity to more than one component).

Next, refine the representation by revisiting the interactions between the components/parts and formalizing what one box provides to the next. Additionally, note failure modes for each component/part and how they affect the interactions. Iterating through this process, experts may change their minds about whether a particular concept is a part or a relationship, and they may remember new parts as they try to explain the functionality of the system.

An alternate approach is to first diagram the timeline of the key events in a system's operation (e.g., fill with water, add detergent, agitate, drain, rinse, drain, spin), then diagram system functions that must occur to make these events happen, and then finally identify components/parts that have to work in order for the functions to transpire [16].

Either of these approaches generates a model structure to work from: a **fault tree** structure if component/part logic is binary or a Bayes network structure if the logic is multistate. Having such a diagrammatic representation of system components/parts and relationships can be a valuable starting point for eliciting point values and distributions.

Seldom do systems engineers have a functional representation of a system in hand and ready for the reliability modeler to use. However, they often have similar representations: process representations like block flow diagrams, IDEF0, or flowcharts; outcome system behavior summaries like event or fault trees; or time logic summaries like timelines or GANT and PERT charts. All of these contain valuable knowledge

about system structure and dependencies that has already been encoded by the expert community. A study of these documents before an elicitation, or a discussion of these documents during an elicitation, can accelerate the process of building an ontology, functionally grouping parts, and understanding relationships and failure modes.

The purpose of eliciting expert judgment and expert knowledge of system dependencies is to provide information about likely system performance and to translate system dependencies into some kind of probabilistic fault structure. As part of any reliability modeling effort, many types of data may need to be stitched together to produce a reliability estimate (e.g., full-system or component test data, computer models, surrogate data, expert opinion, etc.) [17]. Expert opinion is most useful and defensible when it is collected according to scrupulous and documented methods that take into account potential biases and how the experts understand the problem.

References

- [1] Lindley, D.V. & Singpurwalla, N.D. (1986). Reliability (and fault tree) analysis using expert opinions, *Journal of the American Statistical Association* **81**, 87–90.
- [2] Collins, H.M. & Evans, R. (2002). The third wave of science studies: studies of expertise and experience, *Social Studies of Science* **32**, 235–296.
- [3] Meyer, M. & Booker, J. (2001). *Eliciting and Analyzing Expert Judgment: A Practical Guide*, ASA-SIAM, Philadelphia.
- [4] Cooke, R. (1991). *Experts in Uncertainty: Opinion and Subjective Probability in Science*, Oxford University Press, Oxford.
- [5] Helmer, O. (1968). Analysis of the future: the Delphi method. in *Technological Forecasting for Industry and Government*, J. Bright, ed, Prentice-Hall, Englewood.
- [6] O'Hagan, A., Buck, C.E., Daneshkhah, A., Eiser, J.R., Garthwaite, P.H., Jenkinson, D.J., Oakley, J.E. & Rakow, T. (2006). *Uncertain Judgements: Eliciting Experts' Probabilities*, John Wiley & Sons, Chichester and Hoboken.
- [7] Wilson, A. (2004). *Cognitive Factors Affecting Subjective Probability Assessment*, Duke University, ISDS Discussion Paper 94-02.
- [8] Booker, J. & McNamara, L. (2005). Expert knowledge in reliability characterization: a rigorous approach to eliciting, documenting, and analyzing expert knowledge, in *Engineering Design Reliability Handbook*, E. Nikolaidis, D.M. Ghiocel, S. Singhal & N. Nikolaidis, eds, CRC Press, New York.
- [9] Tversky, A. & Kahneman, D. (1974). Judgment under uncertainty: heuristics and biases, *Science* **185**, 1124–1131.
- [10] Saaty, T.L. (1980). *The Analytical Hierarchy Process: Planning, Priority Setting, and Resource Allocation*, McGraw-Hill, New York.
- [11] Spetzler, C.S. & von Holstein, S. (1975). Probability encoding in decision analysis, *Management Science* **22**, 340–358.
- [12] Morgan, G. (1990). *Uncertainty: A Guide to Dealing with Uncertainty in Quantitative Risk and Policy Analysis*, Cambridge University Press, Cambridge.
- [13] Alpert, M. & Raiffa, H. (1982). A progress report on the training of probability assessors, in *Judgment Under Uncertainty: Heuristics and Biases*, D. Kahneman, P. Slovic & A. Tversky, eds, Cambridge University Press, Cambridge, pp. 294–305.
- [14] Bedford, T., Quigley, J. & Walls, L. (2006). Expert elicitation for reliable system design, *Statistical Science* **36**, 428–462.
- [15] Klamann, R. & Koehler, A. (2005). Large-scale qualitative modeling for system prediction, *North American Association for Computational Social and Organizational Systems, Conference Proceedings*, South Bend.
- [16] Wilson, A., McNamara, L. & Wilson, G. (2007). Information integration for complex systems, *Reliability Engineering and System Safety* **92**, 121–130.
- [17] Anderson-Cook, C.M., Graves, T., Hengartner, N., Klamann, R., Koehler, A., Wilson, A.G., Anderson, G. & Lopez, G.S. (Forthcoming). Reliability modeling using both system test and quality assurance data. *Journal of the Military Operations Research Society*.

Related Articles

Subjective Life Models; Fault Trees; Bayesian Networks in Reliability; Fault Tree Analysis for Large Systems.

BENJAMIN H. SIMS, ANDREW KOEHLER
WIEDLEA AND GREGORY D. WILSON

Multivariate Age and Multivariate Renewal Replacement

Introduction

For the past four decades, many researchers have shown interest in the study of **maintenance policies** for systems with stochastic failure. One major reason for this is that these policies can be applied to many areas such as the military, industry, health, and environment.

Generally speaking, in a repairable system, components are repaired or replaced upon failure. There will be cases where the interruption caused in service failure of the system is more costly than the replacement of the components which are in service. Thus, in order to avoid the high cost associated with a failure of the system, an alternative approach is not to wait for components to fail, but to replace or repair them in accordance with a predetermined maintenance schedule. Such preventive maintenance policies reduce unexpected incidents of component or system failure. A large part of reliability theory deals with replacement of single-component systems. Two of the popular **replacement policies** are (a) age- and (b) **block-replacement** policies. Under age replacement a system is replaced only if it has either failed or attained an age, say T . When systems can be replaced in blocks, an example being the replacement of street lamps, a block-replacement policy is used. Here the system is replaced only if either it has failed or the current time is an integer multiple of the block-replacement interval, say T . Many papers have dealt with a stochastic comparison of the number of failures and replacements under these policies. Several authors have also studied different optimal policies under age and block replacement policies. We refer you to Barlow and Proschan [1], Lam [2], Valdez-Flores and Feldman [3] and Shaked and Zhu [4].

As systems become more complicated and require new technologies and methodologies, more sophisticated maintenance policies are needed to solve the maintenance problems. More specifically, it is quite common to find systems consisting of k components which are related. In such cases a legitimate question is: when to replace the components? Obviously,

the classical age-and block-replacement policies are unsuitable for such situations, since they only deal with a single component or system. Here, our focus is strictly on the age replacement policy. Also, our survey is confined to multivariate versions of age replacement policies as well as optimal age replacement policies.

Throughout this article “increasing” means “non-decreasing” and “decreasing” means “nonincreasing”.

Multivariate Age and Multivariate Renewal Replacements

Consider a system consisting of k components that are related. It is assumed that there are many identical components in stock to every component of this system. Below we describe the following two replacement policies, which are extensions of age replacement and ordinary renewal policies to the multivariate case. For more details see Ebrahimi [5].

Definition 1

1. Under a multivariate age replacement (MAR) policy, the component $i, i = 1, \dots, k$, of a system is replaced either at age T_i , or upon its failure. We refer to this as the MAR at $(T_1, \dots, T_k)$;
2. Under a multivariate renewal replacement (MRR), the component $i, i = 1, 2, \dots, k$, is replaced upon its failure.

To avoid complexity, we confine our attention to the case where $k = 2$. We define the following counting processes:

$N_i(t)$ = the number of component i failures in $(0, t]$ under MRR, $i = 1, 2$; $N_A(t; i, T) =$ the number of component i failures in $(0, t]$ under a MAR $(T, T), i = 1, 2$.

We let X_1 and X_2 be two non-negative continuous random variables representing the times to failure of new components 1 and 2, respectively. We also let $\bar{F}(x_1, x_2) = P(X_1 > x_1, X_2 > x_2)$ be the joint survival function of X_1 and X_2 and $\bar{F}_{X_i}(x_i) = P(X_i > x_i)$ be the marginal survival function of $X_i, i = 1, 2$. Here, $\bar{F}(0, 0) = \bar{F}_{X_1}(0) = \bar{F}_{X_2}(0) = 1$ and $\bar{F}(\infty, \infty) = \bar{F}_{X_1}(\infty) = \bar{F}_{X_2}(\infty) = 0$. It is assumed that as soon as a component of the system fails or it reaches the age T it will be replaced by a new and identical component whose lifetime is independent of the replaced component but is

2 Multivariate Age and Multivariate Renewal Replacement

dependent on the lifetimes of the components that are currently in service. It is also assumed that the replacement time (the time it takes to replace either component) is negligible. We denote the conditional survival function of X_1 given $X_2 > y$ and the conditional survival function of X_2 given $X_1 > y$ by

$$\bar{H}_y^{(1)}(x) = P(X_1 > x | X_2 > y) \quad (1)$$

and

$$\bar{H}_y^{(2)}(x) = P(X_2 > x | X_1 > y) \quad (2)$$

respectively.

We let $W_i(j)$ and $Y_i(j)$ be the intervals between the $(j-1)$ th and j th arrivals of the processes $N_i(t)$ and $N_A(t; i, T)$, $i = 1, 2$ respectively. Now, one can show that under MRR

$$\begin{aligned} P(W_i(j) \geq w_i, i = 1, 2 | W_i(\ell) \\ = a_{i\ell}, \ell = 1, \dots, j-1, \\ i = 1, 2) = \bar{F}(w_1, w_2) \end{aligned} \quad (3)$$

and under MAR(T, T) for $\ell_i T \leq y_i < (\ell_i + 1)T$, $i = 1, 2$

$$\begin{aligned} P(Y_i(j) \geq y_i, i = 1, 2 | Y_i(\ell) = a_{i\ell}, \\ \ell = 1, \dots, j-1, i = 1, 2) \\ = (\bar{F}(T, T))^{\ell_1} (\bar{H}_{y_1 - \ell_1 T}^{(2)}(T))^{\ell_2 - \ell_1} \\ \times \bar{F}(y_1 - \ell_1 T, y_2 - \ell_2 T) \\ \text{if } \ell_1 \leq \ell_2 \text{ and} \\ = (\bar{F}(T, T))^{\ell_2} (\bar{H}_{y_2 - \ell_2 T}^{(1)}(T))^{\ell_1 - \ell_2} \\ \times \bar{F}(y_1 - \ell_1 T, y_2 - \ell_2 T) \end{aligned} \quad (4)$$

if $\ell_1 \geq \ell_2$, see Ebrahimi [5] for more details. It should be noted that the equation (4) can be generalized to the case $k > 2$ by defining conditional survival functions similar to equations (2) and (3) for all permutations of $X_1, X_2, \dots, X_k$.

Contrast between MAR and MRR

It is clear that the multivariate age planned replacements take into account the ages of the components. It is therefore of interest to compare it with MRR with respect to the number of failures.

The following result provides a stochastic comparison of failures experienced under MAR(T, T) and MRR.

Theorem 1 *If (a) $\bar{F}(x + t_1, x + t_2) \leq (\geq) \bar{F}(x, x)$ $\bar{F}(t_1, t_2)$ for all $x, t_1, t_2 \geq 0$, and $\bar{H}_y^{(i)}(t_1 + t_2) \leq (\geq) \bar{H}_y^{(i)}(t_1) \bar{H}_y^{(i)}(t_2)$, $i = 1, 2$, for all $y, t_1, t_2 \geq 0$, then $P(N_i(t_i) \leq n_i, i = 1, 2) \leq (\geq) P(N_A(t_i; 1, T) \leq n_i, i = 1, 2)$ for all $n_1, n_2, t_1, t_2 \geq 0$.*

Theorem 1 states that under the conditions (a($\leq$)) and (b($\leq$)) imposed on the joint survival function of X_1 and X_2 the number of failures under MAR is stochastically less than the number of failures under MRR. Next result studies the effect of varying the replacement interval under MAR.

Theorem 2 *If (a) $\bar{F}(x_1 + t_1, x_2 + t_2) \leq \bar{F}(x_1, x_2)$ $\bar{F}(t_1, t_2)$ for all $x_1, x_2, t_1, t_2 \geq 0$ and (b) $\bar{H}_y^{(i)}(t_1 + t_2) \leq \bar{H}_y^{(i)}(t_1) \bar{H}_y^{(i)}(t_2)$, $i = 1, 2$, for all $t_1, t_2, y \geq 0$, then for any $k \geq 1$, $P(N_A(t_i; i, T) \leq n_i, i = 1, 2) \geq P(N_A(t_i; i, kT) \leq n_i, i = 1, 2)$, for all $t_i, n_i, i = 1, 2$.*

Finally, our last result gives us information regarding the weakest component (weakest in the sense that it is more likely to fail).

Theorem 3 *If X_1 is stochastically larger than X_2 (X_1 is said to be stochastically larger than X_2 if $\bar{F}_{X_1}(x) > \bar{F}_{X_2}(x)$ for all $x \geq 0$, see Shaked and Shanthikumar [6]), then (a) $P(N_1(t) \leq n) \geq P(N_2(t) \leq n)$ for all $n, t \geq 0$ and (b) $P(N_A(t; 1, T) \leq n) \geq P(N_A(t; 2, T) \leq n)$, for all $n, t \geq 0$.*

Intuitively speaking, Theorem 3 says that if component 1 has a larger survival function than component 2, then the number of times that component 1 fails is stochastically less than the number of times that component 2 fails under both MAR(T, T) and MRR.

Optimal Age Replacement under MAR(T, T)

In this section we determine an optimal replacement policy T^* under MAR(T, T) such that

$$C(T) = \frac{\text{the expected cost incurred in a cycle}}{\text{the expected length of a cycle}} \quad (5)$$

is minimized. Here a cycle is the time between two consecutive system failures. We note that $C(T)$

is being used as the criterion for evaluating the replacement policies.

Now, consider the following two cases.

Case I Series system (the system works if both components work). Let c_1 be the constant cost of replacement of one or both components when the system fails and let c_2 be the constant cost of a preventive replacement of both components at age T with $0 < c_2 < c_1 < \infty$. Then, equation (5) reduces to

$$C(T) = \frac{1}{\int_0^T \bar{F}(u, u) du} (c_1 - \beta \bar{F}(T, T)) \quad (6)$$

where $\beta = c_1 - c_2 > 0$. We can determine the optimal replacement policy T^* by minimizing the equation (6). The minimization procedure can be done by analytical or numerical methods. See Ebrahimi [5] for more details. In fact, one can show that T^* satisfies the following condition.

$$\frac{g(T)}{\bar{F}(T, T)} \int_0^T \bar{F}(u, u) du + \bar{F}(T, T) = \frac{c_1}{\beta} \quad (7)$$

where $g(T) = -\frac{d}{dt} \bar{F}(t, t)|_{t=T}$.

Case II Parallel System (the system works if at least one component works). Under this case equation (5) reduces to

$$C(T) = \frac{1}{\sum_{i=1}^2 \frac{\int_0^T \bar{F}_{X_i}(u) du}{1 - \bar{F}_{X_i}(T)} - \frac{\int_0^T \bar{F}(u, u) du}{1 - \bar{F}(T, T)}} \times \left(c_1 + c_2 \sum_{i=1}^2 \frac{\bar{F}_{X_i}(T)}{1 - \bar{F}_{X_i}(T)} - \frac{c_2 \bar{F}(T, T)}{1 - \bar{F}(T, T)} \right) \quad (8)$$

We may proceed as in Case 1 and determine the optimal policy T^* by minimizing $C(T)$ in the equation (8).

Applications

This section illustrates applications of our results that are obtained in the previous sections.

Bivariate Gumble Model

Consider X_1 and X_2 with the bivariate Gumble distribution

$$\bar{F}(x_1, x_2) = \exp[-\lambda_1 x_1 - \lambda_2 x_2 - \delta x_1 x_2], \\ x_1, x_2 \geq 0, \lambda_1, \lambda_2, \delta \geq 0 \quad (9)$$

See Gumble [7]. It is clear that $\bar{F}(x_1 + t_1, x_2 + t_2) \leq \bar{F}(x_1, x_2) \bar{F}(t_1, t_2)$. Also, $\bar{H}_y^{(1)}(x) = P(X_1 > x | X_2 > y) = \exp[-\lambda_1 x - \delta x y]$, $\bar{H}_y^{(2)}(x) = P(X_2 > x | X_1 > y) = \exp[-\lambda_2 y - \delta x y]$ and $\bar{H}_y^{(i)}(t_1 + t_2) \leq \bar{H}_y^{(i)}(t_1) \bar{H}_y^{(i)}(t_2)$, $i = 1, 2$, for $t_1, t_2 \geq 0$. Therefore, all the results in Theorems 1 and 2 are true for this family of distributions.

Marshall–Olkin Model

Consider X_1 and X_2 with the joint survival function of Marshall–Olkin bivariate exponential which is given by

$$\bar{F}(x_1, x_2) = \exp(-\lambda_1 x_1 - \lambda_2 x_2 - \delta \max(x_1, x_2)), \\ x_1, x_2 \geq 0, \lambda_1, \lambda_2 > 0 \text{ and } \delta \geq 0 \quad (10)$$

See Marshall and Olkin [8]. It is clear that

$$\bar{F}(x + t_1, x + t_2) \geq \bar{F}(x, x) \bar{F}(t_1, t_2) \quad (11)$$

Also, $\bar{H}_y^{(1)}(x) = \exp(-\lambda_1 x - \delta \max(x, y))$, $\bar{H}_y^{(2)}(x) = \exp(-\lambda_2 x - \delta \max(x, y))$ and $\bar{H}_y^{(i)}(t_1, t_2) \geq \bar{H}_y^{(i)}(t_1) \bar{H}_y^{(i)}(t_2)$, $i = 1, 2$, for $t_1, t_2 \geq 0$. Therefore all the results in Theorem 1 hold.

T^* for Bivariate Gumble and Marshall and Olkin Models

As an application of equation (6), consider the bivariate Gumble family of distributions which was described in the section titled “Bivariate Gumble Model”. For this case, equation (7) reduces to

$$\exp[-(\lambda_1 + \lambda_2)T - \delta T^2] + (\lambda_1 + \lambda_2 + 2\delta T) \\ \int_0^T \exp(-(\lambda_1 + \lambda_2) \\ \times u - \delta u^2) du = \frac{c_1}{\beta} \quad (12)$$

4 Multivariate Age and Multivariate Renewal Replacement

Now, suppose $c_1 = 2$, $c_2 = 1$ and $\lambda_1 = \lambda_2 = \delta = 1$. Then, we get T^* as approximately 0.27. Since $X_i, i = 1, 2$, has an exponential distribution with mean 1, it is clear that under the univariate age replacement policy $T^* = \infty$. Also under the Marshall and Olkin model for series system $T^* = \infty$. Intuitively speaking, this simply means that age replacement policy is not applicable in univariate case of exponential distribution and under the Marshall and Olkin model.

Concluding Remarks

In this article, we described several maintenance policies for multicomponent systems. Our policies are suitable for revealed failures of components, that is, failures that are detected as soon as they occur. There are also examples of units whose failures are not revealed unless some type of inspection is performed. Berrade [9] proposed a policy for unrevealed failures of components. While the results obtained in section titled "Contrast between MAR and MRR" are appropriate for the case in which components are replaced at the same time, $T_1 = T_2 = T$, it will be interesting to see if similar results are obtained when components are replaced at different times, $T_1 \neq T_2$.

References

- [1] Barlow, R. & Proschan, F. (1981). *Statistical Theory of Reliability and Life Testing*, McArdle Press, Silver Spring.
- [2] Lam, Y. (1988). A note on the optimal replacement problem, *Advances in Applied Probability* **20**, 479–482.
- [3] Valdez-Flores, C. & Feldman, R.M. (1989). A survey of preventive maintenance models for stochastically deteriorating single-unit systems, *Naval Research Logistics* **36**, 419–446.
- [4] Shaked, M. & Zhu, H. (1992). Some results on block replacement policies and renewal theory, *Journal of Applied Probability* **29**, 932–946.
- [5] Ebrahimi, N. (1997). Multivariate age replacement, *Journal of Applied Probability* **34**, 1032–1040.
- [6] Shaked, M. & Shanthikumar, J.G. (1994). *Stochastic Orders and Their Applications*, Academic Press, San Diego.
- [7] Gumble, E.J. (1961). Multivariate exponential distributions, *Bulletin of the International Statistical Institute* **39**, 469–475.
- [8] Marshall, A.W. & Olkin, I. (1967). A multivariate exponential distribution, *Journal of the American Statistical Association* **62**, 30–44.
- [9] Berrade, M.D. (1999). Statistical techniques in reliability: aging properties under mixtures and optimal maintenance policies, Ph.D. thesis, University of Zaragoza, Spain.

Related Articles

Age-Dependent Minimal Repair and Maintenance; Block Replacement; Burn-In and Maintenance Policies; Dynamic Programming Methods in Repair and Replacement; General Minimal Repair Models; Group Maintenance Policies; Maintenance and Markov Decision Models; Maintenance Optimization; Maintenance Optimization in Random Environments; Markov Renewal Processes in Reliability Modeling; Multicomponent Maintenance; Multivariate Imperfect Repair Models; Multivariate Stochastic Orders and Aging; Replacement Strategies; Stationary Replacement Strategies; Total Productive Maintenance.

NADER EBRAHIMI

Nonparametric Methods for Analysis of Repair Data

Introduction

Engineering systems and human systems are subject to failure. Rather than replacing the failed system by a new system, repair is implemented because substituting a new item for the failed item is impractical, not feasible, or too expensive. Thus we are led to the situation where a system, undergoing continued maintenance by having some type of repair applied after each failure, yields a data set consisting of the failure times and the type of repair administered at each failure.

A repair model is a stochastic formulation of the joint distribution of the failure times, or equivalently, the joint distribution of the interfailure times. We let $\{S_j\}$ denote the failure times of the system and $\{T_j\}$, the interfailure times, where $T_j = S_j - S_{j-1}$ and $S_0 \stackrel{\text{def}}{=} 0$. Let F be the distribution of the time to first failure of the system. It is of interest to estimate F (or equivalently the survival function $\bar{F} = 1 - F$) nonparametrically from the failure times and knowledge of the type of repairs made after each failure. Other parameters, such as the expected values of the interfailure times and other features of the repair process, can be expressed in terms of F . **Estimation** of F is not straightforward because data that arise from repair processes will typically have dependencies that are induced, for example, by the nature of the repair, other environmental factors, and the specification of the period of observation.

In the section titled “Repair Models” we discuss several of the prominent repair models in use. In the section titled “A General Repair Model” we address nonparametric estimation of F and describe a nonparametric asymptotic simultaneous confidence band for F . The section titled “Additional Nonparametric Methods” considers related nonparametric estimation and testing procedures for **repair data**.

Repair Models

In the *perfect repair model*, a failed system is replaced by a new one having the same stochastic properties

as the original. That is, $\{T_j\}$ are independent and identically distributed (i.i.d.) according to F . The imperfect repair model known as the **minimal repair model**, due to Ascher [1] (see **Imperfect Repair**), specifies that the repaired system has the same distribution as an operational new system with the age of the recently failed system. In symbols, the survival distribution function of the minimally repaired system is given by $\bar{F}_t$, where

$$\bar{F}_t(s) = \frac{1 - F(s + t)}{1 - F(t)}, \quad s > 0 \quad (1)$$

and t denotes the time of the previous failure. The Brown and Proschan (BP) [2] repair model is a generalization of the minimal repair model. The BP model is an imperfect repair model where at the time of each repair, a perfect repair is performed with probability p and a minimal repair is performed with probability $1 - p$. The Block, Borges, and Savits (BBS) [3] model generalizes the BP model by allowing the probability of a perfect repair to depend on the age of the failed system. Kijima [4] introduced two models, models I and II, that allow for repairs to be better than minimal but not as good as perfect repairs. In the Kijima models, the effective age to which the system is restored after repair depends on the effective age just before failure and the degree of repair random variables. Kijima’s models I and II are special cases of a repair model considered by Uematsu and Nishida [5]. Dorado, Hollander, and Sethuraman (DHS) [6] defined a general repair that contains many popular repair models and introduces many others. Because of its generality, we focus on it in the section titled “A General Repair Model”.

A General Repair Model

The DHS [6] model of imperfect repair is based on the family of survival distributions

$$\bar{F}_{a,\theta}(x) \stackrel{\text{def}}{=} \frac{\bar{F}(\theta x + a)}{\bar{F}(a)} \quad (2)$$

where θ is a positive parameter. This family enjoys a stochastic ordering property, namely, $\theta < \theta'$ implies $F_{a,\theta} \stackrel{\text{st}}{\geq} F_{a,\theta'}$ for each a . Thus lower values of θ imply longer remaining life. The DHS model is based on two sequences $\{A_j\}$, $\{\theta_j\}$ where the A ’s are

2 Nonparametric Methods for Analysis of Repair Data

called *effective ages* and the θ 's are known as *life supplements*. The A 's and θ 's are assumed to satisfy

$$\begin{aligned} A_1 &= 0, \theta_1 = 1, A_j \geq 0, \theta_j \in (0, 1] \text{ and} \\ A_j &\leq A_{j-1} + \theta_{j-1}T_{j-1}, j \geq 2 \end{aligned} \quad (3)$$

The DHS model is completely described by the joint distribution of the interfailure times given by

$$\begin{aligned} P(T_j \leq t | A_1, \dots, A_j, \theta_1, \dots, \theta_j, T_1, \dots, T_{j-1}) \\ = F_{A_j, \theta_j}(t) \end{aligned} \quad (4)$$

Note that, for $j \geq 1$, the effective age A_{j+1} of the system after the j th repair is less than the effective age $X_j \stackrel{\text{def}}{=} A_j + \theta_j T_j$ just before the j th failure, and because $\theta_j \leq 1$, X_j is less than the actual age S_j . The DHS model also encompasses an i.i.d. renewal model studied by Peña *et al.* [7] and Hollander and Peña [8].

All of the repair models mentioned in the section titled “Repair Models” are special cases of the DHS model. For example, setting $\theta_j = 1$, $A_j = S_{j-1}$ yields the minimal repair model. For the specifications of the A 's and θ 's that yield the other models of the section titled “Repair Models” and also many new repair models, see Dorado *et al.* [6].

DHS develops asymptotic inferential procedures for model (4) by invoking an analogy to a survival model with censored data. We refer the reader to Dorado *et al.* [6] for details and only summarize the results here.

DHS fixes a $T > 0$ and defines two processes $N(\cdot)$ and $Y(\cdot)$ by

$$N(t) = \sum_j I\{X_j \leq t, S_j \leq t\} \quad (5)$$

and

$$\begin{aligned} Y(t) = \sum_j I\{A_j < t \leq (X_j \wedge [A_j \\ + \theta_j(T - S_{j-1})])\} \end{aligned} \quad (6)$$

where $X_j = A_j + \theta_j T_j$ is the effective age just before the j th failure. DHS assumes that n independent copies of the processes N and Y are observed on

a finite interval. Let N_n and Y_n denote the sum of the n independent copies of N and Y , respectively. Let Λ denote the cumulative hazard function of F . The Nelson–Aalen estimator of Λ of is

$$\hat{\Lambda}_n(t) = \int_0^t \frac{J_n}{Y_n} dN_n \quad (7)$$

where $J_n(t) = I(Y_n(t) > 0)$. The DHS estimator of $\bar{F}$ is

$$\hat{\bar{F}}_n(t) = \prod_{s \leq t} (1 - d\hat{\Lambda}(s)) \quad (8)$$

where $\prod_{s \leq t} (1 - d\hat{\Lambda}(s))$ denotes the product integral (see Gill and Johansen [9]). The estimate $\hat{\bar{F}}_n(t)$ can be shown to be the solution of the Volterra integral equation

$$\hat{\bar{F}}_n(t) = \int_0^t (1 - \hat{\bar{F}}_n(s-)) d\hat{\Lambda}_n(s) \quad (9)$$

DHS obtained weak convergence results for $\hat{\bar{F}}_n(t)$ and also derived a simultaneous confidence band for $\bar{F}$. For $t \in [0, T]$, let $L_n = I(\bar{F}_n(t) < 1)$ and set

$$\begin{aligned} \hat{C}_n(t) &= \int_0^t \frac{J_n L_n}{(Y_n/n)(1 - \hat{\bar{F}}_n)} d\hat{\bar{F}}_n, \\ \hat{K}_n(t) &= \frac{\hat{C}_n(t)}{1 + \hat{C}_n(t)} \end{aligned} \quad (10)$$

For t such $\hat{\bar{F}}_n(t) = 1$, set $\hat{K}_n(t) = 1$. A nonparametric asymptotic simultaneous confidence band for $\bar{F}$ with confidence coefficient at least $100(1 - \alpha)$ is

$$\hat{\bar{F}}_n \pm n^{-1/2} \frac{\lambda_\alpha (1 - \hat{\bar{F}}_n)}{(1 - \hat{K}_n)} \quad (11)$$

where λ_α is such that

$$P\left(\sup_{t \in [0, 1]} |B^0(t)| \leq \lambda_\alpha\right) = 1 - \alpha \quad (12)$$

and B^0 denotes a Brownian bridge on $[0, 1]$.

The following formulas are computational simplifications. Let $X_{(1)}, X_{(2)}, \dots, X_{(r)}$ be the distinct

ordered values of the X 's whose corresponding failure times are within $[0, T]$. Let δ_j be the number of observations with value $X_{(j)}$. Then

$$\hat{F}_n(t) = \prod_{X_{(j)} \leq t} \left(1 - \frac{\delta_j}{Y_n(X_{(j)})}\right) \quad (13)$$

and

$$\hat{C}_n(t) = n \sum_{X_{(j)} \leq t} \frac{\hat{F}_n(X_{(j)}) - \hat{F}_n(X_{(j-1)})}{Y_n(X_{(j)})(1 - \hat{F}_n(X_{(j)}))} \quad (14)$$

Some data sets will lead to $\hat{F}_n(t_0) = 1$ for some $0 < t_0 < T$. For such data sets, the band obtained is a confidence band only on the interval $(0, \sigma)$ where $\sigma = \inf\{t \in [0, T] : \hat{F}_n(t) = 1\}$.

The restriction $\theta_j \in (0, 1]$ imposed by DHS (see equation (3)) does not allow for destructive repair. Last and Szekli (LS) [10] do permit deterioration due to repair by extending Kijima's model II. LS show that their repair model contains many models proposed in the literature including those of Stadjé and Zuckerman [11] and Baxter *et al.* [12].

Additional Nonparametric Methods

Let P denote the probability measure associated with the life distribution F . The restriction P_A of a probability measure P to a set A is defined by

$$P_A(B) = \frac{P(A \cap B)}{P(A)} \text{ if } P(A) > 0 \quad (15)$$

To complete the formal definition when $P(A) = 0$, define $P_A = \mu$ where μ is some conveniently chosen probability measure. Another generalization of the BBS model can be formed by postulating that T_1 has distribution P and that for $n \geq 1$, the distribution of T_{n+1} , given the failure times $T_1, \dots, T_n$ and other environmental factors $E_1, \dots, E_n$ observed at those failure times, is

$$P(T_{n+1} \in B) = P_{A_n}(B) \quad (16)$$

where A_n is a set that depends (measurably) on $T_1, \dots, T_n, E_1, \dots, E_n$. Appropriate choices of the restriction sets $A_1, A_2, \dots$ generate different repair

models. Bayesian nonparametrics for repair model (16) puts a prior distribution on P . For integrated squared-error loss, the Bayes estimator of P is the expected value of the posterior distribution of P , given the data. Sethuraman and Hollander (SH) [13, 14] introduced a new class of priors for Bayesian nonparametrics called *partition-based (PB) priors* and a subclass called *partition-based Dirichlet (PBD) priors*. They utilize these priors to show how to obtain posterior distributions under such priors, and obtain Bayes estimators of F in model (16). When specialized to the BBS model, the SH nonparametric Bayes estimators are competitors to the Whitaker and Samaniego (WS) [15] nonparametric estimator of F (see **Imperfect Repair**).

Kvam *et al.* [16] use isotonic regression techniques to find maximum-likelihood estimators of F and its failure rate when F is known to have an increasing failure rate. Suppose that we have data from a BBS repair consisting of n systems, each observed till the time of a perfect repair. Presnell, Hollander and Sethuraman (PHS) [17] proposed **non-parametric tests** of the minimal repair assumption in the BBS model. Their Mann–Whitney–Wilcoxon-type statistic is of the form $\int \hat{F}_e d\hat{F}$ where $\hat{F}$ is the WS estimator and $\hat{F}_e$ is the empirical distribution based on the initial failure times of n systems. They also studied a Kolmogorov–Smirnov (KS)-type test based on the maximum absolute difference between $\hat{F}$ and $\hat{F}_e$. We briefly describe the KS-type test. A consistent estimator of F is provided by $\hat{F}_e$ and, under the BBS model, $\hat{F}$ also consistently estimates F . If, however, the minimal repair assumption does not hold, $\hat{F}$ may likely diverge from F . Let $\hat{H}$ be the empirical distribution of the time to perfect repair of the n systems. Let

$$D(t) = \int_0^t \frac{dF(s)}{(1 - H(s))(1 - F(s))}$$

$$\hat{D}(t) = \int_0^t \frac{d\hat{F}(s)}{(1 - \hat{H}(s))(1 - \hat{F}(s))}$$

$$L(t) = \frac{1}{(1 - F(t))} - 1 - D(t)$$

$$\hat{L}(t) = \frac{1}{(1 - \hat{F}(t))} - 1 - \hat{D}(t)$$

$$G = \frac{L}{1+L}, \text{ and } \hat{G} = \frac{\hat{L}}{1+\hat{L}} \quad (17)$$

Corollary 2.2 of Presnell *et al.* [17] shows that a test of the minimal repair assumption can be performed by referring

$$S_\tau = \sup_{0 \leq t \leq \tau} \sqrt{n} \frac{(1 - \hat{G}(t))}{(1 - \hat{F}(t))} |\hat{F}(t) - \hat{F}_e(t)| \quad (18)$$

to a table of the supremum of the absolute value of the Brownian bridge over the interval $[0, \hat{G}(\tau)]$ (*cf.* Hall and Wellner [18] and Koziol and Byar [19]).

Hollander *et al.* [20] considered the case where two BBS processes are observed, each until its first perfect repair. They test $H_0: F_1 = F_2$, where F_i is the distribution of the time to first failure for process i , $i = 1, 2$. Their test statistic is $W = \int_0^\infty \hat{F}_1 d\hat{F}_2$ where $\hat{F}_i$ is the WS estimator for the i th process. For the BBS model, Augustin and Peña [21, 22] derived goodness-of-fit tests of $H_0: \Lambda = \Lambda_0$ where Λ_0 is completely specified. Augustin and Peña [23] developed tests for the composite case where the functional form of the failure rate $\lambda_0(\cdot, \xi)$ is known except for the value of ξ . Peña and Hollander (PH) [24] consider a general model for recurrent events that allows for interventions and incorporates covariates. Their general model contains the DHS and LS models as special cases. González *et al.* [25] apply the PH model for modeling intervention effects after cancer relapses.

References

- [1] Ascher, H. (1968). Evaluation of repairable systems using the “bad as old” concept, *IEEE Transactions on Reliability* **17**, 105–110.
- [2] Brown, M. & Proschan, F. (1983). Imperfect repair, *Journal of Applied Probability* **20**, 851–859.
- [3] Block, H., Borges, W. & Savits, T. (1985). Age-dependent minimal repair, *Journal of Applied Probability* **22**, 370–385.
- [4] Kijima, M. (1989). Some results for repairable systems with general repair, *Journal of Applied Probability* **26**, 89–102.
- [5] Uematsu, K. & Nishida, T. (1987). One unit system with a failure rate depending on the degree of repair, *Mathematica Japonica* **32**, 139–147.
- [6] Dorado, C., Hollander, M. & Sethuraman, J. (1997). Nonparametric estimation for a general repair model, *Annals of Statistics* **25**, 1140–1160.
- [7] Peña, E.A., Strawderman, R.L. & Hollander, M. (2001). Nonparametric estimation with recurrent event data, *Journal of the American Statistical Association* **96**, 1299–1315.
- [8] Hollander, M. & Peña, E.A. (2004). Nonparametric methods in reliability, *Statistical Science* **19**, 644–651.
- [9] Gill, R.D. & Johansen, S. (1990). A survey of product-integration with a view toward application in survival analysis, *Annals of Statistics* **18**, 1501–1555.
- [10] Last, G. & Szekli, R. (1998). Asymptotic and monotonicity properties of some repairable systems, *Advances in Applied Probability* **30**, 1089–1110.
- [11] Stadje, W. & Zuckerman, D. (1991). Optimal maintenance strategies for repairable systems with general degree of repair, *Journal of Applied Probability* **28**, 384–396.
- [12] Baxter, L., Kijima, M. & Tortella, M. (1996). A point process model for the reliability of a maintained system subject to general repair, *Communications in Statistics: Stochastic Models* **12**, 37–65.
- [13] Sethuraman, J. & Hollander, M. (2006a). Bayesian methods in repair models, in *Proceedings of the International Conference on Degradation, Damage, Fatigue and Accelerated Life Models in Reliability Testing*, Angers, France, pp. 217–221.
- [14] Sethuraman, J. & Hollander, M. (2006b). Nonparametric Bayes Estimation in Repair Models, Submitted.
- [15] Whitaker, L.R. & Samaniego, F.J. (1989). Estimating the reliability of systems subject to imperfect repair, *Journal of the American Statistical Association* **84**, 301–309.
- [16] Kvam, P.H., Singh, H. & Whitaker, L.R. (2002). Estimating distributions with increasing failure rate in an imperfect repair model, *Lifetime Data Analysis* **8**, 53–67.
- [17] Presnell, B., Hollander, M. & Sethuraman, J. (1994). Testing the minimal repair assumption in an imperfect repair model, *Journal of the American Statistical Association* **89**, 289–297.
- [18] Hall, W.J. & Wellner, J. (1980). Confidence bands for a survival curve from censored data, *Biometrika* **67**, 133–143.
- [19] Koziol, J.A. & Byar, D.P. (1975). Percentage points of the asymptotic distributions of one- and two-sample KS statistics for truncated or censored data, *Technometrics* **17**, 507–510.
- [20] Hollander, M., Presnell, B. & Sethuraman, J. (1992). Nonparametric methods for imperfect repair models, *Annals of Statistics* **20**, 879–896.
- [21] Augustin, Z. & Peña, E.A. (1999). Order statistic properties, random generation, and goodness-of-fit testing for a minimal repair model, *Journal of the American Statistical Association* **94**, 266–272.
- [22] Augustin, Z. & Peña, E.A. (2001). Goodness-of-fit of the distribution of time-to-first-occurrence in recurrent event models, *Lifetime Data Analysis* **7**, 289–306.

- [23] Augustin, Z. & Peña, E.A. (2000). Testing whether the distribution of time to first event occurrence is in a specified family, in *Abstracts Book of the MMR' 2000 Second International Conference on the Mathematical Methods in Reliability: Methodology, Practice and Inference*, Bordeaux, Vol. 1, pp. 47–50.
- [24] Peña, E.A. & Hollander, M. (2004). Models for recurrent phenomena in survival analysis and reliability, in *Mathematical Reliability: An Expository Perspective*, T. Mazzuchi, N. Singpurwalla & R. Soyer, eds, Kluwer, Dordrecht, pp. 105–123.
- [25] González, J.R., Peña, E.A. & Slate, E.H. (2005). Modelling intervention effects after cancer relapses, *Statistics in Medicine* **24**, 3959–3975.

Related Articles

General Minimal Repair Models; Imperfect Repair; Imperfect Repair, Counting Processes; Multivariate Imperfect Repair Models; Repair Data, Sets of: How to Graph, Analyze, and Compare.

MYLES HOLLANDER, FRANCISCO J. SAMANIEGO
AND JAYARAM SETHURAMAN

System Availability

Overview

All mechanical, electrical, nuclear power generation, propulsion, weapons, computer and data storage (both hard and software), communications, and biological/medical systems, and many combinations thereof that are designed, manufactured, tested, operated, and maintained with human skills, engineering, and scientific knowledge, *are subject to failure*. This means that they become unable to perform the intended design function, or mission. *Reliability* is a measure of the propensity of component, subsystem, or system (of systems) to satisfactorily perform a required designed performance function on, say, a military or homeland security mission. This includes effectively and promptly responding to a natural disaster such as a hurricane (Katrina), tornado, or earthquake. During the period of inability to perform, the item is said to be *not available* or *unavailable*. Often this propensity is quantified as a (conditional) probability, where the conditions include stressful aspects of the operational environment and usage. Note that it is not customary to include *vulnerability* to opponent action in measures of military or security system reliability. Many, even most, failed systems are subject to failure correction or rectification, which often requires a nonnegligible total *downtime*, during which time interval they are completely or partially unavailable. Thus an initial informal definition of *availability* can be an estimated probability that, at the moment of demand/need for a system's performance, it will furnish that service, and not be in the process of waiting or being rectified/repaired/replaced. Note that field support of system availability and capability to accomplish a mission profile depends upon *sustainment resources*: personnel, logistics/spare parts, tow vehicles, and so on. These must also be available when needed.

It is often true that the failure propensity measure, or *hazard*, or *hazard rate*, of a component, subsystem, or system eventually increases with time or intensity of usage and "wear" or "age" after being fielded. Typically, a new or modified/upgraded system is tested so as to identify and remove or reduce the effect of *design defects* or *failure modes*; during the early testing period every attempt is made to induce the activation, thus exposure of, *design defects* or *failure*

modes, i.e., to *stimulate* occurrences of failures and remove their causes; see [1]. Some faults inevitably remain, and can be found during routine scheduled maintenance. However, most such maintenance can be unnecessary (time-scheduled maintenance is costly in time and resources, and may well find nothing). It is currently thought to be desirable and increasingly technically feasible, to identify and automatically track *system health* changes: failure-preceding or prognostic events (excess heat, undue vibration, wear of say tire treads, tire, engine oil condition, etc.) that can signal an incipient failure and thereby suggest needed system design or operational changes before total, possibly catastrophic, failure can occur. This procedure is known as *condition-based maintenance* (CBM), and makes use of currently available automated embedded sensors; hence CBM+. However, the CBM+ subsystem is itself subject to malfunction and failure, so false-positive and false-negative diagnoses must be anticipated and made rare in order for overall gains to be made. Thus the reliability/availability of CBM+ must itself be considered as part of the total system package.

Accelerated testing procedures are often used at an early (developmental) testing stage; see [2]. During the testing phase, both developmental and operational, it is intended, likely, and desirable, to *find design faults and, desirably, "fix" them*, so that they will not occur/activate and cause mission-abortive events under active field conditions. It is the usual objective to test and "fix" until times between failures during testing become a stable-in-time random process, often represented by a *time-between-successive-failures distribution*, with the further assumption that the times between failures are statistically independent (*independent identically distributed random variables*, i.i.d.r.v.'s); see [2], and several computer packages, such as ReliaSoft (www.reliasoft.com).

However, initial tests of a newly designed system typically reveal "infant failures" that should be removed before any approximate stability of times between failures occurs. Such removal is referred to as *reliability growth*, and it has been modeled and analyzed extensively, e.g. by Crow [3, 4], e.g. in the software package ReliaSoft; see also [1] for discussion of testing for growth using a sequential (run of successes) stopping rule; such testing is made more complex if the system performance is in sequential stages; failure of an early stage may mask the exposure of later stages, thus requiring

longer tests so as to explore all stages for faults. After (*if*) there is evidence of failure-time stochastic stability it is meaningful to speak of the *mean (operating) time between failures*; this is often specified as a reliability metric, but an estimate of the probability of mission success is usually more relevant, certainly if the item is one-shot, or is usually quiescent in performance: a piece of ammunition or a communication system for emergencies. Caution is required: justification for such a stable failure-time distribution must depend on actual data monitoring, and the currently appropriate particular distributional form (model) may well depend on environmental conditions, including maintenance skill, and usage and environment; e.g., operating times between successive failures for a particular equipment (helicopter) in the Arctic can be expected to differ from those operating under tropical, sandy, desert conditions.

Upon failure, some (but not all, e.g., expendable missiles and ammunition) system components are repaired; this requires maintenance facilities and personnel, and logistics (spare parts); the system may then be again capable of performing the functions it was designed for. Repair or replacement is intended to restore original capability, or possibly upgrade (or in fact unintentionally and occasionally *downgrade*) system function. **Preventive maintenance** (PM), based on the observed or inferred/sensed physical condition of the system, may be used to ward off catastrophic failures. As stated earlier, current attention is being focused on automated information-system-supported *CBM+* (alternatively, integrated vehicle health management, *IVHM*). Such subsystems continuously monitor the “health condition” of other subsystems and provide warning of imminent malfunction/failure. If successful, such a capability should allow more operative mission hours and shorter and less time-consuming and costly repairs. Success depends upon the trustworthiness of the *CBM+/IVHM* systems available; see [5–7].

This chapter considers that times or events between such system failures (during functional usage) are punctuated by periods of “downtime”, during which the system (copy or version thereof) is (a) awaiting maintenance, which may include moving, or being moved, to a repair facility, and awaiting the arrival of parts, instrumentation, and personnel; (b) undergoing diagnosis of faults and subsequent maintenance actions; or (c) awaiting

further operational assignment (in idle standby). Items may fail or degrade during such waiting/idle times, and so preliminary tests may be performed prior to mission assignment: premission testing is often done, especially with aircraft and other platforms. This permits the detection of **imperfect repair** before the mission actually is in progress.

The (classical) system *operational availability* (*Ao*) *parameter* is a measure of the capability of a system to *begin* to perform the designed-for function or mission *on immediate random demand*. Nominally, this would imply that the system is in state (c) above, but too long a sojourn of idleness (in “cold standby”), in a hostile environment can induce *startup unavailability*. *Mission availability*, measuring the capability of a system to perform its function from demand for service until mission termination is usually of far more operational significance than simple *Ao*. The latter is too often treated – inappropriately – as “the long-run fraction of time” (or “constant probability”) that a component or system is “up” or completely mission capable at the time of demand for operation. In reality such a constant value may not prevail; “the” value of *Ao* may well change with time and usage, including maintenance variation between individual copies of a system design. Often no account is taken of system age or condition, or the variable mission environment that may affect it. Further, a system’s availability may be partial: e.g., a multiengine vehicle can still function if degraded by the loss of one out of several engines. However, loss of an engine or, one of several sensors, or a communication link sometimes allows a modified or limited version of the intended mission to proceed until a replacement is furnished. Consequently, availability need not be a binary instantaneous concept.

The object of this chapter is to describe methods in use for describing, measuring, testing, and improving *availability* and to provide cautionary comments similar to those above.

Component and System Operational Classification

Engineered systems of all types are composed of numbers of components (“parts”), subsystems, and entire systems designed to perform cooperative functions. Modern warfare has been said to be conducted by “systems of systems”. Such systems usually operate efficiently and effectively according to correct

design, manufacture, and usage. They inevitably fail or degrade with age and stress; some failure modes are mission and even life threatening, and can occur with little warning. Hence, condition-based monitoring condition-based maintenance (CBM) and PM are required and invoked; time spent in such situations makes the entity temporarily mission unavailable, but can provide greater overall net system mission availability.

Systems, and especially subsystems and components, can be roughly categorized as

- *One shot/disposable* (e.g., fuel, ammunition, missiles, or mines, electronic components, such as computer chips or screens, remote vulnerable sensors, etc.) and/or
- *Repairable (segmentally)/replaceable* (e.g., failure-prone engines and control devices, vehicle chassis, computer hardware, elements of communication networks, etc.) An injured or wounded human can sometimes be included in this category, as being part of an entire system.

In many cases, a failed repairable system becomes less susceptible to repair if too much time elapses, so the *particular item's* availability terminates, but its function is performed by a replacement or substitute; the latter can be a redesigned upgrade; see [1]. Further, the pattern of component/system usage can vary: many items such as some sensors and alarm systems and engines and generators are normally running or “hot” during a mission, so if failure occurs or becomes imminent, symptoms are evident. Others are normally inactive or “cold”, such as an idle aircraft parked on a flight deck or a rescue vehicle.

Many systems are made up of complex assemblies or combinations of one shot/disposable and repairable/replaceable subsystems: a vehicle or platform burns fuel and may fire ammunition, both disposable, but its chassis, propulsion, and steering and sensors and communications are very often repairable/replaceable. Many vehicles operate in groups: convoys of trucks or small boats, task groups, aircraft squadrons, and so on; here the mission may involve all elements of the group, so the timely availability of such a subforce at a particular site can be spoiled if just one of the member elements fails: the entire system may become mission incapable or unavailable if one or more

of such elements fails or is damaged by opponent action (when in military or homeland protection application), or by owner/user mishandling. The tendency for this to occur is enhanced when hostilities occur, and when the sustainment/support forces and logistics are overwhelmed and themselves unavailable.

Modeling Availability

System availability has been characterized probabilistically or stochastically for many years; an initial classic is [8], but aspects of repairable system availability go far back to Erlang, Khinchine, and to Palm in the 1930s and a further few decades, well summarized in [9] as the “repairman problem” or service system. See also [10–11] on survival analysis, essentially summarizing biological (e.g., animal, human “reliability”); failure prone but **repairable systems** resemble biological epidemics: failures of subsystems can stress, or infect, other subsystems that in turn can cause total system collapse if not quickly isolated and mitigated.

The simplest analytical model for a failure prone but repairable system is the alternating renewal process (see, however, practical cautions in the section titled “Overview”). Letting a sequence of uptimes, $\{U_i\}$, and a sequence of downtimes, $\{D_i\}$, be independent sequences of independent random variables with marginal distributions $F_U(u; \underline{\theta})$ and $G_D(v; \underline{\varphi})$, respectively, and $\underline{\theta}$ and $\underline{\varphi}$ representing parameters; then if $A(t)$ is the availability at time t , given that the subsystem starts at the beginning of an up period, U_1 , say, we have the backward recurrence renewal equation

$$A(t) = P\{U_1 > t\} + \int_0^t P\{U_1 + D_1 \in dx\} A(t - x) \quad (1)$$

which can be solved numerically or in terms of Laplace–Stieltjes transforms. Basic renewal theory asymptotics, [12, Chapter 11], shows that if $E[U]$ and $E[D]$ exist/are finite then

$$\lim_{t \rightarrow \infty} A(t) = E[U]/(E[U] + E[D]) = A_o \quad (2)$$

the popular (often misused) operational availability. In cases in which the system must function throughout a mission of length M a more informative measure is mission availability

$$A_m = A_o \frac{P\{U > M\}}{E[U]} \quad (3)$$

an asymptotic renewal theory result.

Complex systems, starting with series/tandem configurations, and extending to series/parallel/redundant combinations can be mathematically analyzed under the initially stated independence assumptions, and with these progressively relaxed.

Statistical Inference on Availability

If the simple assumption of mutually independent sequences of independent identically distributed (i.i.d.) uptime r.v.'s, $\{U_i\}$ and likewise i.i.d. downtimes, $\{D_i\}$ is (provisionally) acceptable then alternating renewal theory shows that the long-run *point availability* is given by $A_o = E[U]/(E[U] + E[D])$, provided expectations exist. A natural estimate of A_o , $\hat{A}_o = \bar{u}/(\bar{u} + \bar{d})$, where $\bar{u}$ and $\bar{d}$ are the sample averages of up and down realizations; assume no censoring; Gaver and Chu assess **confidence intervals** for A_o using jackknifing [13], wherein the logit transform of $\hat{A}_o$ is recomputed, omitting pairs of observations successively. The method can be applied with successfully to various redundant systems and several plausible distributional forms. An alternative would be the **bootstrap** (see [14]). A sensible (semi) nonparametric approach is the nonparametric bootstrap: resample from the empirical distributions of up- and downtimes, provided that preliminary examination of the data does not wildly contradict i.i.d. assumptions; see [15]. In many applied situations, data may be insufficient to validate (or invalidate) such an assumption conclusively. It may well be wise to adaptively smooth up- and downtime means and employ these in the $\hat{A}_o$ formula. Details are available from the references.

Conclusion

Many more, and more complex, results are available from the references. For an excellent general overview see [16].

Acknowledgment

The author thanks Professor Patricia A. Jacobs for her suggestions and contributions.

References

- [1] Gaver, D.P., Jacobs, P.A., Glazebrook, K.D. & Seglie, E.A. (2003). Probability models for sequential-stage system reliability growth via failure mode removal, *International Journal of Reliability, Quality and Safety Engineering* **10**, 15–40.
- [2] Meeker, W.Q. & Escobar, L.A. (1998). *Statistical Methods for Reliability Data*, John Wiley & Sons, New York.
- [3] Crow, L.H. (1974). Reliability analysis for complex repairable systems, in *Reliability and Biometry*, F. Proschan & R.J. Serfling, eds, SIAM, Philadelphia, pp. 379–410.
- [4] Crow, L.H. (2002). Method for achieving an enhanced mission capability, *Proceedings Annual Reliability and Maintainability Symposium*, IEEE, Piscataway, pp. 153–157.
- [5] Macheret, Y. & Koehn, P. (2006). Prognostics and advanced diagnostics for improving steady-state and pulse reliability, *Proceedings 2006 IEEE Aerospace Conference*, IEEE, Piscataway, pp. 1–11.
- [6] Williams, Z. (2006). Benefits of IVHM: an analytical approach, *Proceedings 2006 IEEE Aerospace Conference*, IEEE, Piscataway, pp. 1–8.
- [7] Osborne, B. (2005). Leveraging remote diagnostics data for predictive maintenance, in *Modern Statistical and Mathematical Methods in Reliability*, A.G. Wilson, N. Limnios, S. Keller-McNulty & Y. Armijo, eds, World Scientific Publishing, Singapore, Chapter 25, pp. 353–373.
- [8] Barlow, R.E. & Proschan, F. (1967). *Mathematical Theory of Reliability*, John Wiley & Sons, New York.
- [9] Feller, W. (1968). *An Introduction to Probability Theory and its Applications*, 3rd Edition, John Wiley & Sons, New York, Vol. 1.
- [10] Cox, D.R. (1970). *Renewal Theory*, Methuen & Co, London.
- [11] Cox, D.R. & Oakes, D. (1984). *Analysis of Survival Data*, Chapman & Hall, London.
- [12] Feller, W. (1966). *An Introduction to Probability Theory and its Applications*, John Wiley & Sons, New York, Vol. 2.
- [13] Gaver, D.P. & Chu, B.B. (1979). Jackknife estimates of component and system availability, *Technometrics* **21**, 443–450.
- [14] Efron, B. & Tibshirani, R.J. (1993). *An Introduction to the Bootstrap*, Chapman & Hall, New York.
- [15] Gaver, D.P. & Jacobs, P.A. (1989). System availability: time dependence and statistical inference by (semi) non-parametric methods, *Applied Stochastic Models and Data Analysis* **5**, 357–375.

- [16] Davis, T.P. (2006). Science, engineering, and statistics, *Applied Stochastic Models in Business and Industry* **22**, 401–430.

Growth Modeling; Reliability of Redundant Systems; Renewal Theory; Replacement Strategies; Total Productive Maintenance.

DONALD P. GAVER

Related Articles

Availability Models; Exponentially Weighted Moving Average (EWMA) Control Chart; Group Maintenance Policies; k -out-of- n Systems; Parallel, Series, and Series-Parallel Systems; Reliability

Multivariate Imperfect Repair Models

Introduction

In the study of **maintenance policies** for simple **repairable systems**, i.e. a one-component repairable system with a single repairman, a common assumption is that the system after repair is “as good as new”. This constitutes a perfect repair model. However, because of aging and cumulative wear effects, it is known in practice that repair of a failed system may not yield a functioning system that is as good as new. Barlow and Hunter [1] presented a **minimal repair** model in which repair does not change the age of the system. Brown and Proschan [2] and Berg and Cleroux [3] examined a maintenance action, called the **imperfect repair model**, in which a system is repaired at failure, with probability α , and that it is returned to the as good as new state (perfect repair) and with probability $(1 - \alpha)$ that it is returned to the functioning state, but is only as good as a system of its age at failure (minimal repair). This model has been generalized by Block *et al.* [4] to the case in which α depends on the age of the system at failure. Much work has been performed in this direction since then. See for example, Balaban and Singpurwalla [5], Uematsu and Nishida [6], Langberg [7], Kijima [8], Ebrahimi [9], Last and Szekli [10], Li and Shaked [11], Belzunce and Semeraro [12], Wang and Zhang [13], and many references cited there.

As systems become more complicated and require new technologies and methodologies, more sophisticated repair policies are needed. More specifically, it is quite common to find systems consisting of k components which are related. In such cases, a legitimate question is how to generalize an imperfect repair model to a system with several components. This article deals with multivariate versions of imperfect repair models that may be used for exploring effects of imperfect repairs of some components on the residual lives of components that are still functioning. Throughout this article “increasing” means “nondecreasing” and “decreasing” means “nonincreasing”.

Models, Notations, and Terminology

Consider a system consisting of k components that are related. Let $T = (T_1, \dots, T_k)$ be the vector of lifetimes of components when the components do not undergo any kind of repair. Also, let $T^* = (T_1^*, \dots, T_k^*)$ be the waiting times of components until the first perfect repair starting with new components when they are subjected to imperfect repair. Suppose $\bar{F}$ and $\bar{G}$ are joint survival functions of T and T^* , respectively. That is, $\bar{F}(t_1, \dots, t_k) = P(T_1 > t_1, \dots, T_k > t_k)$ and $\bar{G}(t_1, \dots, t_k) = P(T_1^* > t_1, \dots, T_k^* > t_k)$. We use F and G for joint distribution functions of T and T^* , respectively. For $k = 1$, Brown and Proschan [2] observed that if $\lambda(t) = \frac{f(t)}{\bar{F}(t)}$ is the original hazard function of T_1 , then the resulting hazard function of T_1^* , $\lambda^*(t) = \frac{g(t)}{\bar{G}(t)}$ is $\lambda^*(t) = \alpha\lambda(t)$. Hence $\bar{G}(t) = (\bar{F}(t))^\alpha$. Here f and g are the probability density functions of T_1 and T_1^* , respectively.

Consider models in which, if at some time t_0 , component i_0 fails and is imperfectly repaired, then i_0 is as good as an identical functioning component of age t_0 (minimal repair), and the remaining functioning components “do not know” about the component i_0 ’s failure and imperfect repair. Specifically, suppose that components $j_1, \dots, j_\ell$ have been perfectly repaired at times $t_{j_1}, \dots, t_{j_\ell}$ respectively, and $t_{j_i} < t_0$, for $i = 1, \dots, \ell$. Also, components $i_1, \dots, i_m$ are still functioning at time t_0 and the component i_0 has just been imperfectly repaired. Note that $\{j_1, \dots, j_\ell, i_1, \dots, i_m, i_0\} = \{1, \dots, k\}$. Then, it is clear that the distribution of the times to next failure of components $i_0, i_1, \dots, i_m$ is

$$P\left\{T_{i_0} \leq t_{i_0}, T_{i_1} \leq t_{i_1}, \dots, T_{i_m} \leq t_{i_m} | T_{i_0} > t_0, T_{i_1} > t_0, \dots, T_{i_m} > t_0; T_{j_1} = t_{j_1}, \dots, T_{j_\ell} = t_{j_\ell}\right\} \quad (1)$$

On the other hand, if the component i_0 at time t_0 has been perfectly repaired, then

$$P\left\{T_{i_1} \leq t_{i_1}, \dots, T_{i_m} \leq t_{i_m} | T_{i_1} > t_0, \dots, T_{i_m} > t_0; T_{j_1} = t_{j_1}, \dots, T_{j_\ell} = t_{j_\ell}; T_{i_0} = t_0\right\} \quad (2)$$

Assuming that $\bar{F}(t_1, \dots, t_k)$ is absolutely continuous, i.e., no more than one component can fail at a

2 Multivariate Imperfect Repair Models

time, below we describe two multivariate imperfect repair models:

Model 1 k components start to function at (the same) time zero. Upon failure, a component undergoes a repair. With probability α the repair is unsuccessful and the component is scrapped (perfect repair). With probability $(1 - \alpha)$, the repair is imperfect repair. It is clear that the joint distribution of the times to next failure of the functioning components after an imperfect (perfect) repair is given by equations (1) and (2).

Model 2 k components start to function at time zero. Upon failure, a component undergoes a repair. If n components ($n = 0, 1, \dots, k - 1$) have already been perfectly repaired, with probability α_{n+1} the repair of the component is perfect and with probability $1 - \alpha_{n+1}$ the repair is imperfect repair. It is clear that the joint distribution of the times to next failure of functioning components after an imperfect (perfect) repair is given by equations (1) and (2). For more details about these models, see Shaked and Shanthikumar [14] and Natvig [15].

If we consider the case in which there is a positive probability that at least two of the original components can fail at the same time, that is, $\bar{F}(t_1, \dots, t_k)$ has singularities on sets of the form $\{(t_1, \dots, t_k); t_{i_1} = \dots = t_{i_\ell}\}$ for some $1 \leq i_1 < i_2 < \dots < i_\ell \leq k$. Then, the following two interpretations of model 1 are possible.

Interpretation 1 If two or more components fail at the same time, all the failed components are scrapped with probability α or all the failed components are minimally repaired with probability $(1 - \alpha)$.

Interpretation 2 If two or more components fail at the same time, then each of them, independently of the others, is either successfully minimally repaired with probability $(1 - \alpha)$ or scrapped with probability α .

The next model gives a very general repair policy when $\bar{F}(t_1, \dots, t_k)$ has singularities. To describe this model, let S denote the set of vectors $(s_1, s_2, \dots, s_k)$, where $s_j = 0$ or 1 for $j = 1, \dots, k$. Here $s_j = 0$ means that the component j is minimally repaired and $s_j = 1$ means that the component j is scrapped.

Model 3 k components start at time zero. Upon failure, a component undergoes a repair. There are several sources which cause the failures of components $1, 2, \dots$, and k . Source $\ell(j)$ causes the failure of component j ; with probability $\alpha(j)$ the repair

of this component is perfect and with probability $(1 - \alpha(j))$ the repair is imperfect, $j = 1, \dots, k$. For given i and $j, i \neq j, i, j = 1, \dots, k$, source $\ell(i, j)$ causes the failures of both components i and j , with probability $\alpha(i, j; s_i, s_j)$ the component i (component j) is minimally repaired if $s_i = 0$ ($s_j = 0$). Otherwise it is scrapped. For example, if $s_i = s_j = 1$, then $\alpha(i, j, 0, 0)$ stands for the probability that both components are minimally repaired. If $s_i = 0, s_j = 1$, then $\alpha(i, j, 1, 0)$ stands for the probability that component j is scrapped and component i is minimally repaired. Here $\alpha(i, j, 1, 1) + \alpha(i, j, 0, 0) + \alpha(i, j, 1, 0) + \alpha(i, j, 0, 1) = 1$. For a given $i, j, m; i \neq j \neq m, i, j, m = 1, \dots, k$, the source $\ell(i, j, m)$ causes the failure of three components i, j , and m ; with probability $\alpha(i, j, m; s_i, s_j, s_m)$ either one or two or three are minimally repaired depending on whether s_i, s_j , or s_m are zero or one. Again here, $\sum_{s_i} \sum_{s_j} \sum_{s_m} \alpha(i, j, m, s_i, s_j, s_m) = 1$. Continuing this, suppose source $\ell(1, 1, \dots, 1)$ causes the failures of all k components. Then $\alpha(1, 1, \dots, 1, s_1, \dots, s_k)$ is the probability that either one, two, ... or all k components are minimally repaired depending on whether $s_1, s_2, \dots$, or s_k are zero or one. Note that $\sum_S \alpha(1, 1, \dots, 1, s_1, s_2, \dots, s_k) = 1$. It is clear that $k = 2$, then $\alpha(1)$ and $\alpha(2)$ are the probabilities of perfect repairs of components 1 and 2, respectively. Also $\alpha(1, 1, 0, 0)$, $\alpha(1, 1, 0, 1)$, $\alpha(1, 1, 1, 0)$, and $\alpha(1, 1, 1, 1)$ are the probabilities that both components are minimally repaired, component 1 is minimally repaired, component 2 is minimally repaired, and both components are scrapped. This coincides with the model (3.5) of Shen and Griffith [16].

Main Results

In this section, we obtain the distribution of G in terms of F . We treat the bivariate case ($k = 2$) to keep the notations manageable. The results, however, can be easily extended to the k -dimensional case.

The following theorem gives the distribution of G in terms of F under Models 1 and 2. See Shaked and Shanthikumar [14] for more details.

Theorem 1 Suppose that (T_1, T_2) has the joint density function f and (T_1^*, T_2^*) has the joint density function g :

(a) Under Model 1,

$$\begin{aligned}
 g(t_1, t_2) &= f(t_1, t_2) \alpha^2 (\bar{F}(t_1, t_2))^{\alpha-1} \\
 &\times \left[\frac{\bar{F}_2(t_2|t_1)}{\bar{F}_2(t_1|t_1)} \right]^{\alpha-1} \text{ if } 0 \leq t_1 \leq t_2 \quad (3) \\
 &= f(t_1, t_2) \alpha^2 (\bar{F}(t_1, t_2))^{\alpha-1} \\
 &\times \left[\frac{\bar{F}_1(t_1|t_2)}{\bar{F}_1(t_2|t_2)} \right]^{\alpha-1} \text{ if } 0 \leq t_2 \leq t_1 \quad (4)
 \end{aligned}$$

where $\bar{F}_1(x|y) = P(T_1 > x | T_2 = y)$ and $\bar{F}_2(x|y) = P(T_2 > x | T_1 = y)$.

(b) Under Model 2,

$$\begin{aligned}
 g(t_1, t_2) &= f(t_1, t_2) \alpha_1 \alpha_2 (\bar{F}(t_1, t_2))^{\alpha_1-1} \\
 &\times \left[\frac{\bar{F}_2(t_2|t_1)}{\bar{F}_2(t_1|t_1)} \right]^{\alpha_2-1} \text{ if } 0 \leq t_1 \leq t_2 \\
 &= f(t_1, t_2) \alpha_1 \alpha_2 (\bar{F}(t_1, t_2))^{\alpha_1-1} \\
 &\times \left[\frac{\bar{F}_1(t_1|t_2)}{\bar{F}_1(t_2|t_2)} \right]^{\alpha_2-1} \text{ if } 0 \leq t_2 \leq t_1 \quad (5)
 \end{aligned}$$

From Theorem 1, under Model 2, one can easily show that

$$\begin{aligned}
 P(T_1^* \leq T_2^*) &= \int_0^\infty \alpha_1 (\bar{F}(t, t))^{\alpha_1-1} \\
 &\times \bar{F}_2(t|t) f_1(t) dt \quad (6)
 \end{aligned}$$

where f_1 is the marginal density of T_1 . Also, the marginal density function of T_1^* is

$$\begin{aligned}
 g_1(t) &= \alpha_1 (\bar{F}(t, t))^{\alpha_1-1} f_1(t) \bar{F}_2(t|t) \\
 &+ \alpha_1 \alpha_2 \int_0^t \left\{ f(t, x) (\bar{F}(x, x))^{\alpha_1-1} \right. \\
 &\times \left. \left[\frac{\bar{F}_1(t|x)}{\bar{F}_1(x|x)} \right]^{\alpha_2-1} \right\} dx \quad (7)
 \end{aligned}$$

The density of T_2^* can be written in a similar manner. Suppose that

$$\bar{F}(t_1, t_2) = P(U_1 > t_1, U_2 > t_2,$$

$$\begin{aligned}
 U_{12} &> \max(t_1, t_2)) \\
 &= \bar{F}_1(t_1) \bar{F}_2(t_2) \bar{F}_{12}(\max(t_1, t_2)) \quad (8)
 \end{aligned}$$

where U_1, U_2 , and U_{12} are independent random variables having cumulative distribution functions F_1, F_2 , and F_{12} , respectively.

Given equation (8), under Model 3, the following theorem gives the distribution of G in terms of F_1, F_2 , and F_{12} . See Shen and Griffith [16] for more details.

Theorem 2 Under Model 3,

$$\begin{aligned}
 \bar{G}(t_1, t_2) &= (\bar{F}_1(t_1))^{\alpha(1)} (\bar{F}_{12}(\max(t_1, t_2)))^{\alpha(1,1,0)} \\
 &\times (\bar{F}_2(t_2))^{\alpha(2)} (\bar{F}_{12}(\max(t_1, t_2)))^{\alpha(1,1,0,1)} \\
 &\times (\bar{F}_{12}(\max(t_1, t_2)))^{\alpha(1,1,1,1)} \quad (9)
 \end{aligned}$$

Examples

In this section, we apply our results to two families of distributions.

Marshall–Olkin Distribution

The random vector (T_1, T_2) is said to have the bivariate Marshall–Olkin exponential distribution, see Marshall and Olkin [17], if

$$\begin{aligned}
 \bar{F}(t_1, t_2) &= \exp(-\lambda_1 t_1 - \lambda_2 t_2 - \delta \max(t_1, t_2)), \\
 t_1, t_2 &\geq 0, \lambda_1, \lambda_2 > 0, \delta \geq 0 \quad (10)
 \end{aligned}$$

Given that $\alpha(1) = \alpha(2) = \frac{1}{2}$ and $\alpha(1, 1, 1, 0) = \alpha(1, 1, 1, 1) = \alpha(1, 1, 0, 1) = \alpha(1, 1, 0, 0) = \frac{1}{4}$. From Theorem 2,

$$\begin{aligned}
 \bar{G}(t_1, t_2) &= \exp \left(-\frac{\lambda_1}{2} t_1 - \frac{\lambda_2}{2} t_2 \right. \\
 &\quad \left. - \frac{3}{4} \delta \max(t_1, t_2) \right) \quad (11)
 \end{aligned}$$

4 Multivariate Imperfect Repair Models

Bivariate Gumble Model

Consider T_1 and T_2 with the bivariate Gumble distribution, see Gumble [18],

$$\begin{aligned}\bar{F}(t_1, t_2) &= \exp(-\lambda_1 t_1 - \lambda_2 t_2 - \delta t_1 t_2), t_1, t_2 \geq 0, \\ \lambda_1, \lambda_2 &> 0, \delta \geq 0\end{aligned}\quad (12)$$

It is clear that

$$\begin{aligned}f(t_1, t_2) &= [(\lambda_1 + \delta t_2)(\lambda_2 + \delta t_1) - \delta] \\ &\quad \times \exp(-\lambda_1 t_1 - \lambda_2 t_2 - \delta t_1 t_2)\end{aligned}\quad (13)$$

$$\begin{aligned}\bar{F}_2(x|y) &= P(T_2 > x | T_1 = y) \\ &= \frac{(\lambda_1 + \delta x)}{\lambda_1} \exp(-\lambda_2 x - \delta x y)\end{aligned}\quad (14)$$

and

$$\begin{aligned}\bar{F}_1(x|y) &= P(T_1 > x | T_2 = y) \\ &= \frac{(\lambda_2 + \delta x)}{\lambda_2} \exp(-\lambda_1 x - \delta x y)\end{aligned}\quad (15)$$

Now, suppose that $\alpha_1 = \alpha_2 = 1/2$, then from Theorem 1

$$\begin{aligned}g(t_1, t_2) &= \frac{1}{4} [(\lambda_1 + \delta t_2)(\lambda_2 + \delta t_1) - \delta] \\ &\quad \times \exp\left(-\frac{1}{2}(\lambda_1 t_1 + \lambda_2 t_2 + \delta t_1 t_2)\right) \\ &\quad \times \frac{\lambda_1 + \delta t_2}{\lambda_1 + \delta t_1} \left[\frac{\exp(-\lambda_2 t_2 - \delta t_1 t_2)}{\exp(-\lambda_2 t_1 - \delta t_1^2)} \right], \\ 0 &\leq t_1 \leq t_2\end{aligned}\quad (16)$$

and

$$\begin{aligned}g(t_1, t_2) &= \frac{1}{4} [(\lambda_1 + \delta t_2)(\lambda_2 + \delta t_1) - \delta] \\ &\quad \times \exp\left(-\frac{1}{2}(\lambda_1 t_1 + \lambda_2 t_2 + \delta t_1 t_2)\right) \\ &\quad \times \frac{\lambda_2 + \delta t_1}{\lambda_2 + \delta t_2} \left[\frac{\exp(-\lambda_1 t_1 - \delta t_1 t_2)}{\exp(-\lambda_1 t_2 - \delta t_2^2)} \right], \\ 0 &\leq t_2 \leq t_1\end{aligned}\quad (17)$$

Concluding Remarks

In this article, we describe several extensions of univariate imperfect repair models to multivariate imperfect models. These models are applicable for situations where a system has k components that are related. It should be noted that proposed models in this article do not incorporate any maintenance schedule or cost. For such models, we refer you to Li and Xu [19], Wang and Zhang [13], and many references cited there.

References

- [1] Barlow, R.E. & Hunter, L.C. (1960). Optimum preventive maintenance policy, *Operations Research* **8**, 851–859.
- [2] Brown, M. & Proschan, F. (1983). Imperfect repairs, *Journal of Applied Probability* **20**, 851–859.
- [3] Berg, M. & Cleroux, R. (1982). A marginal cost analysis for an age replacement policy with minimal repair, *Infor* **20**, 256–263.
- [4] Block, H., Borges, W. & Savits, W.S. (1985). Age dependent minimal repair, *Journal of Applied Probability* **22**, 370–385.
- [5] Balaban, H.S. & Singpurwalla, N.D. (1984). Stochastic properties of a sequence of interfailure times under minimal repair and survival, in *Reliability Theory and Models*, M.A. Abdel-Hameed, E. Cinlar & J. Quinn, eds, Academic press, New York, Vol. 65–80.
- [6] Uematsu, K. & Nishida, T. (1987). One unit system with a failure rate depending upon the degree of repair, *Mathematica Japonica* **32**, 139–147.
- [7] Langberg, N. (1988). Comparison of replacement policies, *Journal of Applied Probability* **25**, 780–788.
- [8] Kijima, M. (1989). Some results for repairable systems, *Journal of Applied Probability* **26**, 89–102.
- [9] Ebrahimi, N. (1993). Modeling repairable systems, *Advances in Applied Probability* **25**, 926–938.
- [10] Last, G. & Szekli, R. (1998). Stochastic comparison of repairable systems by coupling, *Journal of Applied Probability* **35**, 348–370.
- [11] Li, H. & Shaked, M. (2003). Imperfect repair model with preventive maintenance, *Journal of Applied Probability* **40**, 1043–1059.
- [12] Belzunce, F. & Semeraro, P. (2004). Preservation of positive and negative orthant dependence concepts under mixture applications, *Journal of Applied Probability* **41**, 961–974.
- [13] Wang, G.J. & Zhang, Y.L. (2006). Optimal periodic preventive repair and replacement policy assuming geometric process repair, *IEEE Transactions on Reliability* **55**, 118–122.
- [14] Shaked, M. & Shanthikumar, J.G. (1986). Multivariate imperfect repair, *Operations Research* **34**, 437–448.

- [15] Natvig, B. (1990). On Information based minimal repair and the reduction in remaining system lifetime due to the failure of specific module, *Journal of Applied Probability* **27**, 365–375.
- [16] Shen, S.H. & Griffith, W.S. (1992). Multivariate imperfect repair, *Journal of Applied Probability* **29**, 947–956.
- [17] Marshall, A.W. & Olkin, I. (1967). A multivariate exponential distribution, *Journal of the American Statistical Association* **62**, 30–44.
- [18] Gumble, E.J. (1960). Bivariate exponential distribution, *Journal of the American Statistical Association* **55**, 698–707.
- [19] Li, H. & Xu, S. (2006). A Multivariate Cumulative Damage Shock Model with Block Preventive Maintenance, *Stochastic models* **22**, 341–360.

Related Articles

Block Replacement; General Minimal Repair Models; Markov Renewal Processes in Reliability Modeling; Multivariate Age and Multivariate Renewal Replacement; Multivariate Stochastic Orders and Aging; Replacement Strategies.

NADER EBRAHIMI

General Minimal Repair Models

Introduction

We consider a system subject to random deterioration and failure. At failures, repairs or replacements are performed and the system resumes operation. To model this situation we have to describe the condition of the system after the maintenance (repair/replacement) action has been performed. Some typical cases are:

- The **maintenance action** is a major repair that brings the unit to a condition that is considered to be as good as new, or equivalently, the unit is physically replaced by a new and identical unit. This maintenance action is referred to as a *replacement* [1, 2].
- The maintenance action is such that the unit is considered to be as good as it was immediately before the failure occurred. This maintenance action is referred to as a (statistical) *minimal repair*.
- The maintenance action brings the unit to a condition that is somewhere between the result of a replacement and a **minimal repair**.
- The result of the maintenance action is unsuccessful, for example in the way that the wrong part is replaced/repared or that some damage is inflicted on the unit during the maintenance. This could result in a higher failure proneness, and is often referred to as an *imperfect repair* [3].
- In some occasions information about the cause of failure makes it possible to improve the unit. The result of the maintenance action could then be an improved failure intensity of the unit.

Of course, when developing a model of a system it is important to strike a balance between the following two desired properties. The model needs to be simple enough so that it can be used to study the system by mathematical/statistical methods, and it must be a sufficiently accurate representation of the system.

In this article we restrict attention to the minimal repairs. Using the minimal repair concept it is possible to describe in a simple way the fact that many repairs in real life bring the system to a condition which is basically the same as it was just before the

failure occurred. Such a repair may be used to model a system where a component of the system is replaced or repaired. Of course, the purpose of the repair action is not to bring the system to the exact same condition. Rather the purpose is to bring the system back to operation as soon as possible. But by looking at the condition of the system after the repair, it is a reasonable assumption to say that the system state has not changed.

To formalize this idea, assume that the system is installed at time $t = 0$ and is subject to failure at a random time T_1 . The system has a failure rate (intensity) at time t given by the function $\lambda(t)$, meaning that the probability that the system fails during the interval $(t, t + h)$ is $\lambda(t)h$ if the system has survived until time t . Here h is a small number. Suppose the repair is minimal in the sense that the state or condition of the system is as good as it was immediately before the failure occurred. The minimal repair means that the age of the system is not disturbed by the failures. Consequently the failure rate at time t is $\lambda(t)$ independent of the number of failures occurred up to time t . We ignore the duration of the repairs. If N_t represents the number of failures in the time period $[0, t]$, it follows that N_t is a nonhomogeneous **Poisson process** with **intensity function** $\lambda(t)$.

This is the basic minimal repair model presented in 1960 by Barlow and Hunter [4]. This model has been extended in many ways since then, see for example Aven [5, 6], Aven and Jensen [7, 8], Phelps [9], Bergman [10], Block *et al.* [11], Stadjé and Zuckerman [12], Shaked and Shanthikumar [13], Beichelt [14], Zhang and Jardine [15], Finkelstein [16], and Zequeira and Berenguer [17]. A Bayesian approach is presented and discussed in Mazzuchi and Soyer [18].

Many special cases of the basic minimal repair model have been addressed in the literature. Two of the most frequently studied special cases are

$$\lambda(t) = \lambda\beta(\lambda t)^{\beta-1} \text{ (power law, [19])} \quad (1)$$

$$\lambda(t) = \lambda e^{\beta t} \text{ (log-linear model)} \quad (2)$$

In this article we focus on extensions of the basic model, and give special attention to the general minimal repair models presented and discussed by Aven [5] and Aven and Jensen [8].

2 General Minimal Repair Models

The minimal repair models are used to describe and predict the performance of **repairable systems**. A number of applications relate to optimal maintenance and replacement policies, see the references cited above and [1, 2]. A simple application is presented in the section titled “An Application to Optimal Replacement”.

In the literature it is often distinguished between a physical minimal repair and a statistical minimal repair (also referred to as a *black box minimal repair*) are often considered different from each other, see for example Aven and Jensen [7] and Boland and El-Newehi [20]:

- If a component of a complex system is replaced, the system as a whole is in approximately the same condition as it was immediately before the failure occurred (physical minimal repair).
- The system is repaired and the repair brings the system back into the same state it was in just prior to the failure (statistical minimal repair).

If we ignore the information about the state of the component in the physical minimal repair case, the repair is an approximate statistical minimal repair. This shows the importance of clarifying the information level. It makes a difference whether all components of a complex system are observed or only failure of the whole system is recognized. In the first case we have information about the condition of a specific component of the system, whereas in the second case the only information about the condition of the system is the age. Hence a minimal repair in the statistical case would mean replacing the system by another one of the same age that has not yet failed.

Example 1 We consider a simple two-component parallel system with independent distributed component lifetimes X_1, X_2 and allow for exactly one minimal repair. For both components the lifetime distribution is exponential with parameter 1 (denoted $\text{Exp}(1)$).

- Physical minimal repair (statistical repair at component level). After failure at $T_1 = X_1 \vee X_2$, that is at the largest lifetime, the component that caused the system to fail is repaired minimally (statistical minimal repair). Since the component lifetimes are exponentially distributed, the additional lifetime is given by an $\text{Exp}(1)$ random

variable X_3 independent of X_1 and X_2 . The total lifetime $T_1 + X_3$ has distribution

$$P(T_1 + X_3 > t) = e^{-t}(2t + e^{-t}) \quad (3)$$

- Statistical minimal repair at system level. The lifetime $T_1 = X_1 \vee X_2$ until the first failure of the system has distribution $P(T_1 \leq t) = (1 - e^{-t})^2$ and failure rate $\lambda(t) = 2 \frac{1 - \exp(-t)}{2 - \exp(-t)}$. The additional lifetime $T_2 - T_1$ until the second failure is assumed to have conditional distribution

$$\begin{aligned} P(T_2 - T_1 \leq x | T_1 = t) &= P(T_1 \leq t + x | T_1 > t) \\ &= 1 - e^{-x} \frac{2 - e^{-(t+x)}}{2 - e^{-t}} \end{aligned} \quad (4)$$

Integrating leads to the distribution of the total lifetime T_2 :

$$\begin{aligned} P(T_2 > t) &= e^{-t}(2 - e^{-t}) \\ &\quad \times (1 + t - \ln(2 - e^{-t})) \end{aligned} \quad (5)$$

It is (perhaps) no surprise that the total lifetime after a statistical minimal repair is stochastically greater than that after a physical minimal repair:

$$P(T_2 > t) \geq P(T_1 + X_3 > t), \text{ for all } t \geq 0 \quad (6)$$

We see that there is a need for being precise about the level of information. This leads us to the modern framework of point processes, which allow us to specify the failure intensity as a function of the history of the system and incorporate information about the condition of the system for example through measurements of wear characteristics and damage inflicted on the system. This framework provides a general setup for analyzing repairable systems and minimal repairs in particular. We refer to works by Aven [5, 6, 21], Aven and Jensen [7, 8], Bergman [10], Arjas and Norros [22], and Natvig [23].

A General Minimal Repair Model

Let X_t , $t \geq 0$, be an observable stochastic process, possibly a vector process, representing the condition of the system at time t . We define N_t as the number of failures in $[0, t]$. The failure intensity process, which is denoted λ_t , may depend on X_s , $0 \leq$

$s \leq t$. Often we can formulate the relation in the following way:

$$\lambda_t = v(X_t) \quad (7)$$

where $v(x)$ is a positive deterministic function. The interpretation of λ_t is that given the history of the system up to time t , the probability that the system fails in the interval $(t, t+h)$ is approximately $\lambda_t h$. In other words, λ_t is the expected number of failures per unit of time, given the information up to time t .

If the failure intensity process depends only on the state process X_t and not on the failure process N_t , we can interpret the repairs as minimal: a repair that changes neither the age of the system nor the information about the condition of the system. In this case, the running information about the condition of the system can be thought to be related to a system that is always functioning.

A way of formalizing this is to assume that given the state process X , the failure process follows a nonhomogeneous Poisson process with intensity $v(X_t)$, that is N is a doubly stochastic Poisson process (Cox process) with intensity $v(X_t)$. Hence if the whole X process is known, the intensity at time t is $v(X_t)$, independent of previous failures. This process models a system which is minimally repaired, as a repair does not change the information about the condition of the system. If we know that the state is x at time t , the failure intensity is $v(x)$, independent of the history of the failure process.

If $X_t = t$ for all t , we are back to the basic minimal repair model, where N_t is a nonhomogeneous Poisson process.

Example 2 Shock model [24]. Assume that shocks occur to the system at random times, each shock causes a random amount of damage, and these damages accumulate additively. At a shock the system fails with a given probability. A system failure can occur only at the occurrence of a shock.

Let V_t denote the number of shocks in $[0, t]$, and let Y_i denote the amount of damage caused by the i th shock. We assume that V_t is a Poisson process with rate ν , and that the Y_i 's are independent and identically distributed random variables with a distribution $H(y)$. Let X_t denote the accumulated

damage in $[0, t]$, that is,

$$X_t = \sum_{i=1}^{V_t} Y_i \quad (8)$$

Now, if the system is active before time t , and the accumulated damage equals x , and a jump of size y occurs at t , then the probability of failure at this point is $p(x+y)$, where $0 < p(x) < 1$ for all x .

This model is a special case of the general setup described above. The failure intensity process of the counting process N_t equals

$$\nu \int_0^\infty p(X_t + y) dH(y) \quad (9)$$

For a formal proof of this result, see [7, 8]. These references also include extensions of this model, by allowing the probability p to depend on the number of failures occurred. However, the repairs are then not minimal repairs.

Example 3 Monotone system. Consider a binary, monotone system ϕ of n independent components. We refer to [7, 25, 26] for the definition of such a system. Let $N_i(t)$ denote the number of failures of component i in $[0, t]$, and N_t the number of system failures in the same interval. The counting process $N_t(i)$ is assumed to have an intensity process $\lambda_t(i)$. Hence the failure process of the system N_t has an intensity λ_t given by

$$\lambda_t = \sum_{i=1}^n \lambda_t(i) X_t(i) (1 - \phi(0_i, \mathbf{X}_t)) \phi(\mathbf{X}_t) \quad (10)$$

where $\phi(\cdot, \mathbf{x}) = \phi(x_1, \dots, x_{i-1}, \cdot, x_{i+1}, \dots, x_n)$.

Observe that $X_t(i)(1 - \phi(0_i, \mathbf{X}_t))\phi(\mathbf{X}_t)$ is either 0 or 1, and equals 1 if and only if the system is functioning, component i is functioning and the system fails if component i fails.

Assume now that if the system fails, it is minimally repaired in the following sense: If a component fails and causes system failure, then this component is minimally repaired (statistical minimal repair). A component that fails without causing system failure is not repaired.

We assume that component i when not being repaired has a lifetime R_i with distribution function $F_i(t)$ and failure rate equal to $r_i(t)$. The n components are assumed to be independent.

4 General Minimal Repair Models

It follows that we have a special case of the general setup with λ_t having the form equation (10) with $\lambda_t(i) = r_i(t)X_t(i)$. The process $X_t(i)$, $t \geq 0$, is in this case either identical to 1, or 1 up to R_i and then 0. If component i is in series with the rest of the system, then $X_t(i) \equiv 1$.

Formal Mathematical Setup and Definition of a Minimal Repair

Let (T_n) , $n \in \mathbb{N}$, where $\mathbb{N} = \{1, 2, \dots\}$, be a point process on a basic probability space $(\Omega, \mathcal{F}, P)$ describing the failure times at which instantaneous repairs are carried out, that is, (T_n) is an increasing sequence of positive random variables which may also take the value $+\infty$: $0 < T_1 \leq T_2 \leq \dots$. The inequality is strict unless $T_n = \infty$. We always assume that $T_\infty = \lim_{n \rightarrow \infty} T_n = \infty$. The corresponding counting process $N = (N_t)$, $t \in \mathbb{R}_+$,

$$N_t(\omega) = \sum_{n \geq 1} I(T_n(\omega) \leq t) \quad (11)$$

where $I(\cdot)$ denotes the indicator function, represents the number of failures up to time t . Here $\mathbb{R}_+ = [0, \infty)$.

The information up to time t is represented by the pre- t -history (σ algebra) $\mathcal{F}_t$, which contains all events of $\mathcal{F}$ that can be distinguished up to and including time t . The filtration $\mathbb{F} = (\mathcal{F}_t)$, $t \in \mathbb{R}_+$, is assumed to follow the usual conditions of completeness and right continuity. The main assumption now is that the failure counting process is integrable and admits an $\mathbb{F}$ -intensity $\lambda = (\lambda_t)$, that is a decomposition

$$N_t = \int_0^t \lambda_s ds + M_t \quad (12)$$

with a mean zero martingale $M = (M_t)$. The question now is which counting processes can be classified as minimal repair processes (MRPs) and which not.

Different types of repair processes are characterized by different intensities, λ . The repairs are minimal if the intensity λ is not affected by the occurrence of failures or in other words, if one cannot determine the failure time points from the observation of λ . More formally, minimal repairs can be characterized as follows.

Definition 1. Let (T_n) , $n \in \mathbb{N}$, be a point process with an integrable counting process N and corresponding $\mathbb{F}$ intensity λ . Suppose that $\mathbb{F}^\lambda = (\mathcal{F}_t^\lambda)$, $t \in$

$\mathbb{R}_+$, is the filtration generated by λ : $\mathcal{F}_t^\lambda = \sigma(\lambda_s, 0 \leq s \leq t)$. Then the point process (T_n) is called an MRP if none of the variables T_n , $n \in \mathbb{N}$, for which $P(T_n < \infty) > 0$, is an $\mathbb{F}^\lambda$ stopping time, that is for all $n \in \mathbb{N}$ with $P(T_n < \infty) > 0$ there exists $t \in \mathbb{R}_+$ such that $\{T_n \leq t\} \notin \mathcal{F}_t^\lambda$.

This is a rather general definition which comprises a lot of special cases. It is easily verified that the nonhomogeneous Poisson process with a time-dependent deterministic function $\lambda_t = \lambda(t)$ is an MRP, because here $\mathcal{F}_t^\lambda = \{\Omega, \emptyset\}$ for all $t \in \mathbb{R}_+$, so clearly the failure times T_n are not $\mathbb{F}^\lambda$ stopping times.

If the intensity is not deterministic but a random variable $\lambda(\omega)$ which is known at the time origin (λ is $\mathcal{F}_0$ -measurable) or, more generally, $\lambda = (\lambda_t)$ is a stochastic process such that λ_t is $\mathcal{F}_0$ -measurable for all $t \in \mathbb{R}_+$, that is $\mathcal{F}_0 = \sigma(\lambda_s, s \in \mathbb{R}_+)$ and $\mathcal{F}_t = \mathcal{F}_0 \vee \sigma(N_s, 0 \leq s \leq t)$, then the process is called a doubly stochastic Poisson process or a Cox process. The failure (minimal repair) times are not $\mathbb{F}^\lambda$ -stopping times, since $\mathcal{F}_t^\lambda = \sigma(\lambda) \subset \mathcal{F}_0$ and T_n is not $\mathcal{F}_0$ -measurable.

The physical minimal repair in Example 1 can be extended to a point process with intensity $\lambda_t = I(X_1 \wedge X_2 \leq t)$. In this case $\mathbb{F}^\lambda$ is generated by the minimum of X_1 and X_2 . The first failure time of the system, T_1 , equals $X_1 \vee X_2$, which is not an $\mathbb{F}^\lambda$ stopping time so that this point process is an MRP according to the definition.

In the following we give another characterization of an MRP.

Theorem 1. Assume that $P(T_n < \infty) = 1$ for all $n \in \mathbb{N}$ and that there exist versions of conditional probabilities $F_t(n) = E[I(T_n \leq t) | \mathcal{F}_t^\lambda]$ such that for each $n \in \mathbb{N}$ $(F_t(n))$, $t \in \mathbb{R}_+$, is a ($\mathbb{F}^\lambda$ progressive) stochastic process.

1. Then the point process (T_n) is an MRP if and only if for each $n \in \mathbb{N}$ there exists some $t \in \mathbb{R}_+$ such that

$$P(0 < F_t(n) < 1) > 0 \quad (13)$$

2. If furthermore $(F_t) = (F_t(1))$ has P -a.s. continuous paths of bounded variation on finite intervals, then

$$1 - F_t = \exp\left\{-\int_0^t \lambda_s ds\right\} \quad (14)$$

See Aven and Jensen [8] for the proof.

Example 4 We return to Example 1. We now allow for repeated minimal repairs on the component level causing system failure. Let $(X_k), k \in \mathbb{N}$, be a sequence of independent identically distributed (i.i.d.) random variables following an exponential distribution with parameter 1: $X_k \sim \text{Exp}(1)$. Then we define

$$T_1 = X_1 \vee X_2, T_{n+1} = T_n + X_{n+2}, n \in \mathbb{N} \quad (15)$$

The intensity of the corresponding counting process N is then $\lambda_t = I(X_1 \wedge X_2 \leq t)$. By elementary calculations it can be seen that

$$E[I(T_1 > t) | \mathcal{F}_t^\lambda] = P(T_1 > t | X_1 \wedge X_2 \wedge t) \quad (16)$$

is continuous and nonincreasing. According to Theorem 1 it follows that (T_n) is an MRP and that the first time to failure has conditional distribution

$$\begin{aligned} 1 - F_t &= \exp \left\{ - \int_0^t I(X_1 \wedge X_2 \leq s) ds \right\} \\ &= \exp \{ -(t - X_1 \wedge X_2)^+ \} \end{aligned} \quad (17)$$

An Application to Optimal Replacement

Consider the setup of the section titled ‘‘A General Minimal Repair Model’’. Suppose now a planned replacement of the system is scheduled at time T , which may depend on the condition of the system, that is on the process X_t . The replacement time T is a *stopping time* in the sense that the event $\{T \leq s\}$ depends on the process X_t up to time s . There is no planned replacement if $T = \infty$.

The following simple cost structure is assumed: a planned replacement of the system costs $K (> 0)$ and a repair/replacement at system failure costs $c (> 0)$.

It is assumed that the systems generated by replacements are stochastically independent and identical, the same replacement policy is used for each system and the replacement and repairs take negligible time.

The problem is to determine a replacement time minimizing the long-run (expected) cost per unit time.

Let M^T and S^T denote the expected cost associated with a replacement cycle and the expected length of a replacement cycle, respectively. We restrict attention to T 's having $M^T < \infty$ and $S^T < \infty$. Then using Theorem 3.16 of [27], the long-run (expected) cost per unit time can be written as follows:

$$B^T = \frac{M^T}{S^T} = \frac{cEN_T + K}{ET} \quad (18)$$

Using equation (12) (see also [7, 28]), it follows from (18) that

$$B^T = \frac{cE \int_0^T \lambda_t dt + K}{E \int_0^T dt} \quad (19)$$

We note that the optimality criterion is in the same form as analyzed by Aven and Bergman [29]. The main results obtained in [29] are summarized below.

Introduce $a_t = c\lambda_t$ and assume that a is nondecreasing in t . Define the replacement time T_δ by the first point in time the process a_t exceeds δ , that is $a_t \geq \delta$. We assume $ET_\delta < \infty$. It can be seen that T_δ minimizes

$$M^T - \delta S^T = E \int_0^T [c\lambda_t - \delta] dt + K \quad (20)$$

The results of [29] follow. Let $B(\delta) = B^{T_\delta}$.

The stopping time T_{δ^*} , where $\delta^* = \inf_T B^T$, minimizes B^T . The value δ^* is given as the unique solution of the equation $\delta = B(\delta)$. Moreover, if $\delta > \delta^*$, then $\delta > B(\delta)$, if $\delta < \delta^*$, then $\delta < B(\delta)$; $B(\delta)$ is non-increasing for $\delta \leq \delta^*$, nondecreasing for $\delta \geq \delta^*$, and $B(\delta)$ is left-continuous.

It follows from equation (19) and Fubini's theorem that

$$B(\delta) = \frac{c \int_0^\infty E[I(a_t < \delta)\lambda_t] dt + K}{\int_0^\infty EI(a_t < \delta) dt} \quad (21)$$

Hence if

$$a_t = cv(X_t) \quad (22)$$

where $v(x)$ is a deterministic function in x , and $Q_t(\cdot)$ is the distribution of X_t , we may write

$$B(\delta) = \frac{c \int_0^\infty \int I(cv(x) < \delta)v(x)Q_t(dx) dt + K}{\int_0^\infty \int I(cv(x) < \delta)Q_t(dx) dt} \quad (23)$$

Note that if X_t is a vector process, then one of the components of X_t may be the time t .

Examples of applications can be derived from the models presented in Examples 2 and 3 above (Aven [30]). For example, for the shock model the $v(x)$ function is given by

$$v(x) = v \int_0^{\infty} p(x+y) dH(y) \quad (24)$$

References

- [1] EQR 103: Stationary replacement strategies.
- [2] Popova, E. & Popova, I. (2007). Replacement strategies, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons, Chichester.
- [3] Hollander, M., Samaniego, F.J. & Sethuraman, J. (2007). Imperfect repair, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons, Chichester.
- [4] Barlow, R. & Hunter, L. (1960). Optimum preventive maintenance policies, *Operations Research* **8**, 90–100.
- [5] Aven, T. (1983). Optimal replacement under a minimal repair strategy - a general failure model, *Advances in Applied Probability* **15**, 198–211.
- [6] Aven, T. (1987). A counting process approach to replacement models, *Optimization* **18**, 285–296.
- [7] Aven, T. & Jensen, U. (1999). *Stochastic Models of Reliability*, Springer, New York.
- [8] Aven, T. & Jensen, U. (2000). A general minimal repair model, *Journal of Applied Probability* **37**, 187–197.
- [9] Phelps, R. (1983). Optimal policy for minimal repair, *Journal of Operational Research* **34**, 425–427.
- [10] Bergman, B. (1985). On reliability theory and its applications, *Scandinavian Journal of Statistics* **12**, 1–41.
- [11] Block, H., Borges, W. & Savits, T. (1985). Age-dependent minimal repair, *Journal of Applied Probability* **22**, 370–385.
- [12] Stadje, W. & Zuckerman, D. (1991). Optimal maintenance strategies for repairable systems with general degree of repair, *Journal of Applied Probability* **28**, 384–396.
- [13] Shaked, M. & Shanthikumar, G. (1986). Multivariate imperfect repair, *Operational Research* **34**, 437–448.
- [14] Beichelt, F. (1993). A unifying treatment of replacement policies with minimal repair, *Naval Research Logistics Quarterly* **40**, 51–67.
- [15] Zhang, F. & Jardine, A.K. (1998). Optimal maintenance models with minimal repair, periodic overhaul and complete renewal, *IIE Transactions* **30**, 1109–1119.
- [16] Finkelstein, M.S. (2004). Minimal repair in heterogeneous populations, *Journal Of Applied Probability* **41**, 281–286.
- [17] Zequeira, R.I. & Berenguer, C. (2006). Periodic imperfect preventive maintenance with two categories of competing failure modes, *Reliability Engineering and System Safety* **91**, 460–468.
- [18] Mazzuchi, T.A. & Soyer, R. (1996). A Bayesian perspective on some replacement strategies, *Reliability Engineering and System Safety* **51**, 295–303.
- [19] EQR 363: Intensity functions for Nonhomogeneous Poisson processes.
- [20] Boland, P.J. & El-Newehi, E. (1998). Statistical and information based (physical) minimal repair of n systems, *Journal of Applied Probability* **35**, 731–740.
- [21] Aven, T. (1996). Condition based replacement times - a counting process approach. Reliability engineering and system safety, *Special issue on Maintenance and Reliability* **51**, 275–292.
- [22] Arjas, E. & Norros, I. (1989). Change of life distribution via hazard transformation: an inequality with application to minimal repair, *Mathematics of Operations Research* **14**, 355–361.
- [23] Natvig, B. (1990). On information-based minimal repair and the reduction in remaining system lifetime due to the failure of a specific module, *Journal of Applied Probability* **27**, 365–375.
- [24] EQR 115: Optimal replacement in shock models.
- [25] EQR 340: Coherent systems.
- [26] Barlow, R.E. & Proschan, F. (1975). *Statistical Theory of Reliability and Life Testing*, Holt, Rinehart Winston, New York.
- [27] Ross, S.M. (1970). *Applied Probability Models with Optimization Applications*, Holden-Day, San Francisco.
- [28] Brémaud, P. (1981). *Point Processes and Queues. Martingale Dynamics*, Springer, New York.
- [29] Aven, T. & Bergman, B. (1986). Optimal replacement time, a general set-up, *Journal of Applied Probability* **23**, 432–442.
- [30] Aven, T. (1996). Optimal replacement of monotone repairable systems. Chapter in lecture notes, *Reliability and Maintenance of Complex Systems*, Springer-Verlag, NATO ASI, New York.

Related Articles

Imperfect Repair; Nonparametric Methods for Analysis of Repair Data; Multivariate Imperfect Repair Models; Age-Dependent Minimal Repair and Maintenance; Repair Data, Sets of: How to Graph, Analyze, and Compare; Dynamic Programming Methods in Repair and Replacement.

TERJE AVEN

Analysis of Recurrent Events from Repairable Systems

Recurrent Events

In many disciplines, the event of interest occurs on a recurring basis. For example, in the engineering, quality, and reliability setting, recurrent events include the failure of a mechanical or electrical component, the discovery of bugs in a computer software, the appearance of cracks in a concrete edifice, the occurrence of industrial accidents in a plant, or the breakdown of fibers in a fiber bundle. In the biomedical and public health setting, examples are the occurrence of epileptic seizures, hospitalization of patients with chronic diseases, and emergence of tumors. In the social and political sciences, civil disturbances (e.g., race riots, strikes by union members), international conflicts, or brushes with the law by a juvenile delinquent are recurrent events. More generally, monitoring the status of a repairable system leads to a recurrent event. Ascher and Feingold [1] defined a repairable system as a system which, after failing to perform one or more of its functions satisfactorily, can be restored to fully satisfactory performance by any method other than the replacement of the entire system. Examples of real recurrent event data sets in the literature are the air-conditioning data examined in Proschan [2], the gastroenterology data in Aalen and Husebye [3], and the data in Table 2.7 in Blischke and Murthy [4], which is depicted in Figure 1 which provides the times between successive failures of hydraulic components of load-haul-dump (LHD) machines that were used in moving ore and rock in underground mines in Sweden.

Because the recurrent event can potentially occur more than once within an experimental unit, it is important to take into account associations among the interevent times. Realistically, the distribution of an interevent time will also depend on the history up to the time of the last event occurrence. For example, knowing when a mechanical component has failed thus far, and the number of observed failures, will affect the time to the next failure of the component. Therefore, the usual assumption of independent and identically distributed (i.i.d.)

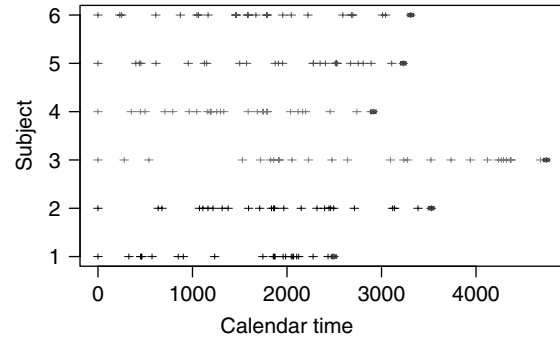


Figure 1 A pictorial depiction of the load-haul-dump (LHD) machine recurrent event data set from Table 2.7 of Blischke and Murthy [4]

interevent times will typically be untenable. Furthermore, interventions (e.g., corrective measures) will usually be performed at each event occurrence. Going back to the mechanical component, upon its breakdown, the intervention will either come in the form of repair or a replacement of this component. The type of intervention performed will therefore partially affect the distribution of the time to the next breakdown of the component. When modeling recurrent events, it becomes imperative that these aforementioned factors be taken into account to arrive at so-called dynamic models.

Software reliability provides a poignant setting for such dynamic recurrent event models. The simplest and classic model in this area is the Jelinski and Moranda (JM) [5] model which postulates that the distribution of the interfailure time of a software depends on the number of bugs remaining in the software and an unknown parameter encoding the contribution of each remaining bug to the interfailure time. The mathematical description of this model will be described later as it is a special case of a general dynamic recurrent event model. The JM model was extended in Agustin and Peña [6] where a series system of m component software's is considered, so the failure of a component software leads to the failure of the software system. In the Agustin and Peña [6] paper, the goal is to perform a debugging of the software system to develop a deployable system which can successfully complete, at a prespecified reliability or success probability, a given task whose completion time may be random. In the debugging stage of this software system during its development, there

are $m + 1$ types of recurrent events: the m failures attributable to each of the component software, and the completion of the assigned task. The JM model and the extension in [6] are examples of dynamic recurrent event models. They go beyond the simplistic i.i.d. model for interevent times.

In this article, we provide a review of dynamic models for recurrent events, and also discuss statistical inference procedures for these dynamic models. Inference procedures will be those pertaining to parameter **estimation**, as well as those that deal with goodness-of-fit (GOF) issues for the class of models. The section titled “Models for Recurrent Events” will discuss recurrent event models, while the section titled “Inference Procedures” will deal with the inferential aspects. A demonstration of the procedures will be performed using the LHD data set. Concluding remarks will be provided in the last section.

Models for Recurrent Events

In this article, we impose the simplifying assumption that repairs take negligible time and hence our focus is just on the operational time. Consider a component placed on test at (calendar) time 0. Let $0 = S_0 < S_1 < S_2 < \dots$ denote the successive times of event occurrences and let $T_i = S_i - S_{i-1}$, $i = 1, 2, \dots$, denote the interevent times. The history of the process will be denoted by $\mathbf{F} = \{\mathcal{F}_s : s \geq 0\}$, which is an increasing, right-continuous collection of σ fields. This history is also called a *filtration*, and $\mathcal{F}_s$ contains information about the component which includes the number of recurrent event occurrences up to time s , **covariate** information, and the type of intervention that were performed. Owing to the dynamic nature in which recurrent events are monitored and the need to cogently incorporate the effects of performed interventions, it is more natural and beneficial to use hazard rate functions instead of density functions in modeling. Recall that for a continuous interevent time $T \geq 0$ with density function $f(t)$ and distribution function $F(t)$, its hazard rate function and cumulative hazard function are defined, respectively, via

$$\lambda(t) = f(t)/[1 - F(t)] \quad \text{and} \quad \Lambda(t) = -\log[1 - F(t)] = \int_0^t \lambda(v) dv \quad (1)$$

The hazard rate function has the intuitive interpretation that for small positive h , $\lambda(t)h$ is approximately the probability that $t \leq T \leq t + h$, given the information that $T \geq t$. In addition to their ability to describe qualitative aspects of a failure process over time, hazard rate functions enable the use of the modern stochastic process approach to event-time analysis, which is the mathematical framework adopted in this article. This framework is anchored in the theory of counting processes, martingales, and stochastic integration as seminally introduced by Aalen [7]. Excellent presentation of this framework can be found in Andersen *et al.* [8] and Fleming and Harrington [9].

The number of occurrences of the recurrent event is represented by the stochastic process $\{N(s) : s \geq 0\}$, where $N(s) = \sum_{i=1}^{\infty} I\{S_i \leq s\}$. On the other hand, the process $\{A(s) : s \geq 0\}$ is a nondecreasing process measurable with respect to $\mathcal{F}_{s-}$ such that $\{M(s) = N(s) - A(s), s \geq 0\}$ is a square-integrable zero-mean martingale. The process $A(s)$ is called the *compensator* of N . It is usually assumed to have the form

$$A(s) = \int_0^s \lambda(w) dw = \int_0^s Y(w)\alpha(w) dw \quad (2)$$

where $\{\alpha(s) : s \geq 0\}$ is a measurable process with respect to $\mathcal{F}_{s-}$ and $\{Y(s) : s \geq 0\}$ is an at-risk process with $Y(s) = 1$ if the component is still under observation at time s and $Y(s) = 0$ otherwise. The process $\lambda(s) = Y(s)\alpha(s)$ is referred to as the *intensity rate process*. The process $\alpha(s)$ may depend on some baseline hazard rate function.

A heuristic way to view the above model is as follows. Consider an infinitesimal interval $[s, s + ds)$ and suppose we are given all the information that has accrued just before time s , formally encapsulated by the σ field $\mathcal{F}_{s-}$. Then the change in N over this interval, denoted by $dN(s) = N((s + ds)-) - N(s-)$, could be viewed as a Bernoulli random variable with success probability of $dA(s) = Y(s)\alpha(s) ds$. As such,

$$E\{dN(s)|\mathcal{F}_{s-}\} = Y(s)\alpha(s) ds \quad \text{and} \quad (3)$$

$$\begin{aligned} \text{Var}\{dN(s)|\mathcal{F}_{s-}\} &= (Y(s)\alpha(s) ds)(1 - Y(s)\alpha(s) ds) \\ &\approx Y(s)\alpha(s) ds \end{aligned} \quad (4)$$

One of the early dynamic models for **repairable systems** is the imperfect repair model of Brown and

Proschan [10]. In this model, a component that fails at time s is repaired either perfectly with probability p , or minimally with probability $1 - p$, where $p \in [0, 1]$. A perfect repair restores the component to the good-as-new state, whereas a **minimal repair** restores the component to its state just prior to failure. Hence, the time to the next failure of a minimally repaired component is stochastically equivalent to that of a working component of the same age. If we take the extreme case where $p = 1$, which is akin to replacing the failed component with a new one, then the interevent times T_i 's are i.i.d. and therefore a renewal process can be used to model the interevent times. On the other hand, if $p = 0$, then a time-transformed nonhomogeneous **Poisson process** is the model of choice. This **imperfect repair** model was further extended by Block, Borges, and Savits [11] to the so-called BBS model by letting the probability of perfect repair p to be time dependent. Hence, upon failure at time s , a component is repaired perfectly with probability $p(s)$, or minimally with probability $1 - p(s)$. For both models, the effective age of a component just after a repair will be one of two possible values: (a) zero, if the component was repaired perfectly or (b) identical to the effective age just prior to failure, if minimal repair was performed.

To accommodate modes of repair other than the two extremes of minimal and perfect, Dorado *et al.* [12] introduced a general repair model where the effective age of a component after repair can possibly be less than its effective age just prior to failure. To describe their model, for a component, define sequences $\{A_j : j = 1, 2, \dots\}$ and $\{\phi_j : j = 1, 2, \dots\}$, which are referred to as the *effective ages* and *life supplements*, respectively, and such that $A_1 = 0$, $\phi_1 = 1$, $A_j \geq 0$, $\phi_j \in (0, 1]$ with

$$A_j \leq A_{j-1} + \phi_{j-1}T_{j-1}, j = 1, 2, \dots \quad (5)$$

Defining the effective age just prior to the j th failure as $X_j = A_j + \phi_j T_j$ enables us to write equation (5) as $A_j \leq X_{j-1}$. It is clear that smaller values of ϕ_j result in smaller effective ages, hence the term life supplement. By choosing A_j and ϕ_j appropriately, other known models for repairable systems such as the Brown and Proschan model as well as the models proposed by Kijima [13] can be obtained. For instance, if we let $\phi_j = 1$ for each j and

$$A_j = \sum_{k=1}^{j-1} \left[\prod_{i=k}^{j-1} D_i \right] T_k = D_{j-1}(A_{j-1} + T_{j-1}), \quad j = 1, 2, \dots \quad (6)$$

where D_j is a random variable such that $D_j = 0$ with probability p and $D_j = 1$ with probability $1 - p$, then we get the Brown and Proschan [10] model.

Recognizing that it is not only the type of intervention that affects the distribution of successive interevent times, Peña and Hollander [14] developed a more general dynamic recurrent event model. In particular, they noted that other concomitant variables may have an effect, together with unobservable variables which may induce dependencies among the interevent times. Their model uses an intensity process that depends on the interaction of four factors, namely, the effect of the interventions undertaken, the number of event occurrences, the presence of covariates, and the impact of unobservable frailty components. Due to the interplay of these various factors, it is not necessary for a repaired component to be better than its state just prior to failure, that is, a repair may put the component in a worse state relative to the state just before the most recent failure. An example of this is imperfect debugging in software reliability where more bugs may be introduced in the process of removing a bug.

The model by Peña and Hollander postulates an intensity rate process for a component with a covariate vector $\mathbf{X} = (X_1, \dots, X_m)^t$ and a positive frailty component Z of form

$$\lambda(s|Z, \mathbf{X}) = Y(s)Z\lambda_0(\mathcal{E}(s))\rho(N(s-); \alpha)\psi(\boldsymbol{\beta}^t\mathbf{X}) \quad (7)$$

In this model, Z has distribution $H(z; \xi)$, $\lambda_0(\cdot)$ is an unknown baseline hazard function, $\rho(\cdot; \cdot)$ is a monotone function from $\mathbf{N} = \{0, 1, 2, \dots\}$ to $\mathbf{R}_+ = [0, \infty)$ with $\rho(0; \alpha) = 1$ and α is some unknown parameter, $\psi(\cdot)$ is a known nonnegative link function from $\mathbf{R}$ to $\mathbf{R}_+$, $\boldsymbol{\beta} = (\beta_1, \dots, \beta_m)^t$ is an unknown regression coefficient vector, and $\mathcal{E}(s)$ is an observable $\mathcal{F}_{s-}$ -measurable process and satisfies (a) $\mathcal{E}(0) = e_0$, where e_0 is a nonnegative real number; (b) $\mathcal{E}(s) \geq 0$; and (c) $\mathcal{E}(s)$ is monotone and differentiable for $s \in [S_{k-1}, S_k)$ with $\mathcal{E}'(s) \in (0, 1]$. The process $\mathcal{E}(s)$ is referred to as the *effective age process* and models the

impact of the type of intervention performed at each event, whereas the function $\rho(\cdot; \cdot)$ contains the effect of the accumulating event occurrences. A value of $\rho(\cdot; \cdot)$ greater than unity indicates that the accumulation of event occurrences results in accelerating future event occurrences. On the other hand, $\rho(\cdot; \cdot) < 1$ implies a deceleration of future event occurrences. In most situations, $\rho(\cdot; \alpha)$ is a monotone function.

To illustrate the generality of this model, we present some examples to show that with appropriate choices for $\mathcal{E}(\cdot)$ and $\rho(\cdot; \cdot)$ in equation (7), some of the well-known models are obtained. Note that in these examples we do not have the frailty component, that is, $Z = 1$.

Example 1 If $\mathcal{E}(s) = s - S_{N(s-)}$, the backward recurrence time, and $\rho(k; \alpha) = \alpha - k$ for some positive integer parameter α , the JM software reliability model is recovered from the Peña and Hollander model. The parameter α represents the initial number of bugs in the software at the beginning of the debugging process. An additional feature is the presence of a covariate which, in practice, could depend on a number of factors such as the ability of the individual performing the debugging of the software.

Example 2 Let $D_1, D_2, \dots$ be an i.i.d. sequence of Bernoulli random variables with success probability p . Let $v(s) = \sum_{i=1}^{N(s)} D_i$ be the number of “successes” observed up to time s and define $0 \equiv \tau_0 < \tau_1 < \tau_2 < \dots$ via $\tau_k = \min\{j > \tau_{k-1} : D_j = 1\}$, $k = 1, 2, \dots$. By setting $\rho(k; \alpha) \equiv 1$ and $\mathcal{E}(s) = s - S_{\tau_{v(s-)-}}$, the Brown and Proschan minimal repair model is obtained. Recall that a “success” is defined as a perfect repair. In addition, setting the success probability p as a function of the calendar time s results in the BBS model.

Example 3 Let $\{A_j : j = 1, 2, \dots\}$ and $\{\phi_j : j = 1, 2, \dots\}$ be two sequences satisfying equation (5). If we let $\rho(k; \alpha) \equiv 1$ and $\mathcal{E}(s) = A_{N(s-)} + \phi_{N(s-)}(s - S_{N(s-)})$ then the model reduces to the general repair model of Dorado, Hollander, and Sethuraman, with the additional feature of incorporating covariates.

For more detailed discussions of the generality of the Peña and Hollander dynamic model, refer to their paper [14] and Peña *et al.* [15].

Inference Procedures

In this section, we discuss two important inferential aspects regarding dynamic recurrent event models. We focus on model parameter estimation for the Peña and Hollander model, whereas for the GOF testing problem, we use the BBS minimal repair model. The GOF testing problem for the more general dynamic model is currently being researched.

Estimation of Parameters

The approach to estimating model parameters is *via* the maximum-likelihood principle. This entails the construction of an appropriate likelihood function, given sample data. We suppose that n independent experimental units are observed over the interval $[0, s^*]$ yielding the sample data

$$\mathbf{D}_i(s^*) = \{(N_i(v), Y_i(v), \mathcal{E}_i(v), \mathbf{X}_i(v)), v \in [0, s^*]\}, \\ i = 1, 2, \dots, n \quad (8)$$

For the i th unit there is also an unobservable frailty variable Z_i , where $Z_1, Z_2, \dots, Z_n$ are assumed to be i.i.d. from a gamma distribution $H(\cdot|\xi)$ with unit mean and variance $1/\xi$, where ξ is unknown. The case where ξ is made to approach ∞ coincides with the model without frailty. The model parameters are α, β, ξ , and $\lambda_0(\cdot)$. For brevity, we let $\delta = (\lambda_0(\cdot), \alpha, \beta)$.

Given $Z_i = z_i$, and invoking the heuristic conditionally Bernoulli interpretation of the model, with $B_i(v; \delta) = Y_i(v) \lambda_0[\mathcal{E}_i(v)] \rho[N_i(v-); \alpha] \psi(\mathbf{X}_i(v)^t \beta)$ the likelihood contribution of the i th unit is given by

$$L_{C,i}(\delta; s^*) \\ = \prod_{v=0}^{s^*} [z_i B_i(v; \delta)]^{\Delta N_i(v)} [1 - z_i B_i(v; \delta)]^{1 - \Delta N_i(v)} \\ = \left\{ \prod_{v=0}^{s^*} [z_i B_i(v; \delta)]^{\Delta N_i(v)} \right\} \exp \left\{ -z_i \int_0^{s^*} B_i(v; \delta) dv \right\} \quad (9)$$

The likelihood function, conditional on $Z_i = z_i$, $i = 1, 2, \dots, n$, becomes $L_C(\delta; s^*) = \prod_{i=1}^n L_{C,i}(\delta; s^*)$. Multiplying this conditional likelihood by $\prod_{i=1}^n h(z_i|\xi)$, where $h(z|\xi) = \xi^\xi z^{\xi-1} \exp(-\xi z) / \Gamma(\xi)$ for $z \geq 0$ is the gamma density function associated with $H(z|\xi)$, then integrating over $z_i \in [0, \infty)$, $i =$

$1, 2, \dots, n$, yields the unconditional likelihood function given by

$$L(\delta, \xi; s^*) = \prod_{i=1}^n \left\{ \left(\frac{\xi}{\xi + \int_0^{s^*} B_i(v; \delta) dv} \right)^\xi \times \prod_{v=0}^{s^*} \left[\frac{(N_i(v-) + \xi) B_i(v; \delta)}{\xi + \int_0^{s^*} B_i(v; \delta) dv} \right]^{\Delta N_i(v)} \right\} \quad (10)$$

This likelihood simplifies under the no-failty setting by letting $\xi \rightarrow \infty$ or setting $z_i = 1, i = 1, 2, \dots, n$, to

$$L(\delta; s^*) = \prod_{i=1}^n \left[\left\{ \prod_{v=0}^{s^*} [B_i(v; \delta)]^{\Delta N_i(v)} \right\} \times \exp \left\{ - \int_0^{s^*} B_i(v; \delta) dv \right\} \right] \quad (11)$$

Two cases need to be considered in the estimation procedure. The first case is when the unknown baseline hazard rate function $\lambda_0(\cdot)$ is assumed to belong to some known parametric family of hazard rate functions $\mathcal{C} = \{\lambda(t|\theta) : \theta \in \Theta\}$ with the parameter θ unknown, such as for example when $\mathcal{C}$ is the Weibull class of hazard rate functions with members $\lambda(t|\theta_1, \theta_2) = (\theta_1 t)^{\theta_2}$ for positive parameters θ_1 and θ_2 . The model parameters to be estimated are therefore $\delta = (\theta, \alpha, \beta)$ and ξ . The maximum likelihood ML estimates $(\hat{\delta}, \hat{\xi})$ of these parameters could be obtained by direct maximization of the likelihood function $L(\delta, \xi; s^*)$ in equation (10) using iterative methods, or an expectation-maximization (EM) (*cf.*, Dempster *et al.* [16]) approach could be employed as demonstrated in Stocker and Peña [17], to which we refer the reader for details. In the EM implementation, the frailty values z_i 's are considered as missing.

In the second case, the baseline hazard rate function $\lambda_0(\cdot)$ is nonparametrically specified, so it is only assumed to be some continuous hazard rate function on $\mathfrak{N}_+ = [0, \infty)$. The model parameter therefore has an infinite-dimensional component, the $\lambda_0(\cdot)$, and finite-dimensional components, which are α, β , and ξ , hence this is referred to as a *semiparametric model*. The estimation algorithm in this semiparametric setting was presented in Peña *et al.* [15]; *see also* Kvam and Peña [18], which deals with a load-sharing dynamic semiparametric model which

is a special case of the general dynamic model. Gonzalez *et al.* [19] developed an R package (Ihaka and Gentleman [20]) called `gcmrec` which implements the algorithm in [15] for estimating the baseline cumulative hazard function $\Lambda_0(t)$ and the finite-dimensional parameters (α, β, ξ) . Properties of the resulting estimators were also examined in [15], to which we refer the reader for details of the semiparametric estimation procedure. We do not go further into details concerning the estimators in this setting, except to mention that the estimator of $\Lambda_0(t)$ is a generalization of the Aalen–Breslow–Nelson estimator of the cumulative hazard function with right-censored data. The resulting estimator of the distribution function associated with $\Lambda_0(\cdot)$ is a generalization of the Kaplan–Meier [21] or the product-limit estimator of the distribution function with right-censored data. The **EM algorithm** plays a central role in the estimation of the finite-dimensional parameters in this semiparametric setting, with the estimates obtained from a partial likelihood function which generalizes the partial likelihood under the **Cox proportional hazards** model in Cox [22] and Andersen and Gill [23].

Example 4 Let us now consider the LHD machine data set in Blischke and Murthy [4], which was first presented and analyzed in Kumar and Klefsjö [24]. Data were collected throughout the 2-year development phase from six machines of varying ages labeled as old, medium, and relatively new. We treat age as a covariate and use two dummy variables X_1 and X_2 to code the age classification as follows: old machines are coded as (0, 0), medium-aged machines as (1, 0), and relatively new machines as (0, 1). We assume that the effective age process, the accumulating event function, and the link function are of the forms $\mathcal{E}(s) = s - S_{N(s)}, \rho[N(s-); \alpha] = \alpha^{N(s-)},$ and $\psi(y) = \exp(y)$, respectively. Note that the given effective age process implies that a perfect repair is performed at each failure, while the given link function corresponds to that in the Cox proportional hazards model [22]. The relevant parameters were estimated by Stocker and Peña [17] under the assumption of a two-parameter Weibull model for the baseline hazard function, while Peña *et al.* [15] considered a nonparametric specification for the baseline hazard function. The resulting estimates of the finite-dimensional parameters are presented in Table 1. In both parametric and semiparametric scenarios, $\hat{\xi}$

6 Analysis of Recurrent Events from Repairable Systems

Table 1 Parameter estimates for the LHD data set under both parametric and nonparametric specifications of the baseline hazard function

Parameter of interest	Estimates	
	$\Lambda_0(t) = (\theta_1 t)^{\theta_2}$	Nonparametric specification for $\Lambda_0(t)$
α	1.030	1.026
β_1	-0.138	-0.076
β_2	-0.079	-0.054
ξ	16 419.81	2.53×10^{28}
θ_1	0.006	NA
θ_2	0.969	NA

turned out to be very large, thereby indicating that there is no need for a frailty component in the model.

Goodness-of-Fit Testing

Another facet of inference is **hypothesis testing**. GOF testing is one form of hypothesis testing. We examine this GOF issue in the recurrent event setting, but for the more restricted BBS model incorporating covariates. The problem of interest is to ascertain, based on sample data, whether the baseline hazard rate function $\lambda_0(\cdot)$ belongs to some specified parametric family of hazard rate functions. A solution to this GOF problem for the BBS model was presented by Agustin and Peña [25, 26] and Agustin [27]. These papers are the basis of the procedure described below, which uses the smooth embedding idea of Neyman [28] but adapted to hazard rate functions. See Rayner and Best [29] for the classical smooth GOF theory and Peña [30, 31] for its hazard-based adaptation.

Consider the situation where n repairable systems are monitored with the repair or intervention scheme following the BBS model with perfect repair probability function given by $p(\cdot)$. We assume that each system is followed until the first perfect repair. For the i th system, we denote by v_i the number of failure events observed until the first perfect repair, so by letting $S_{i1} < S_{i2} < \dots$ be the successive times of failures, the length of time until the first perfect repair is S_{iv_i} . The vector of random observables for the i th system is

$$\mathbf{D}_i = \left\{ v_i, (S_{ij}, \quad j = 1, 2, \dots, v_i), \right. \\ \left. (\mathbf{X}_i(v), v \in [0, S_{iv_i}]) \right\} \quad (12)$$

We convert this to a stochastic process notation by defining for the i th system the processes

$$N_i(s) = \sum_{j=1}^{v_i} I\{S_{ij} \leq s\} \quad \text{and} \quad Y_i(s) = I\{S_{iv_i} \geq s\} \quad (13)$$

so $\mathbf{D}_i$ is equivalent to observing $\{(N_i(v), Y_i(v), \mathbf{X}_i(v)) : v \in [0, \infty)\}$. Recall that under the BBS model with covariates we have, in differential notation,

$$\mathbf{P}\{dN_i(s) = 1 | \mathcal{F}_{s-}\} = Y_i(s) \lambda_0(s) \exp\{\mathbf{X}_i(s)^t \boldsymbol{\beta}\} ds \quad (14)$$

for some unknown baseline hazard rate function $\lambda_0(\cdot)$ and an unknown regression parameter vector $\boldsymbol{\beta}$. We require that v_i be finite with probability one, and a sufficient condition for this to hold is $\int_0^\infty p(v) \lambda(v; \xi) \exp\{\mathbf{X}_i(v)^t \boldsymbol{\beta}\} dv = \infty$.

The GOF testing problem is to test the null hypothesis (H_0) that the baseline hazard rate function $\lambda_0(\cdot)$ belongs to a prespecified parametric family of hazards $\mathcal{C} = \{\lambda_0(\cdot; \xi) : \xi \in \Xi\}$, where the functional form of $\lambda_0(\cdot; \cdot)$ is known, *versus* the alternative hypothesis (H_1) that $\lambda_0(\cdot)$ does not belong to $\mathcal{C}$. The idea behind Neyman's embedding approach, adapted to the hazard function framework, is to imbed $\mathcal{C}$ in a larger class of hazard rate functions. For this embedding, Agustin and Peña [25, 26], following Peña [30, 31], used the larger class given by

$$\mathcal{A}_k = \{\lambda_k(\cdot; \boldsymbol{\theta}, \xi) = \lambda_0(\cdot; \xi) \\ \times \exp\{\boldsymbol{\theta}^t \boldsymbol{\psi}(\cdot; \xi)\} : \xi \in \Xi, \boldsymbol{\theta} \in \mathfrak{R}^k\} \quad (15)$$

for some prespecified positive integer k and functions or processes $\boldsymbol{\psi}(\cdot; \xi) = (\psi_1(\cdot; \xi), \dots, \psi_k(\cdot; \xi))^t$. The class $\mathcal{C}$ then obtains from $\mathcal{A}_k$ by setting $\boldsymbol{\theta} = (\theta_1, \dots, \theta_k)^t = \mathbf{0}$, so that the GOF problem simplifies to testing the hypotheses

$$\begin{aligned} H_0 : \boldsymbol{\theta} = \mathbf{0}, \xi \in \Xi, \boldsymbol{\beta} \in \mathfrak{N}^m \quad \text{versus} \\ H_1 : \boldsymbol{\theta} \neq \mathbf{0}, \xi \in \Xi, \boldsymbol{\beta} \in \mathfrak{N}^m \end{aligned} \quad (16)$$

Because the parameters ξ and $\boldsymbol{\beta}$ are completely unspecified, these hypotheses are composite and ξ and $\boldsymbol{\beta}$ are nuisance.

Following the same conditionally Bernoulli argument utilized in the preceding subsection, under the class $\mathcal{A}_k$, the likelihood function associated with $\mathbf{D} = (\mathbf{D}_1, \dots, \mathbf{D}_n)$ is

$$\begin{aligned} L(\boldsymbol{\theta}, \xi, \boldsymbol{\beta}) = \prod_{i=1}^n \left\{ \left[\prod_{v=0}^{\infty} [B_i(v; \boldsymbol{\theta}, \xi, \boldsymbol{\beta})]^{\Delta N_i(v)} \right] \right. \\ \left. \times \exp \left\{ - \int_0^{\infty} B_i(v; \boldsymbol{\theta}, \xi, \boldsymbol{\beta}) dv \right\} \right\} \end{aligned} \quad (17)$$

where $B_i(v; \boldsymbol{\theta}, \xi, \boldsymbol{\beta}) = Y_i(v) \lambda_0(v; \xi) \exp\{\boldsymbol{\theta}^t \boldsymbol{\psi}(\cdot; \xi) + \mathbf{X}_i(v)^t \boldsymbol{\beta}\}$. The estimated score vector appropriate for the GOF problem is obtained from the likelihood function in equation (17) to be

$$\begin{aligned} \hat{\mathbf{U}} &= \frac{\partial}{\partial \boldsymbol{\theta}} \log L(\boldsymbol{\theta}, \xi, \boldsymbol{\beta})|_{(\boldsymbol{\theta}, \xi, \boldsymbol{\beta}) = (\mathbf{0}, \hat{\xi}, \hat{\boldsymbol{\beta}})} \\ &= \sum_{i=1}^n \int_0^{\infty} \boldsymbol{\psi}(v; \hat{\xi}) \left\{ dN_i(v) - B_i(v; \mathbf{0}, \hat{\xi}, \hat{\boldsymbol{\beta}}) dv \right\} \end{aligned} \quad (18)$$

In equation (18), $\hat{\xi}$ and $\hat{\boldsymbol{\beta}}$ are the estimators of ξ and $\boldsymbol{\beta}$, respectively, under the restriction that H_0 holds. The estimator $\hat{\boldsymbol{\beta}}$ is obtained from an appropriate partial likelihood function and it is the $\boldsymbol{\beta}$ solving the estimating equation

$$\sum_{i=1}^n \int_0^{\infty} [\mathbf{X}_i(v) - \mathbf{E}(v; \boldsymbol{\beta})] dN_i(v) = \mathbf{0} \quad (19)$$

In this estimating equation, using the notation $\mathbf{a}^{\otimes m}$ to represent 1, $\mathbf{a}$, $\mathbf{a}\mathbf{a}^t$ for $m = 0, 1, 2$, respectively, for a

vector $\mathbf{a}$, we define

$$\begin{aligned} S_{(m)}(v; \boldsymbol{\beta}) &= \frac{1}{n} \sum_{i=1}^n \mathbf{X}_i(v)^{\otimes m} Y_i(v) \\ &\times \exp(\mathbf{X}_i(v)^t \boldsymbol{\beta}), \quad m = 0, 1, 2 \end{aligned} \quad (20)$$

and $\mathbf{E}(v; \boldsymbol{\beta}) = S_{(1)}(v; \boldsymbol{\beta})/S_{(0)}(v; \boldsymbol{\beta})$. On the other hand, the estimator $\hat{\xi}$ of ξ solves the estimating equation

$$\begin{aligned} \sum_{i=1}^n \int_0^{\infty} \left\{ \frac{\partial}{\partial \xi} \log \lambda_0(v; \xi) \right\} \\ \times \left\{ dN_i(v) - B_i(v; \mathbf{0}, \xi, \hat{\boldsymbol{\beta}}) dv \right\} = \mathbf{0} \end{aligned} \quad (21)$$

Note that solving these two estimating equations will typically require iterative methods.

From the estimated score vector $\hat{\mathbf{U}}$ in equation (18), the GOF test statistic, based on the fixed smoothing order k , is the quadratic form

$$\hat{Q}_k = \frac{1}{n} \hat{\mathbf{U}}^t \hat{\mathbf{\Gamma}}^{-} \hat{\mathbf{U}} \quad (22)$$

where $\hat{\mathbf{\Gamma}}$ is the estimate of the limiting **covariance** matrix of $\hat{\mathbf{U}}/\sqrt{n}$ under H_0 , and for a matrix $\mathbf{M}$, $\mathbf{M}^{-}$ represents its generalized inverse. The specific expression for $\hat{\mathbf{\Gamma}}$, which is somewhat involved, can be found in Agustin and Peña [26]. The test proceeds by rejecting H_0 if $\hat{Q}$ exceeds $\chi_{k^*, \alpha}^2$, the $100(1 - \alpha)$ th percentile of a central χ^2 distribution with k^* **degrees of freedom**, where k^* is the rank of $\hat{\mathbf{\Gamma}}$. We mention that there is also the issue of properly choosing the smoothing order k , and whether this could be chosen adaptively or in a data-dependent manner. We do not dwell on this issue here, but refer instead to the paper [26].

There is a variety of choices for the vector of smoothing functions or processes $\boldsymbol{\psi}(\cdot; \xi)$. One possibility is a polynomial form given by $\boldsymbol{\psi}(v; \xi) = (\Lambda_0(v; \xi), \Lambda_0^2(v; \xi), \dots, \Lambda_0^k(v; \xi))^t$, where $\Lambda_0(\cdot; \xi)$ is the cumulative hazard function associated with $\lambda_0(\cdot; \xi)$. For time-independent covariates we obtain closed-form expressions for the score statistic. The ℓ th component of the estimated score

statistic $\hat{\mathbf{U}}$ is

$$\hat{U}_\ell = \frac{1}{\sqrt{n}} \sum_{i=1}^n \left[\sum_{j=1}^{v_i} R_{ij}^\ell - \exp \left\{ \mathbf{X}_i^\top \hat{\boldsymbol{\beta}} \right\} (R_{i v_i})^{\ell+1} / (\ell + 1) \right],$$

$$\ell = 1, 2, \dots, k \quad (23)$$

where, for $i = 1, 2, \dots, n$ and $j = 1, \dots, v_i$, we have $R_{ij} = \Lambda_0(S_{ij}; \hat{\xi})$, the model's generalized **residuals**. Some other forms $\psi(\cdot; \xi)$ are discussed in Peña [30, 31] for single-spell data and Agustin [27] for the repairable system setting. Through an appropriate choice of the $\psi(\cdot; \xi)$ function or process, generalizations of GOF tests by Akritas [32], Barlow *et al.* [33], and Moses' test (*cf.* Hollander and Proschan [34]; Gatsonis *et al.* [35]; Horowitz and Neumann [36]) are obtainable from the above framework.

Example 5 Let us revisit the LHD data set presented in Example 4. Agustin and Peña [26] analyzed this data set assuming a generalized BBS model and using the polynomial specification for $\psi(\cdot)$ mentioned above. For demonstration purposes, all the intermediate repairs are assumed to be minimal and the final repair a perfect one, which is in contrast to the assumption in the estimation example where all repairs were assumed perfect. The age of the machines is considered a covariate but the coding is slightly different with (1, 0) being associated with the old machines, (0, 1) with medium-aged ones, and (0, 0) with relatively new ones. The GOF problem entails testing if a constant baseline hazard rate function is appropriate for the data, that is, whether the baseline hazard rate is associated with an exponential distribution function. Hence, the hypotheses of interest are $H_0: \lambda_0(\cdot) \in \mathcal{C} \equiv \{\lambda_0(\cdot; \xi) = \xi : \xi \in (0, \infty)\}$ versus $H_1: \lambda_0(\cdot) \notin \mathcal{C}$. The estimates for the regression coefficients were obtained to be $(\hat{\beta}_1, \hat{\beta}_2) = (0.0541, -0.1028)$, while that for ξ is $\hat{\xi} = 0.0077$. The resulting values of the score test statistic for different choices of k , the smoothing order, together with their corresponding p values are $\hat{Q}_1 = 7.4503$ ($p = 0.0063$), $\hat{Q}_2 = 9.9551$ ($p = 0.0069$), $\hat{Q}_3 = 9.8112$ ($p = 0.0202$), $\hat{Q}_4 = 9.9291$ ($p = 0.0416$). All of the p values are smaller than

0.05; hence these results indicate that the constant baseline hazard model is not tenable in the light of the data.

Concluding Remarks

The general dynamic model for recurrent events discussed in this article shows promise in terms of being able to capture the various factors affecting recurrent event occurrences. The modeling and analysis of recurrent events remains a rich area for further research. For instance, the asymptotic properties of the estimators under the frailty model with the nonparametric specification of the baseline hazard rate function is still under investigation. In addition, validation techniques for ascertaining the appropriateness of the model need to be developed, possibly in analogy to the approach used for the BBS model. In the GOF problem, the problem of adaptively determining the smoothing order is also important. The data-driven approaches in Ledwina [37] and Kallenberg and Ledwina [38] may provide avenues for a successful resolution of this problem. A Bayesian framework may also be considered as an alternative approach to the analysis of recurrent events. In particular, Sinha [39] proposed a Bayesian structure for a proportional intensity model that incorporates a frailty component. A more recent work by Merrick *et al.* [40] presents a model that is useful for developing optimal **maintenance policies**.

Acknowledgments

The research support provided by NIH Grant 1 R01 GM56182, NIH COBRE Grant RR17698, and EPA Grant RD-83241901 is gratefully acknowledged.

References

- [1] Ascher, H. & Feingold, H. (1984). *Repairable Systems - Modeling, Inferences, Misconceptions and Their Causes*, Marcel Dekker, New York.
- [2] Proschan, F. (1963). Theoretical explanation of observed decreasing failure rate, *Technometrics* **5**, 375–383.
- [3] Aalen, O. & Husebye, E. (1991). Statistical analysis of repeated events forming renewal processes, *Statistics in Medicine* **10**, 1227–1240.
- [4] Blischke, W. & Murthy, P. (2000). *Reliability: Modeling, Prediction and Optimization*, Wiley-Interscience, New York.

- [5] Jelinski, J. & Moranda, P. (1972). *Software Reliability Research*, Academic Press, New York, pp. 465–484.
- [6] Agustin, M. & Peña, E. (1999). A dynamic competing risks model, *Probability in the Engineering and Informational Sciences* **13**, 333–358.
- [7] Aalen, O. (1978). Statistical inference for a family of counting processes, *Annals of Statistics* **6**, 701–726.
- [8] Andersen, P., Borgan, O., Gill, R. & Keiding, N. (1993). *Statistical Models Based on Counting Processes*, Springer-Verlag, New York.
- [9] Fleming, T. & Harrington, D. (1991). *Counting Processes and Survival Analysis*, John Wiley & Sons, New York.
- [10] Brown, M. & Proschan, F. (1983). Imperfect repair, *Journal of Applied Probability* **20**, 851–859.
- [11] Block, H., Borges, W. & Savits, T. (1985). Age-dependent minimal repair, *Journal of Applied Probability* **22**, 370–385.
- [12] Dorado, C., Hollander, M. & Sethuraman, J. (1997). Nonparametric estimation for a general repair model, *Annals of Statistics* **25**, 1140–1160.
- [13] Kijima, M. (1989). Some results for repairable systems with general repair, *Journal of Applied Probability* **26**, 89–102.
- [14] Peña, E. & Hollander, M. (2004). Mathematical reliability: an expository perspective, in *Models for Recurrent Events in Reliability and Survival Analysis*, R. Soyer, T. Mazzuchi & N. Singpurwalla, eds, Kluwer Academic Publishers, Chapter 6, pp. 105–123.
- [15] Peña, E., Slate, E. & Gonzalez, J. (2007). Semiparametric inference for a general class of models for recurrent events, *Journal of Statistical Planning and Inference* **137**, 1727–1747.
- [16] Dempster, A., Laird, N. & Rubin, D. (1977). Maximum likelihood from incomplete data via the EM algorithm, *Journal of the Royal Statistical Society. Series B* **39**, 1–38.
- [17] Stocker, R. & Peña, E. (2007). A general class of parametric models for recurrent event data, *Technometrics* **49**, 210–221.
- [18] Kvam, P. & Peña, E. (2005). Estimating load-sharing properties in a dynamic reliability system, *Journal of the American Statistical Association* **100**, 262–272.
- [19] Gonzalez, J., Slate, E. & Peña, E. The gcmrec package. <http://www.cran.r-project.org>: The Comprehensive R Archive Network 2005.
- [20] Ihaka, R. & Gentleman, R. (1996). R: a language for data analysis and graphics, *Journal of Computational and Graphical Statistics* **5**, 299–314.
- [21] Kaplan, E. & Meier, P. (1958). Nonparametric estimator from incomplete observations, *Journal of the American Statistical Association* **53**, 457–481.
- [22] Cox, D. (1972). Regression models and life-tables (with discussion), *Journal of the Royal Statistical Society. Series B* **34**, 187–220.
- [23] Andersen, P. & Gill, R. (1982). Cox's regression model for counting processes: a large sample study, *Annals of Statistics* **10**, 1100–1120.
- [24] Kumar, U. & Klefsjö, B. (1992). Reliability analysis of hydraulic systems of LHD machines using the power law process model, *Reliability Engineering and System Safety* **35**, 217–224.
- [25] Agustin, Z. & Peña, E. (2001). Goodness-of-fit of the distribution of time-to-first-occurrence in recurrent event models, *Lifetime Data Analysis* **7**, 289–306.
- [26] Agustin, Z. & Peña, E. (2005). A basis approach to goodness-of-fit testing in recurrent event models, *Journal of Statistical Planning and Inference* **133**, 285–303.
- [27] Agustin, Z. (2002). Generalized hazard-based goodness-of-fit tests for a repairable system, *Journal of Statistical Computation and Simulation* **72**, 519–531.
- [28] Neyman, J. (1937). “Smooth test” for goodness of fit, *Skandinavisk Aktuarietidskrift* **20**, 149–199.
- [29] Rayner, J. & Best, D. (1989). *Smooth Tests of Goodness of Fit*, Oxford University Press, New York.
- [30] Peña, E. (1998). Smooth goodness-of-fit tests for the baseline hazard in Cox's proportional hazards model, *Journal of the American Statistical Association* **93**, 673–692.
- [31] Peña, E. (1998). Smooth goodness-of-fit tests for composite hypothesis in hazard-based models, *Annals of Statistics* **26**, 1935–1971.
- [32] Akritas, M. (1988). Pearson-type goodness-of-fit tests: the univariate case, *Journal of the American Statistical Association* **83**, 222–230.
- [33] Barlow, R., Bartholomew, D., Bremner, J. & Brunk, H. (1972). *Statistical Inference Under Order Restrictions*, John Wiley & Sons, New York.
- [34] Hollander, M. & Proschan, F. (1979). Testing to determine the underlying distribution using randomly censored data, *Biometrics* **35**, 393–401.
- [35] Gatsonis, C., Hsieh, H. & Korwar, R. (1985). Simple nonparametric tests for a known standard survival based on censored data, *Communications in Statistics: Theory and Methods* **14**, 2137–2162.
- [36] Horowitz, J. & Neumann, G. (1992). A generalized moments specification test of the proportional hazards model, *Journal of the American Statistical Association* **87**, 234–240.
- [37] Ledwina, T. (1994). Data-driven version of Neyman's smooth test-of-fit, *Journal of the American Statistical Association* **89**, 1000–1005.
- [38] Kallenberg, W. & Ledwina, T. (1995). Consistency and Monte Carlo simulation of a data driven version of smooth goodness-of-fit tests, *Annals of Statistics* **23**, 1594–1608.
- [39] Sinha, D. (1993). Semiparametric Bayesian analysis of multiple event time data, *Journal of the American Statistical Association* **88**, 979–983.
- [40] Merrick, J., Soyer, R. & Mazzuchi, T. (2005). Are maintenance practices for railroad tracks effective? *Journal of the American Statistical Association* **100**, 17–25.

Further Reading

Keiding, N., Andersen, C. & Fledelius, P. (1998). The Cox regression model for claims data in non-life insurance, *Astin Bulletin* **28**, 95–118.

Related Articles

General Minimal Repair Models; Imperfect Repair; Multivariate Imperfect Repair Models; Non-

parametric Methods for Analysis of Repair Data; Reliability Growth Modeling; Replacement Strategies; Software Reliability Modeling and Analysis; Software Reliability: Bayesian Analysis; Warranty Claims and Costs: Statistical Analysis of.

MA. ZENIA N. AGUSTIN, MARCUS A.
AGUSTIN AND EDESEL A. PEÑA

Global and Dynamic Information Measures for Reliability

Introduction

Information theory is a discipline in applied mathematics that deals with a basic challenge in communication, namely quantification of data with the goal of enabling as much data as possible to be reliably stored on a medium and/or communicated over a channel. What might statisticians learn from information theory? Basic concepts such as entropy, Kullback–Leibler discrimination information, and mutual information have certainly played important roles in statistics [1, 2].

There are three information theoretic lines of research that have evolved in connection with reliability analysis. The first line of research is developing information functions that are specifically suitable for reliability analysis. The second area of research is developing various entropy-based diagnostics and tests of distributional hypothesis that are useful for reliability model building. The third information theoretic line of research, which has wide applications in reliability, is developing measures that quantify the amount of information about the immediate future contained in the past. In [3] we have updated these developments in the context of reliability analysis. In this chapter, our focus is strictly on information functions for reliability. We outline some recent developments.

Univariate Information Measures

Consider the problem of discriminating between two probability models F and G for distribution of a random variable X . Suppose that F and G have continuous densities f and g over the support $\mathcal{S}$, and f is absolutely continuous with respect to g . The mean information per observation $x \in \mathcal{S}$ from F for discrimination between F and G is

$$K(f : g) \equiv \int_{\mathcal{S}} f(x) \log \frac{f(x)}{g(x)} dx \geq 0 \quad (1)$$

The equality in equation (1) holds if and only if $f(x) = g(x)$ almost everywhere. $K(f : g)$ has many desirable properties. See Kullback [4] for details and Ebrahimi and Soofi [3] for a list of properties. An important property of $K(f : g)$ is invariance under transformations: If $Y = \phi(X)$ is a nonsingular transformation, then $K(f_Y : g_Y) = K(f_X : g_X)$, where f_Y and g_Y are densities of distributions induced by the transformation.

Shannon's entropy of a distribution F with continuous density f [5] is

$$H(f) \equiv H[f(x)] = - \int_{\mathcal{S}} f(x) \log f(x) dx \quad (2)$$

The entropy (equation 2) is not invariant under nonsingular transformations. If $Y = \phi(X)$, is a one-to-one transformation, then

$$H(Y) = H(X) - E \left(\log \left| \frac{d}{dY} \phi^{-1}(Y) \right| \right) \quad (3)$$

If we think of X as the lifetime of a new system, then $\mathcal{S} = \{x : x \geq 0\}$, and $K(f : g)$ and $H(f)$ are global information measures that do not change with the age of the system, and hence are static. Given the current age of a system, t , a set of interest for reliability analysis is the *residual lifetime*, $\mathcal{S}_t = \{x : x > t\}$. The residual life distributions are $F(x; t) = P(X - t \leq x | X > t)$, and $G(x; t) = P(X - t \leq x | X > t)$, with densities $f(x; t) = (f(x)/\bar{F}(t))$ and $g(x; t) = (g(x)/\bar{G}(t))$, where $\bar{F}(t) = 1 - F(t)$, $\bar{G}(t) = 1 - G(t)$ are the survival functions. The discrimination information function between the two residual life distributions is [6].

$$K(f : g; t) = \int_t^{\infty} f(x; t) \log \frac{f(x; t)}{g(x; t)} dx \quad (4)$$

For simplicity, we take $\bar{F}(0) = \bar{G}(0) = 1$. It is clear that $K(f : g; 0) = K(f : g)$. By equation (4), for each t , $K(f : g; t)$ possesses all the properties of $K(f : g)$. If we consider $\mathcal{S}_t$ as an index set, then $K(f : g; t)$ provides a dynamic discrimination information ranging over $\mathcal{S}_t$.

2 Global and Dynamic Information Measures for Reliability

The entropy of residual life distribution, referred to as *residual entropy*, is

$$\begin{aligned} H(X; t) &\equiv H(f; t) \\ &= - \int_t^\infty f(x; t) \log f(x; t) dx \\ &= 1 - \frac{1}{\bar{F}(t)} \int_t^\infty f(x) \log \lambda_F(x) dx \end{aligned} \quad (5)$$

where $\lambda_F(x) = (f(x)/\bar{F}(x))$ is the hazard function, [7]. Clearly $H(f; 0) = H(f)$.

Noting that the residual life distribution of the uniform distribution is also uniform over $\{x : t < x < b\}$, the residual entropy (equation 5) measures the lack of concentration (uncertainty due to lack of predictability) of the remaining lifetime at age t .

Another set of interest that leads to dynamic information measures is the past lifetime of a system, $\mathcal{S}_{[t]} = \{x : x \leq t\}$ (see Di Crescenzo and Longobardi [8, 9]). The discrimination information function between two past lifetime distributions $F(x; [t]) = P(X \leq x | X \leq t)$ and $G(x; [t]) = P(X \leq x | X \leq t)$ is

$$K(f : g; [t]) = \int_0^t f(x; [t]) \log \frac{f(x; [t])}{g(x; [t])} dx \quad (6)$$

where $f(x; [t]) = (f(x)/F(t))$ and $g(x; [t]) = (g(x)/G(t))$ denote the conditional densities.

The entropy of the past lifetime distribution is defined similarly,

$$\begin{aligned} H(X; [t]) &\equiv H(f; [t]) \\ &= - \int_0^t f(x; [t]) \log f(x; [t]) dx \end{aligned} \quad (7)$$

It is clear that for $t^* = \inf\{x : F(x) = 1\}$, $K(f : g; [t^*]) = K(f, g)$, and $H(f; [t^*]) = H(f)$.

The information decomposition rule gives

$$\begin{aligned} K(f : g) &= K(P_f : P_g; \mathcal{S}_t, \mathcal{S}_{[t]}) \\ &\quad + \bar{F}(t)K(f : g; t) + F(t)K(f : g; [t]) \end{aligned} \quad (8)$$

$$\begin{aligned} H(f) &= H(P_f; \mathcal{S}_t, \mathcal{S}_{[t]}) \\ &\quad + \bar{F}(t)H(X; t) + F(t)H(X; [t]) \end{aligned} \quad (9)$$

where

$$\begin{aligned} K(P_f : P_g; \mathcal{S}_t, \mathcal{S}_{[t]}) &= F(t) \log \frac{F(t)}{G(t)} + \bar{F}(t) \log \frac{\bar{F}(t)}{\bar{G}(t)} \end{aligned} \quad (10)$$

is the discrimination information between two Bernoulli distributions implied by F and G on the partition of $\mathcal{S}$ into $\mathcal{S}_t$ and $\mathcal{S}_{[t]}$, and

$$H(P_f; \mathcal{S}_t, \mathcal{S}_{[t]}) = -F(t) \log F(t) - \bar{F}(t) \log \bar{F}(t) \quad (11)$$

is the entropy of the Bernoulli distribution implied by F on the partitions $\mathcal{S}_t$ and $\mathcal{S}_{[t]}$.

Multivariate Information Measures

To avoid complexity, we will confine ourselves to the bivariate case. Extensions to multivariate information measures are given in Ebrahimi *et al.* [10]. Let (X_1, X_2) be a vector of nonnegative random variables. We may think of X_j , $j = 1, 2$, as the lifetimes of components of a system. For a new system, $f(x_1, x_2)$ and $g(x_1, x_2)$ are the joint probability density functions of the lifetime distributions of the components. The joint discrimination information function $K(f : g)$ and the joint entropy $H(f)$ are given by equations (1) and (2) with $\mathcal{S} = \{(x_1, x_2) : x_j \geq 0\}$, which are global measures that do not change with the age of the system, and hence are static.

When the system is not new, at ages $t_1, t_2, t_j \geq 0$ of its components, the bivariate residual lifetime is $\mathcal{S}_{t_1, t_2} = \{(x_1, x_2) : x_j > t_j \geq 0\}$. For new systems and minimal repairs it is reasonable to assume that $t_1 = t_2 = t$. The general case of (t_1, t_2) includes the equal ages as well as the case when the components are replaced with new components.

The joint residual lifetime distributions are $F(x_1, x_2; t_1, t_2) = P_F(X_1 \leq x_1, X_2 \leq x_2 | X_1 > t_1, X_2 > t_2)$, the residual density functions are $f(x_1, x_2; t_1, t_2) = (f(x_1, x_2)/\bar{F}(t_1, t_2))$ and $g(x_1, x_2; t_1, t_2) = (g(x_1, x_2)/\bar{G}(t_1, t_2))$ for $x_1 > t_1, x_2 > t_2$, where $\bar{F}(t_1, t_2) = P_F(X_1 > t_1, X_2 > t_2)$ and $\bar{G}(t_1, t_2) = P_G(X_1 > t_1,$

$X_2 > t_2$) are the joint survival functions. Here, again for simplicity, we take $\bar{F}(0, 0) = \bar{G}(0, 0) = 1$.

If F is absolutely continuous with respect to G , the discrimination information function between two bivariate residual lifetime distributions is [10]

$$K(f : g; t_1, t_2) = \int_{t_2}^{\infty} \int_{t_1}^{\infty} f(x_1, x_2; t_1, t_2) \times \log \frac{f(x_1, x_2; t_1, t_2)}{g(x_1, x_2; t_1, t_2)} dx_1 dx_2 \quad (12)$$

It is clear that $K(f : g; 0, 0) = K(f : g)$. For each (t_1, t_2) , $K(f : g; t_1, t_2)$ possesses all the properties of the discrimination information function $K(f : g)$. If we consider $\mathcal{S}_{t_1, t_2}$ as a bivariate index set, then $K(f : g; t_1, t_2)$ provides a bivariate dynamic discrimination information ranging over $\mathcal{S}_{t_1, t_2}$.

The joint residual entropy of an absolutely continuous distribution is given by

$$H(X_1, X_2; t_1, t_2) = - \int_{t_2}^{\infty} \int_{t_1}^{\infty} f(x_1, x_2; t_1, t_2) \times \log f(x_1, x_2; t_1, t_2) dx_1 dx_2 \quad (13)$$

The *marginal residual entropy* of X_1 given $X_1 > t_1, X_2 > t_2$ is

$$H(X_1; t_1, t_2) = - \int_{t_1}^{\infty} f_1(x_1; t_1, t_2) \times \log f_1(x_1; t_1, t_2) dx_1 \quad (14)$$

where $f_1(x_1; t_1, t_2) = \int_{t_2}^{\infty} f(x_1, x_2; t_1, t_2) dx_2$, $x_1 > t_1, x_2 > t_2$ is the marginal residual density of X_1 , given $X_1 > t_1, X_2 > t_2$. The marginal residual density and entropy of X_2 given $X_1 > t_1, X_2 > t_2$ are defined similarly.

The conditional residual density of X_j given that $X_j > t_j, X_i = x_i > t_i$ is

$$f_{j|i}(x_j|x_i; t_1, t_2) = \frac{f(x_1, x_2; t_1, t_2)}{f_i(x_i; t_1, t_2)}, \quad x_1 > t_1, x_2 > t_2 \quad (15)$$

Its entropy is denoted by $H[f_{j|i}(x_j|x_i; t_1, t_2)] = H(X_j|x_i; t_1, t_2)$, $i \neq j = 1, 2$, which in general is a function of x_i . The expected uncertainty about $X_j > t_j$ when we know $X_i > t_i, i \neq j$ is given by the *conditional residual entropy*,

$$H(X_j|X_i; t_1, t_2) = \int_{t_i}^{\infty} H(X_j|x_i; t_1, t_2) f_i(x_i; t_1, t_2) dx_i \quad (16)$$

The joint residual entropy can be decomposed as

$$H(X_1, X_2; t_1, t_2) = H(X_i; t_1, t_2) + H(X_j|X_i; t_1, t_2), \quad i \neq j, j = 1, 2 \quad (17)$$

An important question in statistics is to what extent the use of a variable X_i reduces uncertainty about predicting the outcomes of another variable X_j . The worth of an outcome x_i of X_i for predicting X_j , given that $X_1 > t_1, X_2 > t_2$, is assessed by comparing $f_j(x_j; t_1, t_2)$ and $f_{j|i}(x_j|x_i; t_1, t_2)$. Two information measures that may be used for this purpose are $K[f_{j|i}(x_j|x_i) : f_j(x_j); t_1, t_2]$ and the entropy difference

$$\vartheta(X_j|x_i; t_1, t_2) = H(X_j; t_1, t_2) - H(X_j|X_i; t_1, t_2) \quad (18)$$

In general, both $K[f_{j|i}(x_j|x_i) : f_j(x_j); t_1, t_2]$ and $\vartheta(X_j|x_i; t_1, t_2)$ depend on x_i . The nonnegative discrimination function $K[f_{j|i}(x_j|x_i) : f_j(x_j); t_1, t_2]$ quantifies the discrepancy between the marginal and conditional distributions, but it does not indicate which of the two distributions is more informative for the prediction of the outcome of X_j . The entropy difference (equation 18) may be positive or negative depending on which of the two distributions is more informative. Interestingly, the mean values of the two measures are the same and this value defines the mutual information between two variables:

$$\begin{aligned} M(X_1, X_2; t_1, t_2) &= \int_{t_i}^{\infty} \vartheta(X_j|x_i; t_1, t_2) f_i(x_i; t_1, t_2) dx_i \\ &= \int_{t_i}^{\infty} K[f_{j|i}(x_j|x_i) : f_j(x_j); t_1, t_2] \\ &\quad \times f_i(x_i; t_1, t_2) dx_i \end{aligned} \quad (19)$$

4 Global and Dynamic Information Measures for Reliability

The mutual information has the following Kullback–Leibler and entropy representations:

$$\begin{aligned} M(X_1, X_2; t_1, t_2) &= K [f(x_1, x_2; t_1, t_2) : f_1(x_1; t_1, t_2) f_2(x_2; t_1, t_2)] \\ &\quad (20) \end{aligned}$$

$$= H(X_1; t_1, t_2) + H(X_2; t_1, t_2) - H(X_1, X_2; t_1, t_2) \quad (21)$$

$$= H(X_j; t_1, t_2) - H(X_j | X_i; t_1, t_2), \quad i \neq j \quad (22)$$

By equation (20), $M(X_1, X_2; t_1, t_2) \geq 0$ and the residual lifetimes are independent if and only if $M(X_1, X_2; t_1, t_2) = 0$. Thus, $M(X_1, X_2; t_1, t_2)$ measures the extent of dependence between the residual lifetimes. Also, by equation (20), $M(X_1, X_2; t_1, t_2)$ is invariant under one-to-one transformations of each component. For new systems, $M(X_1, X_2; 0, 0) = M(X_1, X_2)$, which is a global measure of dependency, widely used in various fields.

At ages (t_1, t_2) , the bivariate support is partitioned as $\mathcal{S} = \mathcal{S}_{t_1, t_2} \cup \mathcal{S}_{[t_1], [t_2]} \cup \mathcal{S}_{t_1, [t_2]} \cup \mathcal{S}_{[t_1], t_2}$, where $\mathcal{S}_{[t_1], [t_2]} = \{(x_1, x_2) : x_j \leq t_j, j = 1, 2\}$ is the past lifetime and $\mathcal{S}_{t_i, [t_j]}, i \neq j = 1, 2$ is the residual lifetime of component i and past lifetime of component j . Dynamic information measures can be derived for these partitions, which lead to the dynamic decompositions of global bivariate information measures. For example, for the mutual information we have

$$\begin{aligned} M(X_1, X_2) &= M(Z_1, Z_2; t_1, t_2) \\ &\quad + \bar{F}(t_1, t_2) M(X_1, X_2; t_1, t_2) \\ &\quad + F(t_1, t_2) M(X_1, X_2; [t_1], [t_2]) \\ &\quad + P_F(\mathcal{S}_{[t_1], t_2}) M(X_1, X_2; [t_1], t_2) \\ &\quad + P_F(\mathcal{S}_{t_1, [t_2]}) M(X_1, X_2; t_1, [t_2]) \quad (23) \end{aligned}$$

where $M(Z_1, Z_2; t_1, t_2)$ is the mutual information between the components of the bivariate random variable (Z_1, Z_2) whose distribution is multinomial implied by F on the partitions of $\mathcal{S}$; and $M(X_1, X_2; t_i, [t_j]), i \neq j$ is the mutual information function that measures the dependency between the past lifetime of individual i and the remaining lifetime of individual j .

All terms in equation (23) are nonnegative. Thus, if the two lifetimes were originally independent, $M(X_1, X_2) = 0$, and the residual lifetimes as well

as the past lifetimes are independent. However, for two originally dependent lifetimes, the dynamic measures can be larger, smaller, or equal to $M(X_1, X_2)$, reflecting the situations that dependency between the residual and/or past lifetimes of components may increase, decrease, or remain unchanged as they age.

Applications

The residual entropies are useful for assessing whether the residual distributions become more/less informative as the components age. Examples of distributions with some memoryless entropies are given below. Sufficient conditions for monotonicity of $H(X_1, X_2; t_1, t_2)$ are given in Ebrahimi *et al.* [10]. Comparison of $H(X_j; t_1, t_2)$ with the univariate residual entropy $H(X_j; t_j)$ (equation 5) provides an assessment of the effect of considering both ages *vis á vis* only one age t_j on the uncertainty about the residual lifetime of the component j . The difference $H(X_j; t_j) - H(X_j; t_1, t_2)$ can be positive, negative, or zero. That is, considering both ages may decrease or increase the uncertainty about the residual lifetime of a component or leave it unchanged. A definite answer is possible under certain conditions [10].

The residual mutual information is useful for assessing expected information about the lifetime of a component from the knowledge of age of the other component as they age. For many purposes, a single global measure of dependence might be satisfactory for failure time data. But for equally many cases it is useful to consider the dependence in more detail. The dynamic mutual information $M(X_1, X_2; t_1, t_2)$ is a local measure, among other things, that can be used to address the concepts of early/late dependence and short-term and long-term dependence between the components of a system.

Residual information measures may be found in closed form for a few well-known distributions. In general, residual information measures are mathematically unwieldy. For some distributions, residual information measures must be studied using numerical integration.

Series and Parallel Components

Let $X_1, \dots, X_n$ be independent and identically distributed lifetimes of components of a system with distribution F_X . The lifetimes of series and **parallel**

systems are $Y_1 = \min(X_1, \dots, X_n)$ and $Y_n = \max(X_1, \dots, X_n)$.

Global Information Measures. We present a few global information measures based on results of Ebrahimi *et al.* [11].

If $f_X(x) \leq f_X(m) = M < \infty$, so m is the mode of the distribution, then for $i = 1, n$,

$$\begin{aligned} 1 - \log n - \frac{1}{n} - \log M &\leq H(Y_i) \\ &\leq 1 - \log n - \frac{1}{n} + (n-1) \log M + nH(X) \end{aligned} \quad (24)$$

The lower bound is useful for **series systems** and attainable for large systems under some important models for components' lifetime distribution, e.g., exponential and Pareto. The upper bound is useful for parallel systems with a few components [11].

For the lifetime distributions of two systems with identical components, one designed in series and one in parallel, we have

$$K(f_{Y_1} : f_{Y_n}) = K(f_{Y_n} : f_{Y_1}) = (n-1)[\psi(n) - \psi(1)] \quad (25)$$

$$M(Y_1, Y_n) = \log \left(1 - \frac{1}{n} \right) + \frac{1}{n-1} \quad (26)$$

where $\psi(z) = d \log \Gamma(z) / dz$ is the digamma function. Both information measures are free from the lifetime model F_X . The discrimination information function is also symmetric and grows with the size of the systems, $K(f_{Y_1} : f_{Y_n}) \geq K(f_{Y_1} : f_{Y_{n-1}}) \geq K(f_{Y_1} : f_{Y_2}) = 1$. The mutual information decreases with the size of the systems, $M(Y_1, Y_n) \leq M(Y_1, Y_{n-1}) \leq M(Y_1, Y_2) = 1 - \log 2$. That is, as the sizes of the two types of systems grow, their lifetime distributions become farther apart, and hence less dependent.

Discrepancy between the lifetime distribution of a component (a single-component system) and an n -component system may be measured by either of the following discrimination information functions:

$$K(f_X : f_{Y_i}) = n - \log n - 1, \quad i = 1, n \quad (27)$$

$$K(f_{Y_i} : f_X) = \frac{1}{n} - \log \left(\frac{1}{n} \right) - 1, \quad i = 1, n \quad (28)$$

These measures do not depend either on the type of the system or on F_X . Both are increasing in n . However, $K(f_{Y_i} : f_X) < K(f_X : f_{Y_i})$, $i = 1, n$,

and their difference grows rapidly with n . That is, on average, observation of failure of a component of a system is more informative for discrimination between the system and the component lifetime distributions than observation of the system's failure, irrespective of its type.

Dynamic Entropy. The residual entropy of lifetime of a series system with two identical components inherits monotonicity of the joint residual entropy completely. If $H(X_1, X_2; t, t) = 2H(X; t)$ is decreasing (increasing) in t , then $H(Y_1; t)$ is decreasing (increasing) in t . However, the parallel system inherits its monotonicity of the joint residual entropy partially. If $H(X_1, X_2; t, t) = 2H(X; t)$ is decreasing in t , then $H(Y_2; t)$ is decreasing in t [10].

Memoryless Information

Proportional Hazards. The dynamic Kullback-Leibler discrimination information function $K(f; g; t)$ is free from t if and only if the hazard functions are proportional, $\bar{G}(t) = (\bar{F}(t))^\rho$ [6]. An analogous result for $K(f; g; [t])$ is given by Di Crescenzo and Longobardi [9].

Univariate Exponential. The residual entropy $H(f; t)$ is free from t if and only if $f_\lambda(x) = \lambda e^{-\lambda x}$, $x \geq 0, \lambda > 0$. That is, memoryless residual entropy characterizes the exponential distribution [7].

Two exponential distributions with parameters λ_1 and λ_2 have proportional hazard functions. Thus $K(f_{\lambda_1} : f_{\lambda_2}; t) = K(f_{\lambda_1} : f_{\lambda_2})$ is free from t . However, the past life time information measures $K(f; g; [t])$ and $H(f; [t])$ are not free from t .

Bivariate Exponential. In the bivariate case, some residual information, e.g. the mutual information, may be free from the ages without others being so. Also, the entire set of bivariate residual information measures may be free from the ages only in certain direction, e.g. $t_1 = t_2 = t$, with or without the marginals being exponential. The following result is noteworthy. The conditional residual entropies $H(X_j | X_i; t_1, t_2)$, $j = 1, 2$, are constants free from t_1, t_2 if and only if X_1 and X_2 are independently exponentially distributed [10].

Bivariate Lack of Memory. A bivariate distribution is said to have the bivariate lack of memory (BLM) property, if

$$\bar{F}(s_1 + t, s_2 + t) = \bar{F}(s_1, s_2) \bar{F}(t, t) \quad (29)$$

[12, 13]. The BLM property (equation 29) implies that $H(X_j; t, t)$, $H(X_1, X_2; t, t)$, $H(X_j|X_i; t, t)$, and $M(X_1, X_2; t, t)$ are all free from t . However, $H(X_j; t)$ may be time dependent. The BLM property is sufficient, but not necessary for memoryless information, see the bivariate Pareto example below.

Pareto and Related Distributions

Univariate Pareto. Consider the Pareto distribution with survival function $\bar{F}_\alpha(x) = (x + 1)^{-\alpha}$, $x \geq 0, \alpha > 0$. Two Pareto distributions with parameters α_1 and α_2 have proportional hazard functions, so

$$K(f_{\alpha_1} : f_{\alpha_2}; t) = K(f_{\alpha_1} : f_{\alpha_2}) = \rho - \log \rho - 1 \quad (30)$$

where $\rho = (\alpha_2/\alpha_1)$ is free from t .

The global and residual entropies of this Pareto model are

$$H(X) = 1 + \frac{1}{\alpha} - \log \alpha \quad (31)$$

$$H(X; t) = H(X) + \log(1 + t) \quad (32)$$

Bivariate Pareto. Let X_1, X_2 have the bivariate Pareto distribution with joint survival function

$$\bar{F}_\alpha(x_1, x_2) = (1 + x_1 + x_2)^{-\alpha}, \quad x_1, x_2 > 0, \alpha > 0 \quad (33)$$

The joint entropy and mutual information of equation (33) are as follows:

$$H(X_1, X_2) = -\log[\alpha(\alpha + 1)] + (\alpha + 2) \left(\frac{1}{\alpha} + \frac{1}{\alpha + 1} \right) \quad (34)$$

$$M(X_1, X_2) = \log \left(\frac{\alpha + 1}{\alpha} \right) + \frac{1}{\alpha + 1} \quad (35)$$

The residual information measures of equation (33) are as follows:

- [1.] $H(X_j; t_1, t_2) = h(\alpha) + \log(1 + t_1 + t_2)$
- [2.] $H(X_j|x_i; t_1, t_2) = h(\alpha + 1) + \log(1 + t_j + x_i)$
- [3.] $H(X_j|X_i; t_1, t_2) = h(\alpha + 1) - \alpha^{-1} - (\alpha + 1)^{-1} + \log(1 + t_1 + t_2)$
- [4.] $H(X_1, X_2; t_1, t_2) = H(X_1, X_2) + 2 \log(1 + t_1 + t_2)$
- [5.] $M(X_1, X_2; t_1, t_2) = M(X_1, X_2)$

where $h(\alpha) = H(X)$ shown in equation (31).

All the entropies are increasing in t_j . Thus, as the components age, these distributions become less informative about the remaining lifetimes. From 1 and equation (32) we have

$$H(X_j; t_i) - H(X_j; t_1, t_2) = -\log \left(1 + \frac{t_i}{1 + t_j} \right), \quad i \neq j \quad (36)$$

which is negative and decreasing in t_i .

The residual mutual information is free from t_1 and t_2 . Note that the rate of increase of the joint residual entropy in t_1 and t_2 is exactly twice the rate of increase of its marginal entropies. This is an interesting finding because the level of dependency between the residual lifetimes remains unchanged, irrespective of the current ages of the components. This example also indicates that the BLM property is not necessary for the memoryless dependency.

The fact that the Pareto distribution has closed form dynamic information measures is particularly useful because the moment-based measures such as variance and **correlation** coefficient do not exist for the entire family.

Related Distributions. Numerous distributions can be obtained by component-wise one-to-one transformations of random variables X_j , $j = 1, 2$ that have Pareto distribution (see Asadi *et al.* [14]). Therefore, their information functions can be derived and studied *via* the information functions of F_α . The Kullback–Leibler information functions between two members of a family, obtained from F_{α_j} , $j = 1, 2$, are the same as equation (30), where α_j , $j = 1, 2$ are determined by the parameters of the transformed models. For example, Pareto and exponential are related by transformation $\phi(x) = \log(1 + x)$, thus $K(f_{\lambda_1} : f_{\lambda_2}) = K(f_{\alpha_1} : f_{\alpha_2})$. Shannon entropy of these distributions are related to $H(f_\alpha)$ by equation (3), where the expectation is taken with respect to f_α .

The mutual information functions of their bivariate versions are all given by equation (35).

Other Recent Developments

Consideration of age has led to some important insights about lifetime models. Applications to a few important models were shown here. Applications to several other important lifetime models, some tests of distributional hypotheses based on dynamic information measures, and some sample and Bayesian dynamic statistics have also been developed (see Ebrahimi *et al.* [10] and references therein).

A major application of dynamic information functions is derivation of lifetime models based on partial information about the hazard rate and mean residual functions. Asadi *et al.* [15] and [16] proposed maximum dynamic entropy (MDE) and minimum dynamic discrimination information (MDDI) procedures, respectively, for developing lifetime models. The MDE and MDDI procedures extend the maximum entropy and minimum discrimination information principles of inference to the cases when the information is given in terms of hazard rate growth inequality, residual moment inequality, or residual moment growth inequality. The results have led to MDE and MDDI characterizations of many well-known lifetime models and development of some new models. Dynamic information measures based on generalizations of Shannon entropy and Kullback–Leibler function have been also developed by Asadi *et al.* [17], Abraham and Sankaran [18], and Nanda and Paul [19].

References

- [1] Soofi, E.S. (1994). Capturing the intangible concept of information, *Journal of the American Statistical Association* **89**, 1243–1254.
- [2] Soofi, E.S. (2000). Principal information theoretic approaches, *Journal of the American Statistical Association* **95**, 1349–1353.
- [3] Ebrahimi, N. & Soofi, E.S. (2004). Information measures for reliability, in *Mathematical Reliability: An Expository Perspective*, R. Soyer, T.A. Mazzuchi & N.D. Singpurwalla, eds, Kluwer, The Netherlands, pp. 127–159.
- [4] Kullback, S. (1959). *Information Theory and Statistics*, John Wiley & Sons, New York, Reprinted in 1968 by Dover.
- [5] Shannon, C.E. (1948). A mathematical theory of communication, *Bell System Technical Journal* **27**, 379–423.
- [6] Ebrahimi, N. & Kirmani, S.N.U.A. (1996). A characterization of the proportional Hazards model through a measure of discrimination between two residual life distributions, *Biometrika* **83**, 233–235.
- [7] Ebrahimi, N. (1996). How to measure uncertainty in the residual lifetime distributions, *Sankhya Series A* **58**, 48–57.
- [8] Di Crescenzo, A. & Longobardi, M. (2002). Entropy-based measure of uncertainty in past lifetime distributions, *Journal of Applied Probability* **39**, 434–440.
- [9] Di Crescenzo, A. & Longobardi, M. (2004). A measure of discrimination between past life-time distributions, *Statistics and Probability Letters* **67**, 173–182.
- [10] Ebrahimi, N., Kirmani, S.N.U.A. & Soofi, E.S. (2007). Multivariate dynamic information, *Journal of Multivariate Analysis* **98**, 328–349.
- [11] Ebrahimi, N., Soofi, E.S. & Zahedi, H. (2004). Information properties of order statistics and spacings, *IEEE Transactions on Information Theory* **IT50**, 177–183.
- [12] Block, H.B. & Basu, A.P. (1974). A continuous bivariate exponential extension, *Journal of the American Statistical Association* **69**, 1031–1037.
- [13] Marshall, A.W. & Olkin, I. (1967). A multivariate exponential distribution, *Journal of the American Statistical Association* **62**, 30–44.
- [14] Asadi, M., Ebrahimi, N., Hamedani, G.G. & Soofi, E.S. (2006). Information measures for Pareto distributions and order statistics, in *Advances on Distribution Theory and Order Statistics*, N. Balakrishnan, E. Castillo & J.M. Sarabia, eds, Birkhauser, Boston, pp. 207–223.
- [15] Asadi, M., Ebrahimi, N., Hamedani, G.G. & Soofi, E.S. (2004). Maximum dynamic entropy models, *Journal of Applied Probability* **41**, 379–390.
- [16] Asadi, M., Ebrahimi, N., Hamedani, G.G. & Soofi, E.S. (2005a). Dynamic minimum discrimination information models, *Journal of Applied Probability* **42**, 643–660.
- [17] Asadi, M., Ebrahimi, N. & Soofi, E.S. (2005b). Dynamic generalized information measures, *Statistics and Probability Letters* **71**, 85–98.
- [18] Abraham, B. & Sankaran, P.G. (2006). Renyi's entropy for residual lifetime distribution, *Statistical Papers* **47**, 17–29.
- [19] Nanda, A.K. & Paul, P. (2006). Some results on generalized residual entropy, *Information Sciences* **176**, 27–47.

Related Article

Aging and Positive Dependence; Accelerated Life Models; General Minimal Repair Models ; Coherent Systems; Parallel, Series, and Series–Parallel Systems.

NADER EBRAHIMI AND EHSAN S. SOOFI

Cumulative Damage Models Based on Gamma Processes

Introduction

General systems of components typically sustain damage or **degradation** due to usage, shocks to the system, or unusual stresses applied to the system during operation. As an example, modern high-performance materials, such as fibrous composite materials, can be viewed as complex systems in which material strength is a function of the components of the material. The components of such complex systems may not be independent, and accurate modeling of system failure is more difficult than for simpler cases with independent components. In this article, modeling cumulative damage in general systems using stochastic processes is briefly discussed.

In recent papers [1–10], the degradation or accumulating damage to a system has been assumed to behave as a well-known stochastic process, such as a Gaussian process with positive drift, a geometric **Brownian motion** process, or a gamma process. Failure of the system occurs when the cumulative damage to, or degradation of, the system reaches a certain critical level. For example, a crack that begins as a small flaw in a piece of material may grow due to cyclic loads or tensile stresses applied to the material until the crack size reaches a critical value at which the material breaks. The drawback to using a Gaussian process with positive drift representing cumulative damage to a system under study is that the process is not strictly increasing, presenting a problem in interpreting the “negative” damage that such a process models. A more realistic cumulative damage/**degradation model** is the gamma process which will be discussed here in the section titled “The Gamma Process Model”. An accelerating variable is also incorporated into the model; see also [4–7, 11–14] and other references therein. **Estimation** for the model is given in the section titled “Parameter Estimation”, and an application to fatigue failure is discussed in the section titled “Example”. A discrete cumulative damage approach is mentioned in the section titled “Concluding Remarks”.

The Gamma Process Model

In this section, we derive the theoretical statistical distribution of a system lifetime based on the gamma process and provide a simpler approximate distribution for easier computation. We also briefly describe the method of incorporating an accelerating variable into the model.

Suppose that a system or material undergoes continuous cycles of stress loads until it fails or breaks down. Then accumulating damage is increasing as cyclic loads are in progress. This is the typical case for which the cumulative damage process should be always positive and strictly increasing. One such stochastic model is the gamma process, which we use to describe this cumulative damage. Let X_t represent the amount of damage accumulation for a material at cyclic load t . It is assumed that X_t is a gamma process, with the following properties:

1. $X_t - X_s$ has a gamma distribution with shape parameter $\alpha(t - s)$ and scale β ;
2. the increments $X_v - X_u$ and $X_t - X_s$ are independent with $v > u > t > s$.

For convenience, let $Y = X_t - X_s$. Then the random variable Y has the following **probability density function**:

$$g_Y(y) = \frac{1}{\Gamma(\alpha(t-s))\beta^{\alpha(t-s)}} y^{\alpha(t-s)-1} \times \exp\left(-\frac{y}{\beta}\right), \quad y > 0 \quad (1)$$

Let S be the breakdown strength of the system with the critical level C . Then S becomes the *first-passage* time or load for the gamma process to reach the critical value C . That is, system failure occurs at S cycles of loads when X_S reaches a critical value C . With starting value $X_0 = 0$, the increment $X_t - X_0$ is distributed as the gamma with shape parameter αt and scale β . Since X_t is strictly increasing in t , we have

$$\begin{aligned} P[S > t] &= P[X_t < C] \\ &= \int_0^C \frac{1}{\Gamma(\alpha t)\beta^{\alpha t}} x^{\alpha t-1} \exp\left(-\frac{x}{\beta}\right) dx \\ &= \frac{1}{\Gamma(\alpha t)} \int_0^C \xi^{\alpha t-1} e^{-\xi} d\xi \end{aligned} \quad (2)$$

2 Cumulative Damage Models Based on Gamma Processes

Thus, the cumulative distribution function (cdf) of S is given by

$$G_S(s; C) = \frac{\Gamma(\alpha s, C)}{\Gamma(\alpha s)} \quad (3)$$

where $\Gamma(a, z)$ is the incomplete gamma function defined by $\Gamma(a, z) = \int_z^\infty \xi^{a-1} e^{-\xi} d\xi$. Differentiating the above cdf, we have the probability density function (pdf) of S given by

$$g_S(s; C) = \alpha \left\{ \Psi(\alpha s) - \log C \right\} \left\{ 1 - \frac{\Gamma(\alpha s, C)}{\Gamma(\alpha s)} \right\} + \frac{\alpha}{\Gamma(\alpha s)} \frac{C^{\alpha s}}{(\alpha s)^2} \cdot \left\{ 1 + \sum_{k=1}^{\infty} \left(\frac{\alpha s}{\alpha s + k} \right)^2 \frac{(-C)^k}{k!} \right\} \quad (4)$$

For more details about the derivation of the above pdf, refer to Park and Padgett [4].

Although the pdf in equation (4) is the exact pdf of the first load at which damage accumulation passes to the critical value, this pdf is so complex that it is very difficult to use as a model in practice. A simpler approximate distribution will be used for computation instead of equation (4). This approximate distribution can be obtained by considering the discrete version of the damage accumulation. Let N denote the discrete version of the first-passage load at which damage accumulation reaches the critical level. Assume that $Y_i = X_i - X_{i-1}$ are independent and identically distributed (i.i.d.). Then Y_i 's have gamma distributions with shape α and scale β . Since the process X_n is strictly increasing, the probability that the number of cycles to first passage of damage accumulation to the critical value exceeds n is given by

$$P[N > n] = P[X_n < C] \quad (5)$$

Using $X_n = Y_n + Y_{n+1} + \dots + Y_1$, we have

$$P[N > n] = P\left[\sum_{i=1}^n Y_i < C \right] \quad (6)$$

Since $Y_i = X_i - X_{i-1}$ are iid, it follows from the central limit theorem that

$$\begin{aligned} P[N \leq n] &\cong 1 - \Phi\left(\frac{C - \alpha\beta n}{\beta\sqrt{\alpha n}}\right) \\ &= \Phi\left(\sqrt{\alpha n} - \frac{C}{\beta\sqrt{\alpha n}}\right) \end{aligned} \quad (7)$$

where $\Phi(z)$ is the cdf of the standard **normal distribution** with the pdf given by $\phi(z) = 1/\sqrt{2\pi} \cdot e^{-z^2/2}$. After reparameterizing from β to $\xi = C/\beta$, a continuous version of N , denoted by S , can be represented by

$$G_S(s) = \Phi\left(\sqrt{\alpha s} - \frac{\xi}{\sqrt{\alpha s}}\right) \quad (8)$$

To incorporate an accelerating variable, we assume that the shape coefficient α of the gamma process is dependent on an accelerating variable L and thus we denote it as α_L . There are several acceleration (link) functions between the shape coefficient and accelerating variable that can be used. These functions yield a large family of inverse-Gaussian-type models (as in [15]) with a **covariate** or accelerating variable that can be incorporated with failure data. Some popular acceleration functions are: (a) power law ($\alpha_L = \theta_0 L^{\theta_1}$), (b) Arrhenius law ($\alpha_L = \theta_0 e^{\theta_1/L}$), (c) inverse-log law ($\alpha_L = \theta_0 (\log L)^{\theta_1}$), and (d) exponential law ($\alpha_L = \theta_0 e^{\theta_1 L}$). For more details about such accelerated test models with these link functions, refer to [12, 16–17]. See also Park and Padgett [6, 14] who extend these link functions to incorporate several accelerating variables.

Parameter Estimation

Maximum likelihood estimates (MLEs) of the parameters of the proposed model can be obtained as follows. Suppose that subjects are tested at k different values of an accelerating variable $L_1, L_2, \dots, L_k$. At each value L_i , tests are performed for n_i specimens, resulting in observed failure times or observed breaking strengths s_{ij} , $j = 1, 2, \dots, n_i$ for $i = 1, 2, \dots, k$. For estimation of an accelerated test model over all of the k acceleration levels, $L_1, L_2, \dots, L_k$, the MLEs of the unknown parameters in a given model can be obtained by maximizing the log-likelihood function

$$\ell(\alpha_L(\theta_0, \theta_1), \xi) = \sum_{i=1}^k \sum_{j=1}^{n_i} \log g_S(s_{ij}) \quad (9)$$

where $g_S(\cdot)$ is the pdf of equation (8). The maximizing values of θ_0, θ_1, ξ are denoted by $\hat{\theta}_0, \hat{\theta}_1, \hat{\xi}$, and typically must be found numerically, as in the next section.

Example

The data in this example consist of the fatigue lifetimes (in 10^5 cycles) on coupons (rectangular strips) cut from aluminum sheeting that were tested under three levels of maximum stress. The original data set is provided by Birnbaum and Saunders [18]. Three sets of data were obtained, each under a periodic loading scheme with maximum stresses (acceleration levels) of $L_1 = 2.1$, $L_2 = 2.6$, and $L_3 = 3.1$ (in 10^4 psi).

A box-and-whisker plot of the data is presented in Figure 1. A glance at the figure shows a tendency that failure lifetimes decrease as maximum stress L increases. So the accelerated testing model seems appropriate.

To compare the fit of the proposed model with a competing model, the following *overall* mean square errors (MSEs) from the fitted model to the empirical distribution, over all values of L_i for $i = 1, 2, \dots, k$, were used. Letting $\hat{G}_{n_i}(s_{ij})$ denote the empirical cdf at L_i and $G_S(s_{ij}; \hat{\theta})$ denote the fitted cdf using the MLE $\hat{\theta}$ of the parameter vector $\theta = (\theta_0, \theta_1, \xi)$, the MSEs for the fitted model can be calculated as either

$$\text{MSE}_1(G_S(\cdot; \hat{\theta})) = \frac{1}{k} \sum_{i=1}^k \frac{1}{n_i} \sum_{j=1}^{n_i} \{G_S(s_{ij}; \hat{\theta}) - \hat{G}_{n_i}(s_{ij})\}^2 \quad (10)$$

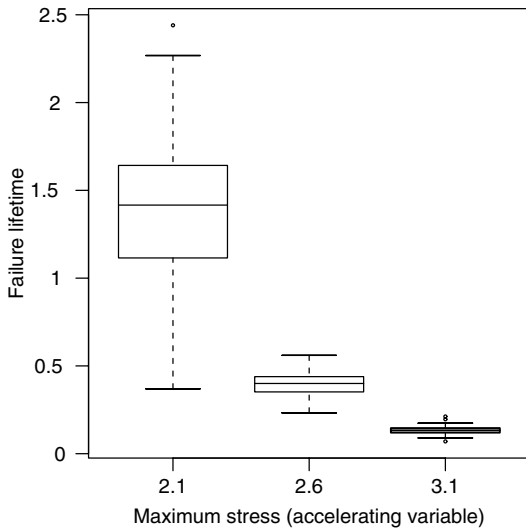


Figure 1 Box-and-whisker plot

or

$$\text{MSE}_2(G_S(\cdot; \hat{\theta})) = \frac{1}{N} \sum_{i=1}^k \sum_{j=1}^{n_i} \{G_S(s_{ij}; \hat{\theta}) - \hat{G}_{n_i}(s_{ij})\}^2 \quad (11)$$

where $N = \sum_{i=1}^k n_i$. Here, the empirical cdf $\hat{G}_{n_i}(s_{ij})$ is obtained by using so-called median rank method.

Numerically maximizing the log-likelihood function (9) when the pdf of S is given by the power-law Weibull model [13, 19] with its cdf

$$F_L(s) = 1 - \exp \left[-L^\theta \left(\frac{s}{\beta} \right)^\alpha \right] \quad (12)$$

the MLEs of the parameters are given by $\hat{\alpha} = 5.285829$, $\hat{\beta} = 162.3285$, and $\hat{\theta} = 32.93402$. For the proposed model with the power-law acceleration link, the MLEs of the parameters are $\hat{\xi} = 19.6607248$, $\hat{\theta}_0 = 0.1771115$, and $\hat{\theta}_1 = 5.9387513$. Using these MLEs, the MSEs are given by $\text{MSE}_1 = 4.823167 \times 10^{-3}$ and $\text{MSE}_2 = 4.818797 \times 10^{-3}$ for the power-law Weibull model, $\text{MSE}_1 = 4.735168 \times 10^{-3}$ and $\text{MSE}_2 = 4.733639 \times 10^{-3}$ for the proposed model. These results show that our model outperforms the power-law Weibull model.

To show these results graphically, we use Weibull plots. For such a plot, if the data are from a Weibull distribution, the estimated distribution is close to the plotted data points which form approximately a straight line. Departures from a straight line indicate departures from a Weibull distribution. Figure 2 gives the Weibull plots of the empirical cdf's (point marks) with the overall fits of the power-law Weibull model (dashed line) and the proposed model based on the gamma process (solid curve). It is easily seen from the figure that the data (empirical cdf) at each maximum stress L_i show a departure from a straight line. This suggests that a non-Weibull distribution is preferable, substantiating the use of our proposed model which more closely fits the data.

Concluding Remarks

A realistic approach to modeling cumulative damage is given here by using a gamma process incorporating an accelerating variable. The discrete version of the damage accumulation process which yielded the simpler failure distribution (8) can be used in a very

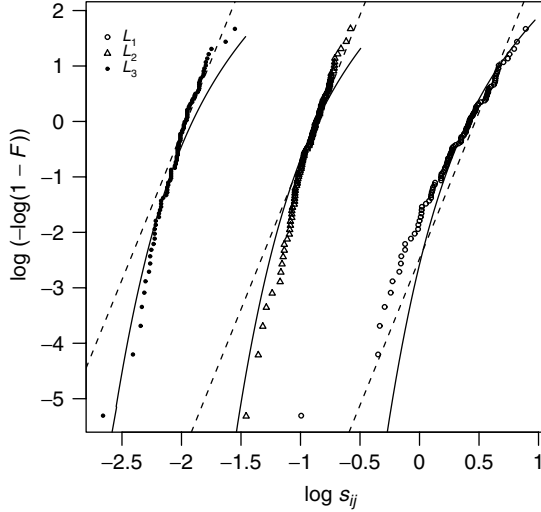


Figure 2 Weibull plot

general cumulative damage framework [5, 7]. This extends Durham and Padgett [11] which was based on Cramér's proportional growth model [20], and is briefly outlined here.

Assume that as the tensile load or stress level on a material specimen or a system is increased, the cumulative damage X_{n+1} after $n + 1$ increments of stress can be described by

$$X_{n+1} = X_n + h(X_n)\Delta_n \quad (13)$$

where Δ_n denotes the damage incurred at the $(n + 1)$ st increment and $h(\cdot)$ is the damage model function. Recently, Park and Padgett [5, 7] developed a more general cumulative damage model,

$$c(X_{n+1}) = c(X_n) + h(X_n)\Delta_n \quad (14)$$

where $c(\cdot)$ is the damage accumulation function. Based on this general framework, Park and Padgett [5, 7] improved and extended existing models. The selection of various forms of the functions $c(\cdot)$ and $h(\cdot)$ along with an appropriate stochastic process Δ_n generated several new models for material strength incorporating the size L as shown in Park and Padgett [5, 7]. The method of obtaining the MLE with an accelerating variable was also addressed in [5, 7]. It is also noteworthy that another approach to generate cumulative damage models (see [7]) can be based on Cauchy's basic functional equation.

In addition, the above discrete damage approach can be taken as the basic framework to develop a *degradation* model. A continuous version of the damage model (14) can be represented by

$$dc(X_u) = h(X_u) dD_u \quad (15)$$

This results in the cumulative damage (or level of degradation) of the system at the stress level (or time) t , given by

$$\int_0^t \frac{1}{h(X_u)} dc(X_u) = \int_0^t dD_u = D_t - D_0 \quad (16)$$

Finally, the aforementioned damage/degradation models have been further generalized by using a hypercuboidal volume approach, which can handle several accelerating variables; refer to Park and Padgett [6, 14].

Acknowledgment

The second author's work was partially supported by the National Science Foundation grant DMS-0243594 to the University of South Carolina.

References

- [1] Bagdonavicius, V. & Nikulin, M. (2000). Estimation in degradation models with explanatory variables, *Lifetime Data Analysis* **7**, 85–103.
- [2] Lawless, J. & Crowder, M. (2004). Covariates and random effects in a gamma process model with application to degradation and failure, *Lifetime Data Analysis* **10**, 213–227.
- [3] Lu, J. (1995). *Degradation processes and related reliability models*, Ph.D. thesis, McGill University.
- [4] Park, C. & Padgett, W.J. (2005). Accelerated degradation models for failure based on geometric Brownian motion and gamma processes, *Lifetime Data Analysis* **11**, 511–527.
- [5] Park, C. & Padgett, W.J. (2005). New cumulative damage models for failure using stochastic processes as initial damage, *IEEE Transactions on Reliability* **54**, 530–540.
- [6] Park, C. & Padgett, W.J. (2007). Cumulative damage models for failure with several accelerating variables, *Quality Technology and Quantitative Management* **4**, 17–34.
- [7] Park, C. & Padgett, W.J. (2006). A general class of cumulative damage models for materials failure, *Journal of Statistical Planning and Inference* **136**, 3783–3801.

- [8] Whitmore, G.A. (1995). Estimating degradation by a Wiener diffusion process subject to measurement error, *Lifetime Data Analysis* **1**, 307–319.
- [9] Whitmore, G.A., Crowder, M.J. & Lawless, J.F. (1998). Failure inference from a marker process based on a bivariate Wiener model, *Lifetime Data Analysis* **4**, 229–251.
- [10] Whitmore, G.A. & Schenkelberg, F. (1997). Modelling accelerated degradation data using Wiener diffusion with a scale transformation, *Lifetime Data Analysis* **3**, 27–45.
- [11] Durham, S.D. & Padgett, W.J. (1997). A cumulative damage model for system failure with application to carbon fibers and composites, *Technometrics* **39**, 34–44.
- [12] Nelson, W. (1990). *Accelerated Testing: Statistical Models, Test Plans, and Data Analyses*, John Wiley & Sons, New York.
- [13] Padgett, W.J., Durham, S.D. & Mason, A.M. (1995). Weibull analysis of the strength of carbon fibers using linear and power law models for the length effect, *Journal of Composite Materials* **29**, 1873–1884.
- [14] Park, C. & Padgett, W.J. (2006). Stochastic degradation models with several accelerating variables, *IEEE Transactions on Reliability* **55**, 379–390.
- [15] Onar, A. & Padgett, W.J. (2000). Inverse Gaussian accelerated test models based on cumulative damage, *Journal of Statistical Computation and Simulation* **66**, 233–247.
- [16] Mann, N.R., Schafer, R.E. & Singpurwalla, N.D. (1974). *Methods for Statistical Analysis of Reliability and Life Data*, John Wiley & Sons, New York.
- [17] Meeker, W.Q. & Escobar, L.A. (1998). *Statistical Methods for Reliability Data*, John Wiley & Sons, New York.
- [18] Birnbaum, Z.W. & Saunders, S.C. (1969). Estimation for a family of life distributions with applications to fatigue, *Journal of Applied Probability* **6**, 328–347.
- [19] Smith, R.L. (1991). Weibull regression models for reliability data, *Reliability Engineering and System Safety* **34**, 55–77.
- [20] Cramér, H. (1946). *Mathematical Methods of Statistics*, Princeton University Press, Princeton.

Related Articles

Accelerated Life Models; Brownian Motion; Degradation and Failure; Degradation Models; Warranty: Usage and Wear Process for.

CHANSEOK PARK AND WILLIAM J. PADGETT

Multistate Reliability Theory

Introduction

An inherent weakness of traditional reliability theory (see **Coherent Systems**) is that the system and the components are always described either as functioning or failed. Early attempts to replace this by a theory of multistate systems for multistate components were made in the late 1970s in [1, 2] and [3]. This was followed up by independent work in [4, 5] and [6] giving proper definitions of a multistate monotone system (MMS) and of multistate **coherent systems** (MCSs) and also of minimal **path and cut** vectors. Furthermore, in [7] upper and lower bounds for the availabilities and unavailabilities, to any level, in a fixed time interval, were obtained for MMSs based on corresponding information on the multistate components. These were assumed to be maintained and interdependent. Such bounds are of great interest when trying to predict the performance process of the system, noting that the exact expressions are obtainable just for trivial systems. Hence, by the mid-1980s the basic multistate reliability theory was established. A review of the early development in this area is given in [8]. Very recently probabilistic modeling of partial monitoring of components with applications to preventive system maintenance has been extended in [9] to MMSs of multistate components.

The theory was applied in [10] to an offshore electrical power generation system for two nearby oilrigs, where the amount of power that may possibly be supplied to the two oilrigs, is considered as system states. This application is also used to illustrate the theory in [9]. In [11] the theory was applied to the Norwegian offshore gas pipeline network in the North Sea, as of the end of the 1980s, transporting gas to Emden in Germany. The system state depends not only on the amount of gas actually delivered, but also to some extent on the amount of gas compressed, mainly by the compressor component closest to Emden. Recently, the first book [12] on multistate system reliability analysis and optimization was printed. The book also contains several examples of application of reliability assessment and optimization methods to real engineering problems.

Basic Concepts and Basic Bounds

Let $S = \{0, 1, \dots, M\}$ be the set of states of the system; then the $M + 1$ states represent successive levels of performance ranging from the perfect functioning level M down to the complete failure level 0. Furthermore, let $C = \{1, \dots, n\}$ be the set of components and S_i ($i = 1, \dots, n$) the set of states of the i th component. We claim $\{0, M\} \subseteq S_i \subseteq S$. Hence, the states 0 and M are chosen to represent the endpoints of a performance scale that might be used for both the system and its components. Let x_i ($i = 1, \dots, n$) denote the state or performance level of the i th component and $\mathbf{x} = (x_1, \dots, x_n)$. It is assumed that the state, ϕ , of the system is given by the structure function $\phi = \phi(\mathbf{x})$. For the following type of multistate systems a series of results can be derived.

Definition 1 A system is an *MMS* iff its structure ϕ satisfies:

- (i) $\phi(\mathbf{x})$ is nondecreasing in each argument
- (ii) $\phi(\mathbf{0}) = 0$ and $\phi(\mathbf{M}) = M$ ($\mathbf{0} = (0, \dots, 0)$, $\mathbf{M} = (M, \dots, M)$).

The first assumption says that improving one of the components cannot harm the system, whereas the second says that if all components are in the complete failure (perfect functioning) state, then the system is in the complete failure (perfect functioning) state.

We now impose some further restrictions on the structure function ϕ . The following notation is required:

$$\begin{aligned} (\cdot_i, \mathbf{x}) &= (x_1, \dots, x_{i-1}, \cdot, x_{i+1}, \dots, x_n), \\ S_{i,j}^0 &= S_i \cap \{0, \dots, j-1\} \text{ and} \\ S_{i,j}^1 &= S_i \cap \{j, \dots, M\} \end{aligned} \quad (1)$$

Definition 2 Consider an MMS with structure function ϕ satisfying

- (i) $\min_{1 \leq i \leq n} x_i \leq \phi(\mathbf{x}) \leq \max_{1 \leq i \leq n} x_i$.

If in addition $\forall i \in \{1, \dots, n\}$, $\forall j \in \{1, \dots, M\}$, $\exists (\cdot_i, \mathbf{x})$ such that

- (ii) $\phi(k_i, \mathbf{x}) \geq j$, $\phi(\ell_i, \mathbf{x}) < j$, $\forall k \in S_{i,j}^1$, $\forall \ell \in S_{i,j}^0$, we have a *multistate strongly coherent system (MSCS)*,

2 Multistate Reliability Theory

- (iii) $\phi(k_i, \mathbf{x}) > \phi(\ell_i, \mathbf{x}) \quad \forall k \in S_{i,j}^1, \quad \forall \ell \in S_{i,j}^0$, we have an MCS,
- (iv) $\phi(M_i, \mathbf{x}) > \phi(0_i, \mathbf{x})$, we have a *multistate weakly coherent system (MWCS)*.

When $M = 1$, all this reduces to the established binary coherent system (BCS) (see **Coherent Systems**). The structure function $\min_{1 \leq i \leq n} x_i$ ($\max_{1 \leq i \leq n} x_i$) is often denoted as the multistate series (parallel) structure.

Now choose $j \in \{1, \dots, M\}$ and let the states $S_{i,j}^0, (S_{i,j}^1)$ correspond to the failure (functioning) state for the i th component, if a binary approach is used. Condition (ii) above, means that for all components i and any level j , there shall exist a combination of the states of the other components, $(\cdot_i, \mathbf{x})$, such that if the i th component is in the binary failure (functioning) state, the system itself is in the corresponding binary failure (functioning) state. In general, modifying [6], condition (ii) says that every level of each component is relevant to the same level of the system, condition (iii) says that every level of each component is relevant to the system, whereas condition (iv) simply says that every component is relevant to the system.

For a BCS one can prove the following, practically very useful principle: **redundancy** at the component level is superior to redundancy at the system level except for a parallel system where it makes no difference. Assuming $S_i = S$ ($i = 1, \dots, n$) then this is also true for an MCS, but not for an MWCS.

We now discuss a special type of an MSCS. Define the indicators ($j = 1, \dots, M$) $I_j(x_i) = 1(0)$ if $x_i \geq j$ ($x_i < j$) and the indicator vector $(\mathbf{I}_j(\mathbf{x})) = (I_j(x_1), \dots, I_j(x_n))$.

Definition 3 An MSCS is said to be a *binary type multistate strongly coherent system (BTMSCS)* iff there exist binary coherent structures ϕ_j , $j = 1, \dots, M$, such that its structure function ϕ satisfies $\phi(\mathbf{x}) \geq j \Leftrightarrow \phi_j(\mathbf{I}_j, \mathbf{x}) = 1$ for all $j \in \{1, \dots, M\}$ and all $\mathbf{x}$.

Choose again $j \in \{1, \dots, M\}$ and let the states $S_{i,j}^0, (S_{i,j}^1)$ correspond to the failure (functioning) state for the i th component, if a binary approach is applied. By the definition above, ϕ_j will, from the binary states of the components, uniquely determine the corresponding binary state of the system.

In what follows $\mathbf{y} < \mathbf{x}$ means $y_i \leq x_i$ for $i = 1, \dots, n$, and $y_i < x_i$ for some i .

Definition 4 Let ϕ be the structure function of an MMS and let $j \in \{1, \dots, M\}$. A vector $\mathbf{x}$ is said to be a *minimal path (cut) vector to level j* iff $\phi(\mathbf{x}) \geq j$ and $\phi(\mathbf{y}) < j$ for all $\mathbf{y} < \mathbf{x}$ ($\phi(\mathbf{x}) < j$ and $\phi(\mathbf{y}) \geq j$ for all $\mathbf{y} > \mathbf{x}$).

Definition 5 The *performance process of the i th component* ($i = 1, \dots, n$) is a stochastic process $\{X_i(t), t \in [0, \infty)\}$, where for each fixed $t \in [0, \infty)$ $X_i(t)$ is a random variable which takes values in S_i . The *joint performance process for the components* $\{X(t), t \in [0, \infty)\} = \{(X_1(t), \dots, X_n(t)), t \in [0, \infty)\}$ is the corresponding vector stochastic process. The *performance process of an MMS* with structure function ϕ is a stochastic process $\{\phi(\mathbf{X}(t)), t \in [0, \infty)\}$, where for each fixed $t \in [0, \infty)$, $\phi(\mathbf{X}(t))$ is a random variable which takes values in S .

Definition 6 The performance processes $\{X_i(t), t \in [0, \infty)\}$, $i = 1, \dots, n$ are *independent* in the time interval I iff, for any integer m and $\{t_1, \dots, t_m\} \subset I$ the random vectors $\{X_1(t_1), \dots, X_1(t_m)\}, \dots, \{X_n(t_1), \dots, X_n(t_m)\}$ are independent.

Definition 7 Let $j \in \{1, \dots, M\}$. The *availability*, $h_\phi^{j(I)}$ and the *unavailability*, $g_\phi^{j(I)}$ to level j in the time interval I for an MMS with structure function ϕ are given by

$$\begin{aligned} h_\phi^{j(I)} &= P[\phi(\mathbf{X}(s)) \geq j \quad \forall s \in I], \\ g_\phi^{j(I)} &= P[\phi(\mathbf{X}(s)) < j \quad \forall s \in I] \end{aligned} \quad (2)$$

Note that $h_\phi^{j(I)} + g_\phi^{j(I)} \leq 1$, with equality for the case $I = [t, t]$.

As an example of the bounds for $h_\phi^{j(I)}$ and $g_\phi^{j(I)}$ given in [7], we give the following theorem by first introducing the $n \times M$ matrices

$$\begin{aligned} \mathbf{P}_\phi^{(I)} &= \{p_i^{j(I)}\}_{j=1, \dots, M}^{i=1, \dots, n}, \\ \mathbf{Q}_\phi^{(I)} &= \{q_i^{j(I)}\}_{j=1, \dots, M}^{i=1, \dots, n} \end{aligned} \quad (3)$$

Note that according to [13] we do not need to assume that each of the performance processes of the components is associated in I .

Theorem 1 Let (C, ϕ) be an MMS with the marginal performance processes of its components being independent in I . Furthermore, for $j \in \{1, \dots, M\}$ let $\mathbf{y}_k^j = (y_{1k}^j, \dots, y_{nk}^j)$, $k = 1, \dots, n^j$ ($\mathbf{z}_k^j = (z_{1k}^j,$

$\dots, z_{nk}^j), k = 1, \dots, m^j)$ be its minimal path (cut) vectors to level j . Define

$$\ell_\phi^{j'}(\mathbf{P}_\phi^{(I)}) = \max_{1 \leq k \leq n^j} \prod_{i=1}^n p_i^{y_{ik}^{j'}(I)}$$

$$\bar{\ell}_\phi^{j'}(\mathbf{Q}_\phi^{(I)}) = \max_{1 \leq k \leq m^j} \prod_{i=1}^n q_i^{z_{ik}^{j'}+1(I)} \quad (4)$$

$$\ell_\phi^{j*}(\mathbf{P}_\phi^{(I)}) = \prod_{k=1}^{m^j} \prod_{i=1}^n p_i^{z_{ik}^{j*}+1(I)}$$

$$\bar{\ell}_\phi^{j*}(\mathbf{Q}_\phi^{(I)}) = \prod_{k=1}^{n^j} \prod_{i=1}^n q_i^{y_{ik}^{j*}(I)} \quad (5)$$

$$B_\phi^j(\mathbf{P}_\phi^{(I)}) = \max_{j \leq k \leq M} \times \left\{ \max \left[\ell_\phi^{k'}(\mathbf{P}_\phi^{(I)}), \ell_\phi^{k*}(\mathbf{P}_\phi^{(I)}) \right] \right\} \quad (6)$$

$$\bar{B}_\phi^j(\mathbf{Q}_\phi^{(I)}) = \max_{1 \leq k \leq j} \times \left\{ \max \left[\bar{\ell}_\phi^{k'}(\mathbf{Q}_\phi^{(I)}), \bar{\ell}_\phi^{k*}(\mathbf{Q}_\phi^{(I)}) \right] \right\} \quad (7)$$

Then

$$B_\phi^j(\mathbf{P}_\phi^{(I)}) \leq h_\phi^{j(I)} \leq \inf_{t \in I} \left[1 - \bar{B}_\phi^j(\mathbf{Q}_\phi^{(I,t)}) \right]$$

$$\leq 1 - \bar{B}_\phi^j(\mathbf{Q}_\phi^{(I)}) \quad (8)$$

$$\bar{B}_\phi^j(\mathbf{Q}_\phi^{(I)}) \leq g_\phi^{j(I)} \leq \inf_{t \in I} \left[1 - B_\phi^j(\mathbf{P}_\phi^{(I,t)}) \right]$$

$$\leq 1 - B_\phi^j(\mathbf{P}_\phi^{(I)}) \quad (9)$$

Here $\prod_{i=1}^n a_i \stackrel{\text{def}}{=} 1 - \prod_{i=1}^n (1 - a_i)$. By specializing $M = 1$ and $I = [t, t]$ the bounds reduce to the familiar ones from binary theory as given in [14].

An Offshore Electrical Power Generation System

The purpose of the offshore electrical power generation system considered in [9] and [10], depicted in Figure 1, is to supply two nearby oilrigs with electrical power. Both oilrigs have their own main generation, represented by equivalent generators A_1 and A_3 , each having a capacity of 50 MW. In addition oilrig 1 has a standby generator A_2 that is switched

into the network in case of outage of A_1 or A_3 . A_2 also has a capacity of 50 MW. The control unit, U , continuously supervises the supply from each of the generators with an automatic control of the switches. If, for instance, the supply from A_3 to oilrig 2 is not sufficient, whereas the supply from A_1 to oilrig 1 is sufficient, U can activate A_2 to supply oilrig 2 with electrical power through the standby subsea cables L .

The components to be considered here are A_1, A_2, A_3, U and L . We let the perfect functioning level M equal 4 and let the set of states of all components be $\{0, 2, 4\}$. For A_1, A_2 and A_3 these states are interpreted as

- 0: The generator cannot supply any power.
- 2: The generator can supply maximum 25 MW power.
- 4: The generator can supply maximum 50 MW power.

Note that as an approximation we have chosen to describe supply capacity of the generators using a discrete scale of three points. The supply capacity is not a measure of the actual amount of power delivered at a fixed point of time.

The control unit U has the states

- 0: U will by mistake switch off the main generators A_1 and A_3 without switching on A_2 ;
- 2: U will not switch on A_2 when needed;
- 4: U is functioning perfectly.

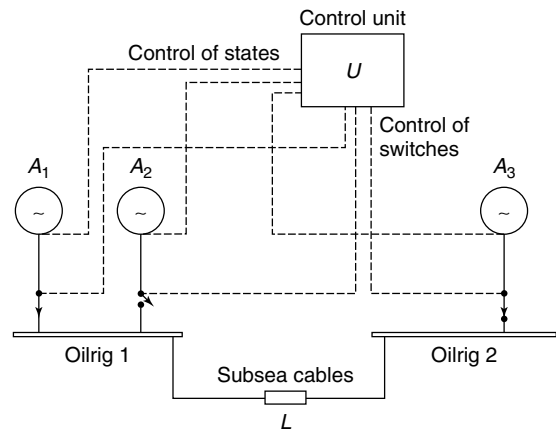


Figure 1 Outline of an offshore electrical power generation system

4 Multistate Reliability Theory

The subsea cables L are actually assumed to be constructed as double cables transferring half of the power through each cable. This leads to the following states of L

- 0: No power is transferred.
- 2: 50% of the power is transferred.
- 4: 100% of the power is transferred.

Let us now for simplicity assume that the mechanism that distributes the power from A_2 to oilrig 1 or 2 is working perfectly, transferring excess power from A_2 to oilrig 2 if oilrig 1 is ensured a delivery corresponding to state 4. Now let $\phi(A_1, A_2, A_3, U, L)$ = the amount of power that can be supplied to oilrig 2. In addition to the states taken by A_1, A_2, A_3 , ϕ can also take the following states

- 1: The maximum amount of power that can be supplied is 12.5 MW.
- 3: The maximum amount of power that can be supplied is 37.5 MW.

Number the components A_1, A_2, A_3, U , and L successively as 1, 2, 3, 4, and 5. Then this leads to

$$\phi(\mathbf{x}) = I(x_4 > 0) \min(x_3 + \max(x_1 + x_2 I(x_4 = 4) - 4, 0)x_5/4, 4) \quad (10)$$

noting that $\max(x_1 + x_2 I(x_4 = 4) - 4, 0)$ is just the excess power from A_2 , which one tries to transfer to oilrig 2. This is obviously a MMS.

References

- [1] Barlow, R.E. & Wu, A.S. (1978). Coherent systems with multistate components, *Mathematics of Operations Research* **4**, 275–281.
- [2] El-Newehi, E., Proschan, F. & Sethuraman, J. (1978). Multistate coherent systems, *Journal of Applied Probability* **15**, 675–688.
- [3] Ross, S. (1979). Multivalued state component reliability systems, *Annals of Probability* **7**, 379–383.
- [4] Griffith, W. (1980). Multistate reliability models, *Journal of Applied Probability* **17**, 735–744.
- [5] Natvig, B. (1982). Two suggestions of how to define a multistate coherent system, *Advances in Applied Probability* **14**, 434–455.
- [6] Block, H.W. & Savits, T.S. (1982). A decomposition for multistate monotone systems, *Journal of Applied Probability* **19**, 391–402.
- [7] Funnemark, E. & Natvig, B. (1985). Bounds for the availabilities in a fixed time interval for multistate monotone systems, *Advances in Applied Probability* **17**, 638–665.
- [8] Natvig, B. (1985). Multistate coherent systems, in *Encyclopedia of Statistical Sciences*, N.L. Johnson & S. Kotz, eds, John Wiley & Sons Vol. 5, pp. 732–735.
- [9] Gåsemyr, J. & Natvig, B. (2005). Probabilistic modelling of monitoring and maintenance of multistate monotone systems with dependent components, *Methodology and Computing in Applied Probability* **7**, 63–78.
- [10] Natvig, B., Sørmo, S., Holen, A.T. & Høgåsen, G.T. (1986). Multistate reliability theory – a case study, *Advances in Applied Probability* **18**, 921–932.
- [11] Natvig, B. & Mørch, H.W. (2003). An application of multistate reliability theory to an offshore gas pipeline network, *International Journal of Reliability, Quality and Safety Engineering* **10**, 361–381.
- [12] Lisnianski, A. & Levitin, G. (2003). *Multi-State System Reliability*, World Scientific, London.
- [13] Natvig, B. (1993). Strict and exact bounds for the availabilities in a fixed time interval for multistate monotone systems, *Scandinavian Journal of Statistics* **20**, 171–175.
- [14] Barlow, R.E. & Proschan, F. (1975). *Statistical Theory of Reliability and Life Testing: Probability Models*, Holt, Rinehart and Winston, New York.

Related Articles

Coherent Systems; Multicomponent Maintenance; Path Sets and Cut Sets in System Reliability Modeling; Reliability of Redundant Systems; System Availability.

BENT NATVIG

Masked Failure Data: Competing Risks

Introduction

The term **competing-risks problem** has come to encompass the study of any failure process in which there is more than one distinct cause or type of failure. Frequent reference is made to possible influences of competing failure types in studies reported in the industrial, engineering as well as clinical, epidemiologic, demographic, and basic science literature. In engineering applications, the causes or risks may signify either multiple modes of failure for a complex unit or multiple components or subsystems which comprise an entire system. Occurrence of a system failure is caused by the earliest onset of any of these component failures. In this respect, the framework is that of a system with components connected in series. An additional complication arises when the causes of failure for a subset of systems are only known to belong to certain subsets of all possible causes; this is generally referred to as the cause of failure being *masked*. Masking is often the manifestation of an attempt to expedite the process of repair by replacing the entire subset of components responsible for failure instead of further investigation toward identifying the specific cause for failure. In practice, one possibility is to carry out a second-stage analysis to uniquely determine the cause; however, statistical inference may be possible even when the cause of failure is masked for a subset of systems.

Examples of failure data obtained under competing risks are abundant in reliability applications in engineering as well as in biomedical applications. In a medical context, the causes of failure (a failure may indicate death or relapse or other time-to-event outcome) typically refer to various potential risk factors for a patient (or animal) observed in a biomedical study. For example, Basu *et al.* [1] consider the risks from breast cancer and heart disease as well as other causes for breast cancer patients.

The literature on competing risks spans the work of Berman [2] to very recent work including the 2006 special issue of (JSPI) Journal of Statistical Planning and Inference [3]; the latter contains a wealth of references. An excellent resource for statistical theory and analysis of competing risks is the book by

Crowder [4]. The more recent book by Pintilie [5] emphasizes practical application of competing-risks theory, mostly in biomedical applications. In engineering applications, Langseth and Lindqvist [6] considered an interesting application of competing-risks models in the repairable systems, whereas Bunea and Mazzuchi [7] studied **competing failure modes in accelerated life testing**. Sun and Tiwari [8] analyzed the failure times of small electric appliances that may fail due to two competing risk whereas Taylor [9] modeled the tensile strength of certain materials known to contain two or more flaw types.

Competing Risks (without masking)

The early approach to analysis of competing-risks data introduced conceptual or latent failure times [2, 10] T_j , $j = 1, \dots, K$, corresponding to time to failure from risks $j = 1, \dots, K$, with associated probability density function $f_j(t)$, marginal survival function $S_j(t) = 1 - F_j(t)$, and net hazard function $h_j(t) = f_j(t)/S_j(t)$. Thus T_j is the potential failure time from cause $C = j$ that would be observed if the possibility of failure from causes other than j were removed. The observed lifetime T is then taken to be $T = \min(T_1, \dots, T_K)$. Let $S(t) = P(T > t)$ be the survival function for the observed lifetime T ; under the assumption that the potential failure times $T_1, \dots, T_K$ from the K competing risks are independent we have $S(t) = \prod_{j=1}^K S_j(t)$. The observation triplet is (T, C, δ) where C is the cause of failure and δ is the censoring indicator. The likelihood for an observation at time t can be expressed as

$$L(t, \delta = 0) = S(t) \text{ for a censored case} \quad (1)$$

$$\begin{aligned} L(t, C = j, \delta = 1) &= f_j(t) \prod_{l \neq j} S_l(t) \\ &= S(t) h_j(t) \text{ if cause is} \\ &\quad \text{identified as } C = j \end{aligned} \quad (2)$$

$$\begin{aligned} L(t, C \in \mathcal{C}, \delta = 1) &= \sum_{j \in \mathcal{C}} f_j(t) \prod_{l \neq j} S_l(t) \\ &= S(t) \sum_{j \in \mathcal{C}} h_j(t) \end{aligned} \quad (3)$$

where the last expression applies when the cause of failure is masked and is narrowed down to a subset $\mathcal{C}$

2 Masked Failure Data: Competing Risks

of $\{1, \dots, K\}$. This construction, of course, assumes independence of the K competing risks.

The random variable T_j used in the latent formulation is often referred to as the *potential failure time* since it represents the time to failure of the system when all other risks except the j th risk are removed from the system. The usefulness and interpretability of this model have been the subjects of debate [11, 12] since in most cases, it is either not feasible or physically impossible to remove all other risks from the system. The alternative cause-specific hazard formulation [11–13] is based on the joint distribution of the observed survival time T and cause of failure C . In particular, let $S(j, t) = P(C = j, T > t)$ be a “subsurvival” function with $S(t) = \sum_{j=1}^K S(j, t)$ denoting the marginal survival function of T . The corresponding subdensity function is denoted as $f(j, t)$. The cause-specific or subhazard function

$$\begin{aligned} h(j, t) &= \lim_{\varepsilon \rightarrow 0} \frac{P(C = j, T \leq t + \varepsilon | T > t)}{\varepsilon} \\ &= \frac{f(j, t)}{S(t)} \end{aligned} \quad (4)$$

is the instantaneous failure rate from cause j at time t after surviving from all risks $1, \dots, K$ up to time t . The cause-specific subhazard functions $h(j, t)$ are thus defined in terms of the joint distribution of the observed time and cause of failure (T, C) and constitute the basis of the cause-specific approach. The overall hazard from all risks combined is $h(t) = \sum_{j=1}^K h(j, t)$. The marginal survival can be written in term of the cause-specific hazards as $S(t) = \exp\left[-\sum \int_0^t h(j, s) ds\right]$ and each of the likelihood terms in equations (1–3) can be reexpressed by replacing the net hazard $h_j(t)$ with the cause-specific subhazard $h(j, t)$. The cumulative incidence function is defined as

$$\begin{aligned} P(T < t, J = j) &= \int_0^t f(j, u) du \\ &= \int_0^t h(j, u) S(u) du \end{aligned} \quad (5)$$

and may serve as a useful descriptive device [13, 14].

The book by Crowder [4] contains detailed comparison of the latent and the cause-specific approaches. Under the assumption of independence of the potential failure times $T_1, \dots, T_K$, the net hazard $h_j(t)$ and the cause-specific hazard $h(j, t)$ are identical (see Crowder [4] and Kalbfleisch and Prentice [12]), and the latent and cause-specific approaches lead to identical statistical inference. Independence assumes that the time of failure from cause j under one set of study conditions in which all K causes are operative is precisely the same as under an altered set of conditions in which all causes except the j th have been removed. However, the elimination of certain risks may well alter the hazards from other causes making the independence assumption of questionable validity (see Lagakos [15]). In addition, the all-too-known identifiability problem [16–20] entails that given a competing-risks model with a specific joint survivor function $\bar{F}(t_1, \dots, t_K)$ and arbitrary dependence between $T_1, \dots, T_K$, there exists a different joint survivor function in which $T_1, \dots, T_K$ are independent such that both models reproduce the same set of subsurvival functions $S(j, t)$. In particular, one cannot distinguish between the dependent model and the proxy independent model on the basis of observations on (C, T) alone. In fact, each such a model has a whole class of proxy models (not necessarily independent). Crowder [19] has extended this result to show that (under mild conditions) there exist infinitely many proxy joint survivor functions which will reproduce both the set of subsurvival functions $S(j, t)$ as well as the set of marginal survival functions $S_j(t)$, $j = 1, \dots, K$. Identifiability, however, can be regained when covariates are present. Heckman and Honore [21] showed that when there are explanatory variables in the model, identification of the joint survivor function $\bar{F}(t_1, \dots, t_K)$ is possible from the subsurvivor functions within a certain framework. Slud [22] proved a result of similar spirit under a different setup.

Moeschberger and Klein [23] provided a thorough review of analytic approaches for handling competing-risks data when the independence assumption is suspect. One approach is to assume a parametric joint distribution for the latent lifetimes T_j , $j = 1, \dots, K$. Parametric models are common in reliability applications and the identifiability problem is much less pervasive here (as opposed to non- or semiparametric models, which are more popular

in biomedical applications). Moeschberger [24] suggested using multivariate Weibull and normal distributions for the joint distribution of the latent lifetimes $(T_1, \dots, T_K)$. Basu and Klein [25] established that there is no identifiability problem in this setting. Basu and Ghosh [26, 27] demonstrated that many parametric distributions, including the Marshall–Olkin and the Gumbel distribution, are identifiable (up to a permutation) from (T, C) , the time and the cause.

As we noted before, the assumption of independence of the potential failure times $T_1, \dots, T_K$ implies that the cause-specific hazard $h(j, t)$ and the net hazard $h_j(t)$ are the same for all j and all $t > 0$. The latter equality $h(j, t) = h_j(t)$, $\forall j, t$ is known as the *Makeham assumption* [4]. While independence implies the Makeham condition, Gail [28] showed that the Makeham condition, in turn, implies that the subsurvival functions $S(j, t)$ determine the marginal survival functions $S_j(t)$. This implies $\bar{F}(t, t, \dots, t) = \prod_{j=1}^K S_j(t)$ or “independence on the diagonal”, but rather interestingly, the Makeham assumption may not imply independence (see William and Lagakos [29], and Crowder [4, 19]).

Competing Risks for Masked Data

Masking refers to the case when the causes of failure for a subset of systems are not exactly identified, but instead, have been narrowed down to subsets of $\{1, \dots, K\}$. Dinse [30] and Kodell and Chen [31] considered bioassays for animal carcinogenicity where masking arises owing to the disagreement on the reliability of cause of death information. They focused on nonparametric estimation of survival functions associated with the cause of interest as well as the other causes. Racine-Poon and Hoel [32] formally accounted for the uncertainty in diagnosing the exact cause of death by assigning a score on the diagnostic probability being correct. Dinse [33] obtained nonparametric maximum-likelihood estimates (MLEs) of cause-specific hazards in the setting of two competing risks and indicated a generalization to the multirisk case. Dinse [34] developed nonparametric estimation of the incidence rate. Lagakos and Louis [35] investigated various nonparametric tests for comparing difference in survival functions for two groups under masking. Goetghebeur and Ryan [36] derived a modified log-rank test for comparing the survival of two populations that was later modified by Dewanji [37]. Goetghebeur and Ryan [38]

developed proportional hazards structure in the context of two competing risks under the assumption that the baseline hazards for the two competing risks are proportional. Lu and Tsiatis [39] utilized multiple imputations to generalize this to the case when the baselines are not proportional. Flehinger *et al.* [40] proposed maximum-likelihood estimation under a similar model. Dewanji and SenGupta [41] developed an **EM algorithm** in the setting of masked but grouped survival data. Lu and Tsiatis [42] compared the two partial-likelihood approaches in masked data.

The parametric models are typically considered in engineering applications and are mostly based on the latent lifetime approach. One of the earliest parametric models was proposed by Miyakawa [43], who derived closed-form estimates for the case of two competing risks, each following exponential failure time distribution. In the same exponential framework, Usher and Hodgson [44] and Lin *et al.* [45] presented iterative algorithms for obtaining maximum-likelihood estimators when three competing risks are present. Guess *et al.* [46] presented a general framework that included an arbitrary number of competing risks. They introduced the term *minimum random subset* (MRS) to denote the subset $\mathcal{C}$ of the causes $\{1, \dots, K\}$ to which the cause of failure can be narrowed down. They also noticed that the likelihood term in equations (1–3) for the masked case actually contains additional terms involving the masking probabilities $Q(\mathcal{C}|C = j, t)$ and is given by

$$L(t, C \in \mathcal{C}, \delta = 1) = S(t) \sum_{j \in \mathcal{C}} \{h_j(t) Q(\mathcal{C}|C = j, t)\} \quad (6)$$

where the masking probabilities are defined as

$$Q(\mathcal{C}|C = j, t) = P(\text{cause is masked in subset } \mathcal{C} | C = j, t) \quad (7)$$

As noted by Flehinger *et al.* [40, 47], Craiu and Reiser [48], Craiu and Lee [49], and Craiu and Duchesne [50], the diagnostic probabilities

$$\pi(C = j|\mathcal{C}, t) = P(\text{actual failure due to cause } j | \text{cause masked in } \mathcal{C}, t) \quad (8)$$

which can be obtained from $Q(\mathcal{C}|C = j, t)$ and the likelihood model *via* Bayes rule are of interest. Guess *et al.* [46] described two “noninformative masking”

and “symmetry” conditions under which one can ignore the $Q(C|C = j, t)$ terms and proceed with likelihood inference. Guess *et al.* [46] and Lin *et al.* [45] cited industrial examples where it is reasonable to assume that these conditions hold. Similar symmetry conditions are assumed in Schäbe’s [51] and Goetghebeur and Ryan’s [38] semiparametric formulations. Lin and Guess [52] and Guttman *et al.* [53] use a proportionality assumption to meet the symmetry condition. Flehinger *et al.* [40] consider an interesting case when further data are available from second-stage autopsy on a subset of masked units. Craiu and Duchesne [50] and Craiu and Reiser [48] used EM-based methods to model and analyze these masking probabilities. Craiu and Lee [49] developed model selection for these models, whereas Craiu and Duchesne [54] compared the performances of EM and Bayesian data augmentation methods in this scenario. Sen *et al.* [55] provided a recent review of competing-risks analysis for masked data. Mukhopadhyay [56] obtained (MLEs) and bootstrap estimates of these masking probabilities in a general setting and established consistency and as well asymptotic normality of the MLEs under suitable regularity conditions. Kuo and Yang [57] developed different models for the masking probabilities and described Bayesian analysis of these models using Gibbs sampling methods. Mukhopadhyay and Basu [58] described a general Bayesian model for the masking probabilities; the details of this model is discussed in the next section.

Bayesian Approaches

Bayesian analysis of masked competing risks has been mostly limited to parametric models and the latent failure time approach, but see Craiu and Duchesne [54]. Mukhopadhyay and Basu [59], in one of the early Bayesian works, assumed that the component net survival distributions are exponential and obtain closed-form Bayesian estimates under a Dirichlet prior. Reiser *et al.* [60] considered two independent latent risks following exponential distributions, derived closed-form Bayes estimators for a noninformative Jeffreys’ prior, and showed that they are identical to the corresponding MLEs. Other prior choices under exponential lifetimes are discussed in Lin *et al.* [61] and Mukhopadhyay and Basu [62]. Guttman *et al.* [53] developed Bayesian analysis relaxing the symmetry condition for two competing

risks with associated exponential net survival distributions. Berger and Sun [63], Mukhopadhyay and Basu [62], and Basu *et al.* [64] considered the case of arbitrary number of competing risks where the latent failure time of each risk follows a Weibull distribution. While Mukhopadhyay and Basu [62] provided an EM algorithm, Basu *et al.* [64] used data augmentation and Gibbs sampling in the Bayesian analysis. Basu *et al.* [65] and Sen *et al.* [55] provided a general framework for parametric Bayesian analysis which encompasses location-scale and log-location-scale distributions. They used Markov chain sampling for the posterior analysis where they introduced latent variables C_i to augment the causes of failure for the masked cases. Basu *et al.* [65] also considered the case of interval censoring (when, instead of actually observing the time to event, events are only known to occur within a specified interval), applied these Bayesian parametric models in two engineering applications and provided Bayesian estimates of the diagnostic probability for the masked cases. They further compared different parametric models by Bayes factors and provide model diagnostics using the cross-validated predictive approach.

Mukhopadhyay and Basu [58] described a general approach for Bayesian analysis with the masking probabilities $Q(C|C = j, t)$. Suppose that $Q(C|C = j, t) = Q(C|C = j)$ are free of t and let $\mathcal{F}_j$ denote the collection of all subsets of $\{1, \dots, K\}$ which includes j . We then have $\sum \{Q(C|C = j) : C \in \mathcal{F}_j\} = 1$ and Mukhopadhyay and Basu [58] assigned Dirichlet distribution prior on this simplex. Such a prior specification, in fact, produces a conditionally conjugate structure, which lead to a straightforward updating scheme for the masking probabilities within the Markov chain sampling. Mukhopadhyay and Basu [58] obtained Bayesian estimates of both the masking and the diagnostic probabilities in the data example. As we noted before, Kuo and Yang [57] also considered Bayesian analysis of the masking probabilities, whereas Craiu and Duchesne [54] compared Bayesian data augmentation with the EM algorithm.

Applications

Crowder [4] and Nelson [66] listed a collection of data sets from engineering applications involving competing risks. They include data on breaking strength of wire connections (2 possible causes of

failure), data on failure time of electric appliances (many possible causes), data from accelerated life tests of industrial heaters and data from low-cycle fatigue test of a superalloy. Langseth and Lindqvist [6] analyzed an interesting dataset from the Off-shore Reliability Database which involved competing risks in repairable systems. Flehinger *et al.* [47] listed data of 172 failure times observed over a period of 4 years during which 10 000 computer hard drives were monitored. There are three possible causes of failure; but they are masked for some of the systems in either $\mathcal{C} = \{1, 3\}$ or $\mathcal{C} = \{1, 2, 3\}$. These data were also analyzed by Craiu and Reiser [48]. Basu *et al.* [65] analyzed data from a 10 000-hr field trial of 4993 circuit-boards reported in Chan and Meeker [67] and considered “infant mortality” and wearout as two competing risks. There is an interval when both modes are simultaneously operative and the cause for a failure within this interval is considered as masked. Reiser *et al.* [60], Basu *et al.* [64, 65], Mukhopadhyay [56], and Mukhopadhyay and Basu [58] analyzed time-to-failure data of computer component systems where failure was attributed to malfunctioning of one of the three components. Out of a total of eight system failures, the causes of three were masked. In addition, 676 systems were progressively censored at various times and this heavy censoring poses an additional challenge to statistical analysis. Basu *et al.* [65] analyzed these data using different parametric models and noted that due to the limited information in the data, the Bayesian inference is expected to depend on the assumed model and prior. Mukhopadhyay [56] and Mukhopadhyay and Basu [58] analyzed these data using a general model for the masking probabilities $Q(\mathcal{C}|C = j)$. The maximum-likelihood analysis in Mukhopadhyay [56], for example, estimate that if the true cause of failure is $C = 3$, then the probability that the cause will be identified $\hat{Q}(3|C = 3) \approx 0.5$, the probability that the cause will be masked in $\{2, 3\}$ is $\hat{Q}(\{2, 3\}|C = 3) \approx 0.5$, whereas $\hat{Q}(\{1, 3\}|C = 3)$ and $\hat{Q}(\{1, 2, 3\}|C = 3)$ are ≈ 0 . The Bayesian estimates in Mukhopadhyay and Basu [58], however, assign approximately 0.20 probability to each of $Q(\{1, 3\}|C = 3)$ and $Q(\{1, 2, 3\}|C = 3)$. The diagnostic probabilities $\pi(C = j|\mathcal{C})$ reported in Basu *et al.* [69] also differ from the estimates reported in Mukhopadhyay and Basu [58]. Such diversity in the results, however, are expected due to the limited information available in these data and the

diversity in the models used for analysis. Mukhopadhyay and Basu, for example, used noninformative priors whereas Basu *et al.* [65] argued for a rather informative prior.

In their review on analysis of masked data in competing risks, Sen *et al.* [55] cited dependent causes, symmetry assumption, and covariate inclusion as research areas that need further exploration. The recent literature has seen substantial progress in each of these areas. Identifiability concerns from competing risks are more pervasive in masked data (when some causes are not completely known); while individual results are scattered in the literature, we are not aware of a general result. The shortage in the availability of rich data from reliability applications is also often a deterrent to methodological development in this area.

Acknowledgment

This material is based upon work partially supported by the National Science Foundation under Grant No. 0306416. Any opinions, findings, and conclusions or recommendations expressed in this material are those of the author(s) and do not necessarily reflect the views of the National Science Foundation.

References

- [1] Basu, S., Tiwari, R.C. & Feuer, E.J. (2007). *Bayesian Analysis of Cancer Survival Data Using Competing Risks and Mixture Cure Models*, Technical Report, Northern Illinois University.
- [2] Berman, S.M. (1963). Notes on extreme values, competing risks, and semi-Markov processes, *Annals of Mathematical Statistics* **34**, 1104–1106.
- [3] Deshpande, J.V. & Cooke, R.M. (2006). Editors, competing risks: theory and applications, *A Special Volume of the Journal of Statistical Planning and Inference* **136**(5), 1569–1746.
- [4] Crowder, M. (2001). *Classical Competing Risks*, Chapman & Hall/CRC, London.
- [5] Pintilie, M. (2006). *Competing Risks: A Practical Perspective*, John Wiley & Sons, Hoboken.
- [6] Langseth, H. & Lindqvist, B.H. (2006). Competing risks for repairable systems: a data study, *Journal of Statistical Planning and Inference* **136**(5), 1687–1700.
- [7] Bunea, C. & Mazzuchi, T.A. (2006). Competing failure modes in accelerated life testing, *Journal of Statistical Planning and Inference* **136**(5), 1608–1620.
- [8] Sun, Y. & Tiwari, R.C. (1997). Comparing cumulative incidence functions of a competing-risks model, *IEEE Transactions on Reliability* **46**, 247–253.

- [9] Taylor, H.M. (1994). The Poisson-Weibull flaw model for brittle fiber strength, in *Extreme Value Theory*, J. Galambos, J. Lechner & Emil Simiu, eds, Kluwer Academic Publishers, Amsterdam, Vol 1, pp. 43–59.
- [10] David, H.A. & Moeschberger, M.L. (1978). *The Theory of Competing Risks*, Griffin's Statistical Monographs & Courses, Vol. 39, Charles W. Griffin, London.
- [11] Prentice, R.L., Kalbfleisch, J.D., Peterson, A.V., Flournoy, N., Farewell, V.T. & Breslow, N.E. (1978). The analysis of failure time data in the presence of competing risks, *Biometrics* **34**, 541–554.
- [12] Kalbfleisch, J.D. & Prentice, R.L. (2002). *The Statistical Analysis of Failure Time Data*, 2nd Edition, John Wiley & Sons, New York.
- [13] Pepe, M.S. & Mori, M. (1993). Kaplan-Meier, marginal or conditional probability curves in summarizing competing risks failure time data? *Statistics in Medicine* **12**, 737–751.
- [14] Gray, R.J. (1988). A class of k -sample tests for comparing the cumulative incidences of a competing risk, *Annals of Statistics* **16**, 1141–1154.
- [15] Lagakos, S.W. (1979). General right censoring and its impact on the analysis of survival data, *Biometrics* **35**, 139–156.
- [16] Cox, D.R. (1959). The analysis of exponentially distributed lifetimes with two types of failures, *Journal of the Royal Statistical Society. Series B* **21**, 411–421.
- [17] Tsiatis, A. (1975). A nonidentifiability aspect of the problem of competing risks, *Proceedings of the National Academy of Sciences of the United States of America* **72**, 20–22.
- [18] Peterson, A.V. (1976). Bounds for a joint distribution function with fixed sub-distribution functions: applications to competing risks, *Proceedings of the National Academy of Sciences* **73**, 11–13.
- [19] Crowder, M. (1991). On the identifiability crisis in competing risks theory, *Scandinavian Journal of Statistics* **18**, 223–233.
- [20] Crowder, M. (1994). Identifiability crises in competing risks, *International Statistical Review* **62**, 379–391.
- [21] Heckman, J.J. & Honore, B.E. (1989). The identifiability of the competing risks model, *Biometrika* **77**, 893–896.
- [22] Slud, E. (1992). Nonparametric identifiability of marginal survival distributions in the presence of dependent competing risks and a prognostic covariate, in *Survival Analysis: State of the Art*, J.P. Klein & P.K. Goel, eds, Kluwer Academic Publishers, Boston, pp. 355–368.
- [23] Moeschberger, M.L. & Klein, J.P. (1996). Statistical methods for dependent competing risks, in *Lifetime Data Models in Reliability and Survival Analysis*, N.P. Jewell, A.C. Kimber, M.-L.T. Lee & G.A. Whitmore, eds, Kluwer Academic Publishers, Norwell, pp. 233–242.
- [24] Moeschberger, M.L. (1974). Life tests under dependent competing causes of failure, *Technometrics* **16**, 39–47.
- [25] Basu, A.P. & Klein, J.P. (1982). Some recent results in competing risks theory, in *Survival Analysis*, J. Crowley & R.A. Johnson, eds, Hayward, California, pp. 216–229.
- [26] Basu, A.P. & Ghosh, J.K. (1978). Identifiability of the multinormal distribution under competing risks models, *Journal of Multivariate Analysis* **8**, 413–429.
- [27] Basu, A.P. & Ghosh, J.K. (1980). Identifiability of distributions under competing risks and complementary risks models, *Communications in Statistics-Theory and Methods* **9**, 1515–1525.
- [28] Gail, M. (1975). A review and critique of some models used in competing risk analyses, *Biometrics* **31**, 209–222.
- [29] Williams, J.S. & Lagakos, S.W. (1977). Models for censored survival analysis: constant-sum and variable-sum models, *Biometrika* **64**, 215–224.
- [30] Dinse, G.E. (1982). Nonparametric estimates for partially-complete time and type of failure data, *Biometrics* **38**, 417–431.
- [31] Kodell, R.J. & Chen, J.J. (1987). Handling cause of death in equivocal cases using the EM algorithm, *Communications in Statistics – Theory and Methods* **16**, 2565–2585.
- [32] Racine-Poon, A.H. & Hoel, D.G. (1984). Nonparametric estimation of survival function when cause of death is uncertain, *Biometrics* **40**, 1151–1158.
- [33] Dinse, G.E. (1986). Nonparametric prevalence and mortality estimators for animal experiments with incomplete cause-of-death data, *Journal of the American Statistical Association* **81**, 328–336.
- [34] Dinse, G.E. (1988). Estimating tumor incidence rates in animal carcinogenicity experiments, *Biometrics* **44**, 405–415.
- [35] Lagakos, S.W. & Louis, T.A. (1988). Use of tumor lethality to interpret tumorigenicity experiments lacking cause-of-death data, *Applied Statistics* **37**, 169–179.
- [36] Goetghebeur, E. & Ryan, L. (1990). A modified logrank test for competing risks with missing failure type, *Biometrika* **77**, 207–211.
- [37] Dewanji, A. (1992). A note on the test for competing risks with missing failure type, *Biometrika* **79**, 855–857.
- [38] Goetghebeur, E. & Ryan, L. (1995). Analysis of competing risks survival data when some failure types are missing, *Biometrika* **82**, 821–833.
- [39] Lu, K. & Tsiatis, A. (2001). Multiple imputation methods for estimating regression coefficients in the competing risks model with missing cause of failure, *Biometrics* **57**, 1191–1197.
- [40] Flehinger, B.J., Reiser, B. & Yashchin, E. (1998). Survival with competing risks and masked causes of failures, *Biometrika* **85**, 151–164.
- [41] Dewanji, A. & Sengupta, D. (2003). Estimation of competing risks with general missing pattern in failure types, *Biometrics* **59**, 1063–1070.
- [42] Lu, K. & Tsiatis, A. (2005). Comparison between two partial likelihood approaches for the competing risks model with missing cause of failure, *Lifetime Data Analysis* **11**, 29–40.
- [43] Miyakawa, M. (1984). Analysis of incomplete data in competing risks model, *IEEE Transactions on Reliability* **33**, 293–296.

- [44] Usher, J.S. & Hodgson, T.J. (1988). Maximum likelihood analysis of component reliability using masked system life-test data, *IEEE Transactions on Reliability* **37**, 550–555.
- [45] Lin, D., Usher, J. & Guess, F. (1993). Exact maximum likelihood estimation using masked system data, *IEEE Transactions on Reliability* **42**, 631–635.
- [46] Guess, F.M., Usher, J.S. & Hodgson, T.J. (1991). Estimating system and component reliabilities under partial information on cause of failure, *Journal of Statistical Planning and Inference* **29**, 75–85.
- [47] Flehinger, B.J., Reiser, B. & Yashchin, E. (2001). Statistical analysis for masked data advances in reliability, in *Handbook of Statistics*, N. Balakrishnan & C. Radhakrishna Rao, eds, Elsevier/North-Holland, Vol. 18, pp. 499–522.
- [48] Craiu, R.V. & Reiser, B. (2006). Inference for the dependent competing risks model with masked causes of failure, *Lifetime Data Analysis* **12**(1), 21–33.
- [49] Craiu, R.V. & Lee, T.C.M. (2005). Model selection for the competing risks model with and without masking, *Technometrics* **47**(4), 457–467.
- [50] Craiu, R.V. & Duchesne, T. (2004a). Inference based on the EM algorithm for the competing risks model with masked causes of failure, *Biometrika* **91**(3), 543–558.
- [51] Schäbe, H. (1994). Nonparametric estimation of component lifetime based on masked system life test data, *Journal of the Royal Statistical Society B* **56**, 251–259.
- [52] Lin, D. & Guess, F.M. (1994). System life data analysis with dependent partial knowledge on the exact cause of system failure, *Microelectronics and Reliability* **34**, 535–544.
- [53] Guttman, I., Lin, D.K., Reiser, B. & Usher, J.S. (1995). Dependent masking and system life data analysis: Bayesian inference for two-component systems, *Lifetime Data Analysis* **1**, 87–100.
- [54] Craiu, R.V. & Duchesne, T. (2004b). Using EM and data augmentation for the competing risks model, in *Applied Bayesian Modeling and Causal Inference from an Incomplete-Data Perspective*, A. Gelman & X.L. Meng, eds, John Wiley & Sons.
- [55] Sen, A., Basu, S. & Banerjee, M. (2001). Statistical analysis of life-data with masked cause-of-failure, in *Handbook of Statistics, Advances in Reliability*, A.P. Basu & N. Balakrishnan, eds, Elsevier Sciences, Amsterdam, Vol. 20, pp. 523–540.
- [56] Mukhopadhyay, C. (2006). Maximum likelihood analysis of masked series system lifetime data, *Journal of Statistical Planning and Inference* **136**, 803–838.
- [57] Kuo, L. & Yang, T.Y. (2000). Bayesian reliability modeling for masked system lifetime data, *Statistics and Probability Letters* **47**, 229–241.
- [58] Mukhopadhyay, C. & Basu, S. (2007). Bayesian analysis of masked series system lifetime data, *Communication in Statistics: Theory and Applications*, **36**, 329–348.
- [59] Mukhopadhyay, C. & Basu, A.P. (1993). *Bayesian Analysis of Competing Risks: K Independent Exponentials*, Technical Report #516, Department of Statistics, The Ohio State University.
- [60] Reiser, B., Guttman, I., Lin, D.K.J., Guess, F.M. & Usher, J.S. (1995). Bayesian inference for masked system lifetime data, *Applied Statistics* **44**, 79–90.
- [61] Lin, D., Usher, J. & Guess, F. (1996). Bayes estimation of component-reliability from masked system-life data, *IEEE Transactions on Reliability* **45**, 233–237.
- [62] Mukhopadhyay, C. & Basu, A.P. (1997). Bayesian analysis of incomplete time and cause of failure data, *Journal of Statistical Planning and Inference* **59**, 79–100.
- [63] Berger, J.O. & Sun, D. (1993). Bayesian analysis of the ply-Weibull distribution, *Journal of the American Statistical Association* **82**, 1412–1418.
- [64] Basu, S., Basu, A.P. & Mukhopadhyay, C. (1999). Bayesian analysis of masked system failure data using non-identical Weibull models, *Journal of Statistical Planning and Inference* **78**, 255–275.
- [65] Basu, S., Sen, A. & Banerjee, M. (2003). Bayesian analysis of competing risks with partially masked cause-of-failure, *Journal of Royal Statistical Society. Series C: Applied Statistics* **52**, 77–93.
- [66] Nelson, W. (1982). *Applied Life Data Analysis*, John Wiley & Sons, New York.
- [67] Chan, V. & Meeker, W.Q. (1999). A failure-time model for infant mortality and wearout failure modes, *IEEE Transactions on Reliability* **TR-48**, 678–682.

Further Reading

Flehinger, B.J., Reiser, B. & Yashchin, E. (1996). Inference about defects in the presence of masking, *Technometrics* **38**, 247–255.

Related Articles

Accelerated Life Tests: Analysis with Competing Failure Modes; Competing Risks; Expectation Maximization Algorithm; Masked Failure Data: Bayesian Modeling.

SANJIB BASU

Nonparametric and Semiparametric Bayesian Reliability Analysis

Introduction

Suppose that T is a nonnegative random variable denoting lifetime. We define the survival function at time t to be $S(t) = P(T > t) = 1 - F(t)$, where F is the cumulative distribution function (**cdf**). If T has density $f(t)$, the hazard rate function $r(t)$ at time t is given by $r(t) = \frac{f(t)}{S(t)}$. This gives the instantaneous failure rate at time t given survival just prior to time t . The corresponding cumulative hazard at time t is $R(t) = \int_0^t r(s) ds$. It is well known that $S(t) = \exp(-R(t))$. Hence, knowing any one of density, cdf, survival function, hazard rate function, or cumulative hazard function is equivalent for inference about the properties of the life distribution. Sometimes, the phrase “reliability function” is used to denote the survival function of a component or a system of components connected in some structure.

In reliability and survival analysis, it is often of interest to estimate the reliability of the system/component from the observed lifetime data. Suppose n components are put to test and $T_1, \dots, T_n$ are the corresponding lifetimes. If all $T_1, \dots, T_n$ are actually observed, we have complete data. Usually, however, some of the T_i 's are censored. Suppose that T_i and C_i are the lifetime and (right) censoring times, respectively, of the i th component. One observes $X_i = \min(T_i, C_i)$ and $\delta_i = I_{T_i \leq C_i}$. On the basis of observed data (X_i, δ_i) , $i = 1, \dots, n$, one would like to infer about the lifetime distribution as a whole.

No matter whether one has complete or censored data, a host of techniques have been developed for their analysis. They can be broadly classified as parametric and nonparametric approaches. The advantages of parametric methods are their simplicity in the sense that only a handful of parameters can be used to explain the behavior. For example, Kvam and Samaniego [1] used a parametric approach for analyzing the life-history data from an r -out-of- k system whereby the component lifetimes were assumed to be exponentially distributed. The unknown parameter was estimated using the maximum-likelihood method. However, as is well known, parametric

models impose certain structural restrictions and are less than optimal if one is unable to verify/justify the model assumptions. An alternative is to use nonparametric methods such as the ones used by Kvam and Samaniego [1] or Chen [2] in estimating **component reliability**.

In this article, we first review some nonparametric and semiparametric Bayesian methods available for the analysis of life-history data. The advantage of using Bayesian methods is the ability to incorporate prior information in the inferential procedure, and semiparametric models allow for partial specification of model structure through parameters. We then present a semiparametric Bayesian approach for estimating the component lifetime distribution on the basis of the system lifetime data, along with additional information (such as the censoring indicator and the number of failed components) for multiple k -out-of- n systems [3].

In the section titled “Nonparametric Methods”, we briefly review past work on the nonparametric Bayesian analysis for survival data. In the section titled “Semiparametric Method”, we present a theoretical development of a semiparametric procedure. Under “Sampling Procedure”, we briefly discuss the sampling procedure that will be used to implement the model. In the section titled “Illustration”, we present the results of a simulation study based on data from Chen [2]. Finally, in the last section, we close with some concluding remarks.

Nonparametric Methods

In Bayesian analysis, it is very important to have as close to correct a likelihood function as possible. Often, the standard parametric models are not rich enough to capture the uncertainty in the observed data. In such situations, nonparametric Bayesian methods provide a more flexible alternative. Nonparametric Bayesian methods in reliability can be broadly classified into three groups, depending on the quantity on which one places a prior distribution: prior on the class of all distributions, prior on the class of all hazard rates, and prior on the class of all cumulative hazards. In each case, the underlying distribution is assumed to be free of any parameters and the prior information is combined with life-history data to obtain the corresponding posterior. Owing to the nonparametric nature, the **prior distributions** are taken

2 Nonparametric and Semiparametric Bayesian Analysis

to be stochastic processes. We provide a brief review of the three methods in the following subsections. For a comprehensive review of various nonparametric methods of estimating the survival function, see Ferguson *et al.* [4]. Also see Sinha and Dey [5] and Singpurwalla [6] for a detailed discussion on nonparametric approaches to reliability **estimation**.

Prior on Distributions

Dirichlet processes, introduced by Ferguson [7] are one of the fundamental concepts in the nonparametric Bayesian analysis literature. A Dirichlet process can be used to provide a nonparametric prior for a distribution function.

Suppose that F is a cdf. We say that F has a Dirichlet process prior, i.e. $F \sim \mathcal{D}(M, F_0)$ if the following happens. For any partition $A_1 \cup A_2 \cup \dots \cup A_k = \mathfrak{R}$,

$$(F(A_1), F(A_2), \dots, F(A_k)) \sim \text{Dirichlet}(MF_0(A_1), MF_0(A_2), \dots, MF_0(A_k)) \quad (1)$$

where $\text{Dirichlet}(\alpha_1, \dots, \alpha_k)$ denotes a Dirichlet distribution with parameters $(\alpha_1, \dots, \alpha_k)$. See Wilks [8] for more on Dirichlet distribution. $M > 0$ is called the *precision parameter* and F_0 is called the *baseline distribution*. F_0 can be thought of as the “average value” of F and M as the amount of “concentration” of F around F_0 . A high value of M signifies that F is very close to F_0 and a low value of M signifies large dispersion around F_0 . However, see Sethuraman and Tiwari [9] for another interpretation of small values of M .

Ferguson [7] showed that if $\theta|F \sim F$ and $F \sim \mathcal{D}(M, F_0)$, then $F|\theta \sim \frac{1}{M+1}[MF_0 + \delta_\theta]$. Here, δ_θ denotes the degenerate part of the distribution at a specific value of θ . Repeated application of this result shows that if $T_1, \dots, T_n$ are observations generated by F , and F has a Dirichlet process prior $\mathcal{D}(M, F_0)$, the Bayes estimator of F is

$$\hat{F}(t) = \frac{M}{M+n}F_0(t) + \frac{1}{M+n}F_n(t) \quad (2)$$

where $F_n(t)$ is the empirical distribution function of the sample.

The above requires complete information about the lifetimes (i.e., no censoring). Susarla and Van Ryzin [10] developed the nonparametric Bayes estimator of the cdf in the presence of right censoring and assuming a Dirichlet process prior on the underlying distribution. Let $(X_i, \delta_i), i = 1, \dots, n$ be the observed data with $X_i = \min(T_i, C_i)$ and $\delta_i = I_{T_i \leq C_i}$. Let $u_1 < \dots < u_k$ be the distinct values among $X_1, \dots, X_n$ and λ_j be the number of censored observations at u_j . Let $k(t)$ be the number of u_j 's that are less than or equal to u_k and h_k be the number greater than u_k . Then, the Bayes estimator of the survival function under squared error loss is given by

$$\frac{M(1 - F_0(t)) + h_k(t)}{M + n} \prod_{j=1}^{k(t)} \frac{M(1 - F_0(u_j)) + h_j + \lambda_j}{M(1 - F_0(u_j)) + h_j} \quad (3)$$

This estimator reduces to the Kaplan–Meier estimator as $M \rightarrow 0$, and under no censoring, reduces to the estimator given earlier.

Dirichlet processes give rise to distributions as priors that are discrete with probability one. In addition, the corresponding Bayes estimator gets complicated for right-censored data. To avoid these difficulties, one may use neutral-to-the-right (NTR) priors. A random distribution function F on the real line is said to be NTR if for every m and $t_1 < t_2 < \dots < t_m$, there exist independent random variables $V_1, \dots, V_m$ such that $(1 - F(t_1), \dots, 1 - F(t_m))$ has the same distribution as $(V_1, V_1 V_2, \dots, \prod_{i=1}^m V_i)$. Doksum [11] and Ferguson and Phadia [12] show that if $X_1, \dots, X_n$ is a sample from F and F is NTR, then the posterior distribution of F given $X_1, \dots, X_n$ is also NTR. It turns out that the censored case is simpler to treat than the uncensored case. The implementation of the results for practical applications gets cumbersome. Damien *et al.* [13] and Walker and Damien [14] have proposed simulation-based approaches for a full Bayesian analysis involving NTR processes.

Prior on Hazard Functions

An alternative approach is to put a nonparametric prior on the class of all hazard functions. Suppose we partition the time axis into $(k+1)$ intervals $0 < t_1 < \dots < t_k < \infty$. Let $p_1 = F(t_1)$, $p_i =$

$F(t_i) - F(t_{i-1})$, $i = 2, \dots, k$. Note that $F(0) = 0$ and $F(\infty) = 1$. Define $Z_1 = p_1$, $Z_i = p_i / (1 - p_1 - \dots - p_{i-1})$ for $i = 2, \dots, k$. Note that Z_i is the failure rate over the interval $[t_{i-1}, t_i)$ and $(Z_1, \dots, Z_k)$ gives the piecewise constant hazard function over $[0, t_k)$. Let n_i denote the number of failures in the interval $[t_{i-1}, t_i)$, $\mathbf{p} = (p_1, \dots, p_k)$ and $\mathbf{d} = (n_1, \dots, n_{k+1})$. The observed data is given by $\mathbf{d}$.

In such a scenario, one may assume that the prior distribution of the Z_i 's is independent, $\beta(v_{1i}, v_{2i})$. This results in a generalized Dirichlet prior for $\mathbf{p}$ given by the probability function

$$f(p_1, \dots, p_k) = \prod_{i=1}^k \frac{\Gamma(v_{1i} + v_{2i})}{\Gamma(v_{1i})\Gamma(v_{2i})} p_i^{v_{1i}-1} \times (1 - p_1 - \dots - p_k)^{\gamma_i} \quad (4)$$

where $\gamma_i = v_{2i} - v_{1,i+1} - v_{2,i+1}$ for $i = 1, \dots, k-1$ and $\gamma_k = v_{2k} - 1$. See Basu and Tiwari [15] for the special case when the generalized Dirichlet prior for $\mathbf{p}$ reduces to the Dirichlet prior. Given the observed count data $\mathbf{d}$, the posterior of the piecewise hazards turns out also to be generalized Dirichlet. Lochner [16] and Tiwari and Rao [17] further discuss the use of this approach to estimate F . This method has several disadvantages. First, it uses count data, leading to loss of information by not using the actual failure times. Second, the inferential conclusions depend on whether the cells are combined at the prior or at the posterior stage. In addition, one needs an excessive number of parameters to specify the generalized Dirichlet distribution. See also Wilks [8] for more on generalized Dirichlet distributions.

The above procedure cannot account for any structural pattern in the hazard rate function. Padgett and Wei [18] and Arjas and Gasbarra [19] take the prior on the hazard rate function to be a **Poisson process** with constant jump size. Mazzuchi and Singpurwalla [20] take the prior on hazard rate function to be ordered Dirichlet. The former ensures that the hazard rate function is nondecreasing and the latter ensures that it is monotone.

Dykstra and Laud [21] introduce an extended gamma process and use it to model the prior distribution of a nondecreasing hazard rate process. They show that when $t_1, \dots, t_n$ are the right-censoring times of n observations and $\Gamma(a(s), \beta(s))$ is the prior on the hazard rate process; the posterior on

the hazard rate process is also an extended gamma process $\Gamma(a(s), \hat{\beta}(s))$ where

$$\hat{\beta}(s) = \frac{\beta(s)}{1 + \beta(s) \sum_{i=1}^n (t_i - s)^+} \quad (5)$$

with $a^+ = \max(a, 0)$. When one observes the actual failure times $t_1, \dots, t_n$, the resulting posterior hazard rate is a mixture of extended gamma processes, which is complicated to calculate. Laud *et al.* [22] approximate this posterior by approximating its random independent increments *via* a Gibbs sampler. Once the posterior hazard rate process is obtained, the corresponding survival function is given by

$$P(T \geq t) = \exp \left[- \int_0^t \log(1 + \hat{\beta}(s)(t - s)) da(s) \right] \quad (6)$$

Other methods of specifying the prior on hazard rates include use of Markov beta and Markov gamma processes considered by Nieto-Barajas and Walker [23]. See also Tiwari and Rao [17], Tiwari and Jammalamadaka [24], and Tiwari and Kumar [25].

Prior on Cumulative Hazard

An alternative to putting a prior on the hazard rate function is to put a prior on the cumulative hazard function. This is particularly attractive especially when there is no density, since the cumulative hazard still exists in such situations. Kalbfleisch [26] proposed using a gamma process as a prior for the cumulative hazard function $R(t)$. Take any partition $0 \equiv t_0 < t_1 < t_2 < \dots < t_{k-1} < t_k \equiv \infty$, of the time points and define $r_i = \log(1 - Z_i)$ where $Z_i = P(T \in [t_{i-1}, t_i) | T \geq t_{i-1})$ is the hazard rate over the interval $[t_{i-1}, t_i)$. Assume that r_i 's are independent gamma random variables with shape $\alpha_i - \alpha_{i-1}$ and scale c , where $\alpha_i = cR^*(t_i)$ and $R^*(t)$ is interpreted as the best guess of $R(t)$. c is the measure of strength of conviction about the guess and large values indicate a strong conviction.

Given n failure times $\tau_1, \dots, \tau_n$, the posterior cumulative hazard will be a process with independent increments. Kalbfleisch [26] has shown that the posterior cumulative hazard has an increment at τ_i

4 Nonparametric and Semiparametric Bayesian Analysis

which is given by a density $A(c + A_i, c + A_{i+1})$ at u , where $A_i = n - i$ and $A(a, b)$ is of the form

$$\frac{\exp(-bu) - \exp(au)}{u \log(a/b)}$$

Between τ_{i-1} and τ_i , the increments are prescribed by a gamma process with shape function $cR^*(\cdot)$ and scale $c + A_i$. The survival function is recovered either by simulation or by approximation using expected value of the process.

The above formulation suffers from some difficulties. First, there is a lack of intuition regarding the assumption of gamma distribution on r_i . Second, the independent increments property of $R(t)$ may not be meaningful, since under aging and wear, the successive Z_i 's would be judged to be increasing. Finally, presence of ties in the failure time data presents problems in the model fitting.

A different model was proposed by Hjort [27] to model the randomness in the cumulative hazard rate. Suppose R has a beta process prior with parameters $c(\cdot)$ and $R_0(\cdot)$. The posterior of R given life-history data $(X_1, \delta_1), \dots, (X_n, \delta_n)$ is also a beta process of the form

$$R|\text{data} \sim \beta \left[c(\cdot) + Y(\cdot), \int_0^{\cdot} \frac{c \, dR_0 + dN}{c + Y} \right] \quad (7)$$

where

$$N(t) = \sum_{i=1}^n I(X_i \leq t, \delta_i = 1) \quad (8)$$

and

$$Y(t) = \sum_{i=1}^n I(X_i \geq t) \quad (9)$$

are two counting processes derived from the data. The Bayes estimator of $R(t)$ under squared error loss is

$$\hat{R}(t) = \int_0^t \frac{c \, dR_0 + dN}{c + Y} \quad (10)$$

and the corresponding estimator of F is

$$\hat{F}(t) = 1 - \prod_{[0,t]} \left[1 - \frac{c \, dR_0 + dN}{c + Y} \right] \quad (11)$$

As $c(\cdot)$ decreases to zero, $\hat{F}$ tends to the Kaplan–Meier estimator.

Semiparametric Method

Often, it may be necessary to impose certain structural restrictions on the underlying survival/reliability model to aid in the physical understanding and interpretation, all the while maintaining considerable generality. In such situations, one may decide to use semiparametric models, whereby parts of the model are parametric (reflecting the desired structural restriction) and the remainder nonparametric. An example of this is the **Cox** model, introduced by Cox [28]. To perform Bayesian analysis, the nonparametric part is assumed to be a realization of a stochastic process. As before in the nonparametric Bayesian analysis case, one can put nonparametric priors on the hazard rates. Different methods have been discussed in Sinha and Dey [5] for dealing with these models.

Recently, Merrick *et al.* [29] developed a semiparametric Bayesian proportional hazards model for reliability and maintenance of machine tools. Their proposed model uses a mixture of Dirichlet process (MDP) prior for the baseline failure rate in the proportional hazards model. Such priors were introduced by Antoniak [30] and have been popularized by various authors such as MacEachern [31] and West *et al.* [32]. Another application of Bayesian semiparametric models in reliability is presented in Merrick *et al.* [33], where the authors use a proportional intensities model and present a semiparametric Bayesian approach for nonhomogeneous Poisson processes. Apart from the above, we were unable to find Bayesian semiparametric models in the context of reliability analysis. This emphasizes the fact that while Bayesian semiparametric models have been popular in survival analysis, they have not been as popular in (engineering) reliability analysis. Below, we present a new model to estimate component reliability using system reliability data from multiple r -out-of- k systems.

In industrial and biological problems, one often has a multicomponent system (consisting of, say, k components) and observes the time of system failure. Owing to the system architecture, failure of the system occurs if and only if at least a certain number of components (say, r) fail. For example, a liquid crystal display (LCD) may be said to correctly function when at least 80% of its pixels function properly. Such a system is called an r -out-of- k system, special cases of which are a series system

($r = 1$) and a parallel system ($r = k$). r -out-of- k systems have been well studied in the reliability context and are a favorite way of increasing system **redundancy**. See Hoyland and Rausand [34] for more details on r -out-of- k systems and associated examples.

Suppose that we have an r -out-of- k system that consists of k components with identical life distributions that act independent of each other. Assume that the component life distribution is given by $F(\cdot|\theta)$ where θ is an unknown p -dimensional parameter. Suppose that T is the system failure time and C is the censoring time. We observe $X = T \wedge C$ and $\delta = I(X = T)$, the censoring indicator. When $\delta = 1$, X is distributed as the r th-order statistic based on a sample of size k from $F(\cdot|\theta)$. When $\delta = 0$ (i.e., the system was alive and censored at X), we may or may not have information on the number of components $s(< r)$ in the system that have failed. Let $\gamma = 1$ indicate that s is observed and $\gamma = 0$ indicate that s is unobserved. Assuming that $F(\cdot|\theta)$ is absolutely continuous with respect to the Lebesgue measure, the likelihood contribution of the system is

$$L(\theta) = \left[r \binom{k}{r} F^{r-1}(x|\theta) f(x|\theta) \{1 - F(x|\theta)\}^{k-r} \right]^\delta \left[\binom{k}{s} F^s(x|\theta) \{1 - F(x|\theta)\}^{k-s} \right]^{(1-\delta)\gamma} \left[\sum_{s=0}^{r-1} \binom{k}{s} F^s(x|\theta) \{1 - F(x|\theta)\}^{k-s} \right]^{(1-\delta)(1-\gamma)} \quad (12)$$

When $\delta = 1$, we take $\gamma = 1$ and $s = r$.

Suppose that for $i = 1, \dots, n$, we have m_i copies of an r_i -out-of- k_i system. The m_i copies are assumed to be independently distributed with the same component distribution $F(\cdot|\theta_i)$. For $i \neq j$, the systems are also assumed to be conditionally independent of each other and the parameters θ_i and θ_j may or may not be equal.

The resulting data X is presented in Table 1.

Denoting $\theta_n = (\theta_1, \dots, \theta_n)$, the likelihood is given by

$$L(\theta_n|X) = \prod_{i=1}^n \prod_{j=1}^{m_i} \left[r_i \binom{k_i}{r_i} \{F(x_{ij}|\theta_i)\}^{r_i-1} f(x_{ij}|\theta_i) \times \{1 - F(x_{ij}|\theta_i)\}^{k_i-r_i} \right]^{\delta_{ij}} \times \left[\binom{k_i}{s_{ij}} \{F(x_{ij}|\theta_i)\}^{s_{ij}} \{1 - F(x_{ij}|\theta_i)\}^{k_i-s_{ij}} \right]^{(1-\delta_{ij})\gamma_{ij}} \times \left[\sum_{s=0}^{r_i-1} \binom{k_i}{s} \{F(x_{ij}|\theta_i)\}^s \{1 - F(x_{ij}|\theta_i)\}^{k_i-s} \right]^{(1-\delta_{ij})(1-\gamma_{ij})} \quad (13)$$

We assume that the θ_i 's are independent and identically distributed from a distribution G with a Dirichlet process prior having baseline distribution G_0 and precision M . Hence the prior distribution of $\theta_1, \dots, \theta_n$ assuming M and G_0 are known is given by (see Antoniak [30], Blackwell and MacQueen [35]):

$$\pi(\theta_n) = \prod_{i=1}^n \left[\frac{M G_0(d\theta_i) + \sum_{j < i} \delta_{\theta_j}(d\theta_i)}{M + i - 1} \right] \quad (14)$$

As mentioned earlier, under the Dirichlet process setup, some of the system parameters θ_i may be identical. This is a reflection of the fact that some of the systems may be built using components from the same manufacturer and thus would have similar behavior.

Our goal will be to estimate the reliability of a component using the predictive approach. Assuming that a future system has a parameter θ_{n+1} and that

Table 1 Failure time data from several r -out-of- k systems

System	r	k	Observations		
1	r_1	k_1	$(X_{11}, \delta_{11}, \gamma_{11}, s_{11})$	$\dots$	$(X_{1m_1}, \delta_{1m_1}, \gamma_{1m_1}, s_{1m_1})$
2	r_2	k_2	$(X_{21}, \delta_{21}, \gamma_{21}, s_{21})$	$\dots$	$(X_{2m_2}, \delta_{2m_2}, \gamma_{2m_2}, s_{2m_2})$
$\dots$	$\dots$	$\dots$	$\dots$	$\dots$	$\dots$
n	r_n	k_n	$(X_{n1}, \delta_{n1}, \gamma_{n1}, s_{n1})$	$\dots$	$(X_{nm_n}, \delta_{nm_n}, \gamma_{nm_n}, s_{nm_n})$

X_{n+1} is the lifetime of a component of the system, we want

$$\begin{aligned}
S(x) &= P(X_{n+1} > x | \mathbf{X}) \\
&= \int P(X_{n+1} > x | \theta_{n+1}, \mathbf{X}, \boldsymbol{\theta}_n) \\
&\quad \times f(\theta_{n+1} | \boldsymbol{\theta}_n, \mathbf{X}) f(\boldsymbol{\theta}_n | \mathbf{X}) d\boldsymbol{\theta}_{n+1} \\
&= \int \bar{F}_{n+1}(x | \theta_{n+1}) \frac{1}{M+n} [MG_0(d\theta_{n+1}) \\
&\quad + \sum_{j=1}^n \delta_{\theta_j}(d\theta_{n+1})] f(\boldsymbol{\theta}_n | \mathbf{X}_n) d\boldsymbol{\theta}_n \\
&= \frac{1}{M+n} \int [M \int \bar{F}(x | \theta) G_0(d\theta) \\
&\quad + \sum_{j=1}^n \bar{F}(x | \theta_j)] f(\boldsymbol{\theta}_n | \mathbf{X}) d\boldsymbol{\theta}_n \quad (15)
\end{aligned}$$

Note that the complicated nature of the likelihood precludes any conjugate choice of baseline prior to simplify the posterior distribution calculations.

Ferguson [36], Kuo [37], and Tiwari and Kumar [25] have used this type of mixture approach for estimating parameters such as density function and reliability. We will use the Gibbs sampler to sample from the posterior and draw inferences. The Gibbs sampler is difficult to implement as is, since the likelihood contributions involve calculations of the cdf, the **probability density function**, and/or sums involving them. While several algorithms have been proposed to deal with such non-conjugate setups for sampling from the posterior in an MDP setting [38], they are computationally intensive.

The problem arises because we have several unobserved lifetimes, essentially making data incomplete. Here, we introduce a data-augmentation technique (see Tanner and Wong [39], Tanner [40]) whereby the observed data are augmented to get the “complete” data $\tilde{\mathbf{X}} = \{X_{ijl}\}$, i.e. the exact failure times of all the components in a system. The likelihood based on the augmented data can then be written as

$$\tilde{L}(\boldsymbol{\theta}_n | \tilde{\mathbf{X}}) = \prod_{i=1}^n \prod_{j=1}^{m_i} \prod_{l=1}^{k_i} f(x_{ijl} | \theta_i) \quad (16)$$

Similar techniques were also used in Kim and Arnold [41]. This simplifies calculations by taking advantage of the conjugate structure of the likelihood and the baseline prior. Note that the posterior distribution of the number of failed components s_{ij} for systems in which such data are missing arises naturally and can be interpreted as the posterior based on a uniform prior on $\{0, \dots, k_i\}$.

Sampling Procedure

1. Generate a set of starting values of $\theta_1, \dots, \theta_n$.
2. If $\delta_{ij} = \gamma_{ij} = 0$, sample s_{ij} from binomial($k_i, F(X_{ij} | \theta_i)$) truncated to $\{0, \dots, r_i - 1\}$.
3. If $\delta_{ij} = 1$, generate $\{X_{ij(l)}\}_{l=1}^{r_i-1}$ as order statistics from

$$g_L(x | X_{ij}, \theta_i) = \frac{f(x | \theta_i)}{F(X_{ij} | \theta_i)}, \quad 0 < x < X_{ij} \quad (17)$$

Also generate $\{X_{ij(l)}\}_{l=r+1}^{k_i}$ as order statistics from

$$g_U(x | X_{ij}, \theta_i) = \frac{f(x | \theta_i)}{1 - F(X_{ij} | \theta_i)}, \quad X_{ij} < x < \infty \quad (18)$$

Also, set $X_{ij(r)} \equiv X_{ij}$.

Note that $g_L(\cdot | X_{ij}, \theta_i)$ is the density of observations from $F(\cdot | \theta_i)$ truncated above at X_{ij} . Similarly, $g_U(\cdot | X_{ij}, \theta_i)$ is the density for observations that are truncated below.

4. If $\delta_{ij} = 0$, generate $\{X_{ij(l)}\}_{l=1}^{s_{ij}}$ as order statistics from $g_L(\cdot | X_{ij}, \theta_i)$ and also generate $\{X_{ij(l)}\}_{l=s_{ij}+1}^{k_i}$ as order statistics from $g_U(\cdot | X_{ij}, \theta_i)$.
5. Having generated the “complete data” $\{X_{ijl}\}$, we update the θ ’s conditionally as

$$\begin{aligned}
\theta_i | \text{rest} &\sim \prod_{j=1}^{m_i} \prod_{l=1}^{k_i} f(X_{ijl} | \theta_i) \\
&\times \left[\frac{MG_0(d\theta_i) + \sum_{j \neq i} \delta_{\theta_j}(d\theta_i)}{M+n-1} \right] \quad (19)
\end{aligned}$$

This is done using the Gibbs sampler. Note that in the above expression, we need the exact failure times X_{ijl} , not the ordered values $X_{ij(l)}$. If the sufficient statistic does not depend on the ordering of the observations, one can replace the exact values by the ordered values. Such is the case, for example, when

one is interested in inferring the mean of a normal or the scale of a Weibull distribution.

6. Go back to 2 and repeat until convergence to the steady state.

Below, we give a special case for illustration:

Lognormal Distribution

If the component failure times are distributed as lognormal(μ, σ^2) ($LN(\mu, \sigma^2)$), we can generate samples from the truncated distributions given by $g_L(\cdot)$ and $g_U(\cdot)$ using

$$X_L = \exp \left[\mu + \sigma \Phi^{-1} \left(U \Phi \left(\frac{\log X - \mu}{\sigma} \right) \right) \right] \quad (20)$$

and

$$X_U = \exp \left[\mu + \sigma \Phi^{-1} \left(U + (1 - U) \Phi \left(\frac{\log X - \mu}{\sigma} \right) \right) \right] \quad (21)$$

where $U \sim U(0, 1)$. Note that the random variables X_L and X_U above are generated by inverting the cdf's corresponding to g_L and g_U , respectively.

In this case,

$$\prod_{j=1}^{m_i} \prod_{l=1}^{k_i} f(x_{ij(l)} | \theta_i) = \frac{1}{(2\pi\sigma^2)^{m_i k_i / 2}} \frac{1}{\prod_{j,l} x_{ij(l)}} \times \exp \left[-\frac{1}{2\sigma^2} \sum_{j,l} (\log x_{ij(l)} - \mu_i)^2 \right] \quad (22)$$

Note that $\sum_{j=1}^{m_i} \sum_{l=1}^{k_i} \log x_{ij(l)} = \sum_{j=1}^{m_i} \sum_{l=1}^{k_i} \log x_{ijl}$ is the sufficient statistic and is not order dependent.

Illustration

We used the system configuration given in Table 2 to generate our data [see 2].

To generate an observation on an r_i -out-of- k_i system, a simple random sample of size k_i is generated from F and another simple random sample of size k_i is generated from G . Let $u_{i(r_i)}$ and $v_{i(r_i)}$ be

Table 2 Configuration of several r -out-of- k systems

i	1	2	3	4	5	6
r_i	1	3	5	1	4	7
k_i	5	5	5	7	7	7
m_i	10	9	11	12	10	13

the r_i th-order statistics from the two samples, respectively. If $u_{i(r_i)} \leq v_{i(r_i)}$, the generated observation is taken as $(X_i, \delta_i, \gamma_i, s_i) = (u_{i(r_i)}, 1, 1, r_i)$. Otherwise, let r be the rank of the largest-order statistic from F that is smaller than $v_{i(r_i)}$. With 90% probability, the generated observation is taken as $(X_i, \delta_i, \gamma_i, s_i) = (v_{i(r_i)}, 0, 1, r)$ and with 10% probability, the generated observation is taken as $(X_i, \delta_i, \gamma_i, s_i) = (v_{i(r_i)}, 0, 0, 0)$. We used lognormal distribution with $\mu_{01} = 1$, $\sigma_{01} = 1$ as the F and a lognormal distribution with $\mu_{02} = \{0.8, 2.8\}$, $\sigma_{02} = 1$ as G . This procedure was seen to give rise to 68% censoring when $\mu_{02} = 0.8$ and about 0% censoring when $\mu_{02} = 2.8$. As indicated earlier, for a censored lifetime, the number of failed components was noted with probability 0.9 in each of the censoring schemes.

Once the dataset was generated according to the above procedure, we estimated the reliability of a single component $S(x)$ as outlined in equation (15). We assumed that the underlying population is $LN(\mu, 1)$ with $\mu \sim \mathcal{D}(M, G_0)$ and $G_0 \sim N(1, 1)$. We also assumed $M \sim \gamma(0.1, 0.1)$ and ran a full Bayesian approach. The updating of M was done along the lines of Escobar and West [42] and for improved mixing, the μ_i 's were updated using "Algorithm 3" in Neal [38]. The predictive reliability is based on 5000 runs of the Gibbs sampler with a thinning of 10, obtained after a **burn-in** of 5000 iterations. Convergence was ascertained using CODA [43]. The results are presented in Figure 1. The precision of the estimate is measured by RMSE (root mean square error), which is the RMSE of the estimate at selected points and is given in Table 3.

We see that as the percentage of censoring increases, the estimated values get farther away from

Table 3 Accuracy of the estimates based on data from various censoring schemes

μ_{02}	RMSE	Censoring proportion (%)
0.8	0.055	67.69
2.8	0.048	0

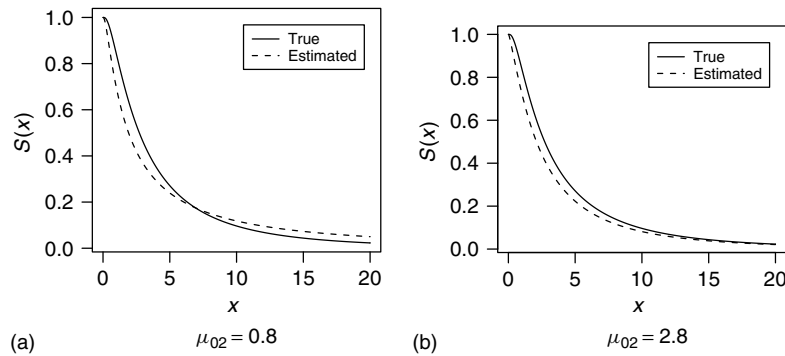


Figure 1 Estimated reliability of a component from the system configuration in Table 2 using different censoring distributions

the true values. In the case where $\mu_{02} = 0.8$, we also kept track of the true (but unobserved) number of component failures at the data generation stage and the estimated values at the simulation stage. The average number of failures turned out to be (0, 0, 2.18, 2.77, 5.70, 5.97), while the true values are (0, 0, 2, 1, 6, 6).

Conclusion

We have presented a brief overview of nonparametric and semiparametric Bayes methods in lifetime data analysis. In the semiparametric method outlined here, one can, as a by-product, infer the number of failed components for a censored system lifetime when it is unobserved. The method of data augmentation can be used when the sufficient statistic is not order dependent—otherwise one can always use the original likelihood and use one of the nonconjugate sampling methods outlined in Neal [38]. This discussion is not exhaustive but is intended to give a flavor of the current state of the art. Owing to advances in computing, realistic models, which were once avoided because of difficulty in implementation, will be becoming more popular.

References

- [1] Kvam, P.H. & Samaniego, F.J. (1993). On maximum likelihood estimation based on ranked set samples with applications to reliability, in *Advances in Reliability*, Elsevier Science B. V., The Netherlands, pp. 215–229.
- [2] Chen, Z. (2003). Component reliability analysis of k -out-of- n systems with censored data, *Journal of Statistical Planning and Inference* **116**, 305–315.
- [3] Boyles, R.A. & Samaniego, F.J. (1987). On estimating component reliability for systems with random redundancy levels, *IEEE Transactions in Reliability* **R-36**, 403–407.
- [4] Ferguson, T.S., Phadia, E.G. & Tiwari, R.C. (1992). Bayesian nonparametric inference, in *Current Issues in Statistical Inference: Essays in Honor of D. Basu, No. 17 in IMS Lecture Notes Monograph Series*, Institute of Mathematical Statistics, Hayward, pp. 127–150.
- [5] Sinha, D. & Dey, D.K. (1997). Semiparametric Bayesian analysis of survival data, *Journal of the American Statistical Association* **92**, 1195–1212.
- [6] Singpurwalla, N.D. (2006). *Reliability and Risk: A Bayesian Perspective*, John Wiley & Sons, New York.
- [7] Ferguson, T.S. (1973). A Bayesian analysis of some nonparametric problems, *The Annals of Statistics* **1**, 209–230.
- [8] Wilks, S.S. (1962). *Mathematical Statistics*, John Wiley & Sons, New York.
- [9] Sethuraman, J. & Tiwari, R.C. (1982). Convergence of Dirichlet measures and the interpretation of their parameter, in *Statistical Decision Theory and Related Topics III*, Academic Press, Vol. 2, pp. 305–315.
- [10] Susarla, V. & VanRyzin, J. (1976). Nonparametric Bayesian estimation of survival curves from incomplete observations, *Journal of the American Statistical Association* **71**, 740–754.
- [11] Doksum, K.A. (1974). Tailfree and neutral random probabilities and their posterior distributions, *The Annals of Probability* **2**, 183–201.
- [12] Ferguson, T.S. & Phadia, E.G. (1979). Bayesian nonparametric estimation based on censored data, *The Annals of Statistics* **7**, 163–176.
- [13] Damien, P., Laud, P.W. & Smith, A.F.M. (1995). Approximate random variate generation from infinitely divisible distributions with applications to Bayesian inference, *Journal of the Royal Statistical Society. Series B* **57**, 547–563.

- [14] Walker, S. & Damien, P. (1998). A full Bayesian non-parametric analysis involving neutral to the right process, *Scandinavian Journal of Statistics* **25**, 669–680.
- [15] Basu, D. & Tiwari, R.C. (1982). A note on the Dirichlet process, in *Statistics and Probability: Essays in Honor of C. R. Rao*, North-Holland, New York, pp. 89–103.
- [16] Lochner, R.H. (1975). A generalized Dirichlet distribution in Bayesian life testing, *Journal of the Royal Statistical Society. Series B* **37**, 103–113.
- [17] Tiwari, R.C. & Rao, J.S. (1983). Bayesian nonparametric estimation of failure rates with censored data, *Calcutta Statistical Association Bulletin* **32**, 79–90.
- [18] Padgett, W.J. & Wei, L.J. (1981). A Bayesian non-parametric estimator of survival probability assuming increasing failure rate, *Communications in Statistics: Theory and Methods* **A10**, 49–63.
- [19] Arjas, E. & Gasbarra, D. (1994). Nonparametric Bayesian inference from right-censored survival data using the Gibbs sampler, *Statistica Sinica* **4**, 505–524.
- [20] Mazzuchi, T.A. & Singpurwalla, N.D. (1985). A Bayesian approach to inference for monotone failure rates, *Statistics and Probability Letters* **3**, 135–142.
- [21] Dykstra, R.L. & Laud, P. (1981). A Bayesian approach to reliability, *The Annals of Statistics* **9**, 356–367.
- [22] Laud, P.W., Smith, A.F.M. & Damien, P. (1996). Monte Carlo methods for approximating a posterior hazard rate process, *Statistics and Computing* **6**, 77–83.
- [23] Nieto-Barajas, L.E. & Walker, S.G. (2002). Markov beta and gamma processes for modelling hazard rates, *Scandinavian Journal of Statistics* **29**, 413–424.
- [24] Tiwari, R.C. & Jammalamadaka, S.R. (1985). Estimation of survival function and failure rates with censored data, *Statistics* **16**, 535–540.
- [25] Tiwari, R.C. & Kumar, S. (1989). Bayes reliability estimation under a random environment governed by a Dirichlet prior, *IEEE Transactions on Reliability* **R-37**, 218–223.
- [26] Kalbfleisch, J.D. (1978). Non-parametric Bayesian analysis of survival time data, *Journal of the Royal Statistical Society. Series B* **40**, 214–221.
- [27] Hjort, N.L. (1990). Nonparametric Bayes estimators based on beta processes in models for life history data, *The Annals of Statistics* **18**, 1259–1294.
- [28] Cox, D.R. (1972). Regression models and life tables, *Journal of the Royal Statistical Society. Series B* **34**, 187–220.
- [29] Merrick, J.R.W., Soyer, R. & Mazzuchi, T.A. (2003). A Bayesian semiparametric analysis of the reliability and maintenance of machine tools, *Technometrics* **45**, 58–69.
- [30] Antoniak, C.E. (1974). Mixtures of Dirichlet processes with applications to Bayesian nonparametric problems, *The Annals of Statistics* **2**, 1152–1174.
- [31] MacEachern, S.N. (1994). Estimating normal means with a conjugate-style Dirichlet process prior, *Communications in Statistics: Simulation and Computation* **23**, 727–741.
- [32] West, M., Muller, P. & Escobar, M.D. (1994). Hierarchical priors and mixture models with application in regression density estimation, in *Aspects of Uncertainty: A Tribute to D. V. Lindley*, John Wiley & Sons, London, pp. 363–368.
- [33] Merrick, J.R., Soyer, R. & Mazzuchi, T.A. (2005). Are maintenance practices for railroad tracks effective? *Journal of the American Statistical Association* **100**, 17–25.
- [34] Hoyland, A. & Rausand, M. (1994). *System Reliability Theory: Models and Statistical Methods*, John Wiley & Sons, New York.
- [35] Blackwell, D. & MacQueen, J.B. (1973). Ferguson distribution via Polya urn schemes, *The Annals of Statistics* **1**, 353–355.
- [36] Ferguson, T.S. (1983). Bayesian density estimation by mixtures of normal distributions, in *Recent Advances in Statistics: Papers in Honor of Herman Chernoff on his Sixtieth Birthday*, Academic Press, New York; London, pp. 287–302.
- [37] Kuo, L. (1986). Computations of mixtures of Dirichlet processes, *SIAM Journal of Scientific Statistical Computing* **7**, 60–71.
- [38] Neal, R.M. (2000). Markov chain sampling methods for Dirichlet process mixture models, *Journal of Computational and Graphical Statistics* **9**, 249–265.
- [39] Tanner, M.A. & Wong, W.H. (1987). The calculation of posterior distributions by data augmentation, *Journal of the American Statistical Association* **82**, 528–540.
- [40] Tanner, M.A. (1993). *Tools for Statistical Inference*, Springer, New York.
- [41] Kim, Y. & Arnold, B.C. (1999). Parameter estimation under generalized ranked set sampling, *Statistics and Probability Letters* **42**, 353–360.
- [42] Escobar, M.D. & West, M. (1995). Bayesian density estimation and inference using mixtures, *Journal of the American Statistical Association* **90**, 577–588.
- [43] Best, N.G., Cowles, M.K. & Vines, K. (1995). *CODA: Convergence Diagnosis and Output Analysis Software for Gibbs Sampling Output, Version 0.30*, Technical Report, MRC Biostatistics Unit, University of Cambridge.

Related Articles

Cumulative Damage Models Based on Gamma Processes; Dirichlet Process, Simulation of; k -out-of- n Systems; Markov Chain Monte Carlo, Introduction; Nonparametric Methods for Analysis of Repair Data; Poisson Processes; Survival Analysis, Nonparametric.

KAUSHIK GHOSH AND RAM C. TIWARI

Degradation Models

Introduction

Reliability testing typically generates product lifetime data, but for some tests, **covariate** information about the wear and tear on the product during the **life test** can provide additional insight into the product's lifetime distribution. This usage, or **degradation**, can be the physical parameters of the product (e.g., corrosion thickness on a metal plate) or merely indicated through product performance (e.g., the luminosity of a light emitting diode). The measurements made across the product's lifetime are *degradation data*, and *degradation analysis* is the statistical tool for providing inference about the lifetime distribution from the degradation data.

Degradation testing and analysis is tied in with accelerated life testing (ALT) because both methods have evolved in recent years to suit reliability tests for which product lifetimes are expected to last far beyond the allotted test time (*see Accelerated Life Models*). ALT is meant to expedite product failure during test intervals by stressing the product beyond its normal use. This helps to bring in more information if the link between the accelerated test environment and the regular use environment is known. The same is true of degradation analysis. If the link between the measure of degradation and lifetime is clearly known, the degradation data provide valuable information about product reliability. Accelerated degradation testing (ADT) combines these two approaches by testing products in harsh environments and measuring the evidence of product degradation during the ALT.

Degradation analysis is especially useful for tests in which soft failures occur; that is, the lifetime of the test item is said to end after the measured performance decreases to a predetermined threshold value that designates a nonfunctioning state or an incipient failure. For a test of material strength, as an example, the degradation measurement might be the increasing size of the largest observed crack in the material, or the amount of corrosion measured on the surface. A failure event can be designated long before the material actually breaks. As another example, electronic components function reliably if their resistance, capacitance and voltage stay within

design limits, and failure can be said to occur when one or more of these parameters degrades beyond a specified limit. Ohring [1] presents a comprehensive array of physical models for electronic devices. Similarly, light emitting devices may be considered to fail only after the luminosity degrades below a fixed measured limit.

Figures 1 and 2 illustrate two different examples of degradation data. Figure 1, from Bogdanoff and Kozin [2] and featured in Meeker and Escobar [3], shows the measured crack size for alloy specimens that were fatigued by rapid cycling. Twelve of the 21 test items failed because their crack sizes exceeded the fixed threshold of 1.6 inches. Nine other test items did not fail, and if the degradation model is sound, more information will be gained from the degradation measurements in these nine observations compared to the respective lifetime measurements, which are right censored.

Figure 2 shows the measured degradation path for light emission of seven vacuum fluorescent displays (VFDs). The VFDs were tested for 200 h in an accelerated failure environment, so this is an example of an ADT. See Bae and Kvam [4] for details. In this case, the degradation path of each test item reaches the failure threshold (defined as the time the luminosity decreases by 50%) where a soft failure is defined.

Degradation Modeling

Models for degradation are generally either data-driven or derived from physical principles via stochastic processes. Although the data-driven model is more commonly applied to analyze degradation data, viewing degradation through stochastic processes helps researchers theoretically characterize the failure process.

Data-Driven Model

The measured degradation path for the i th tested device ($i = 1, \dots, n$) will consist of a vector of m_i measurements made at time points $t_{i1}, \dots, t_{im_i}$. The VFDs in Figure 2, for example, were each measured at the same five time points (0, 24, 72, 100, and 200 h). The measured degradation at t can be modeled as the actual unknown degradation $\eta(t)$

2 Degradation Models

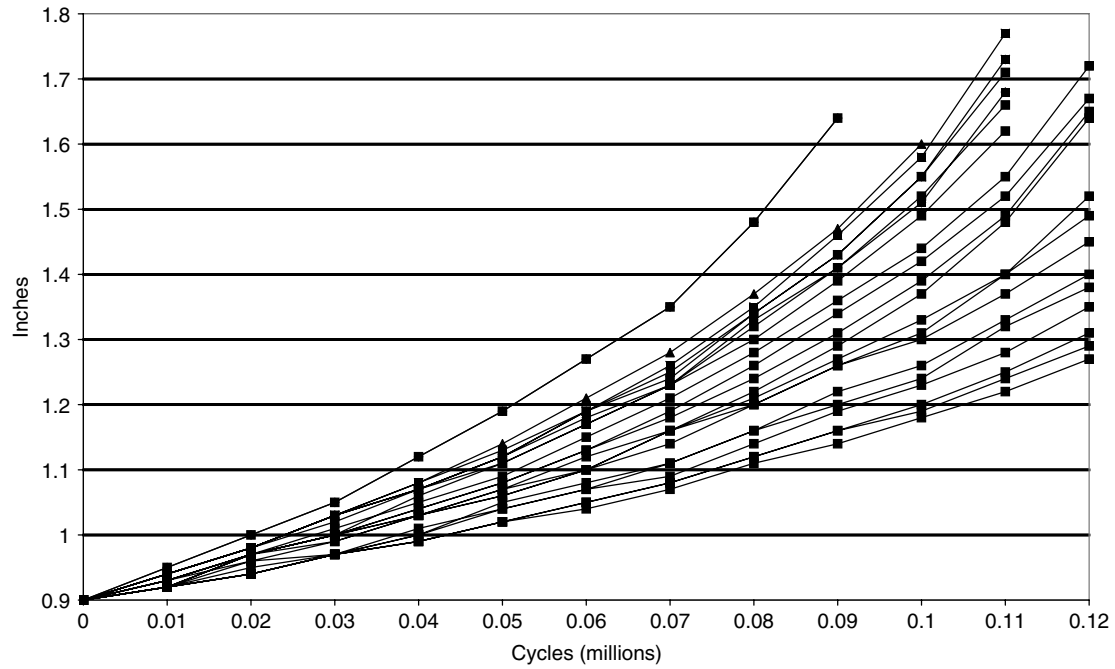


Figure 1 Alloy fatigue crack size (in inches), observations, measured in millions of cycles

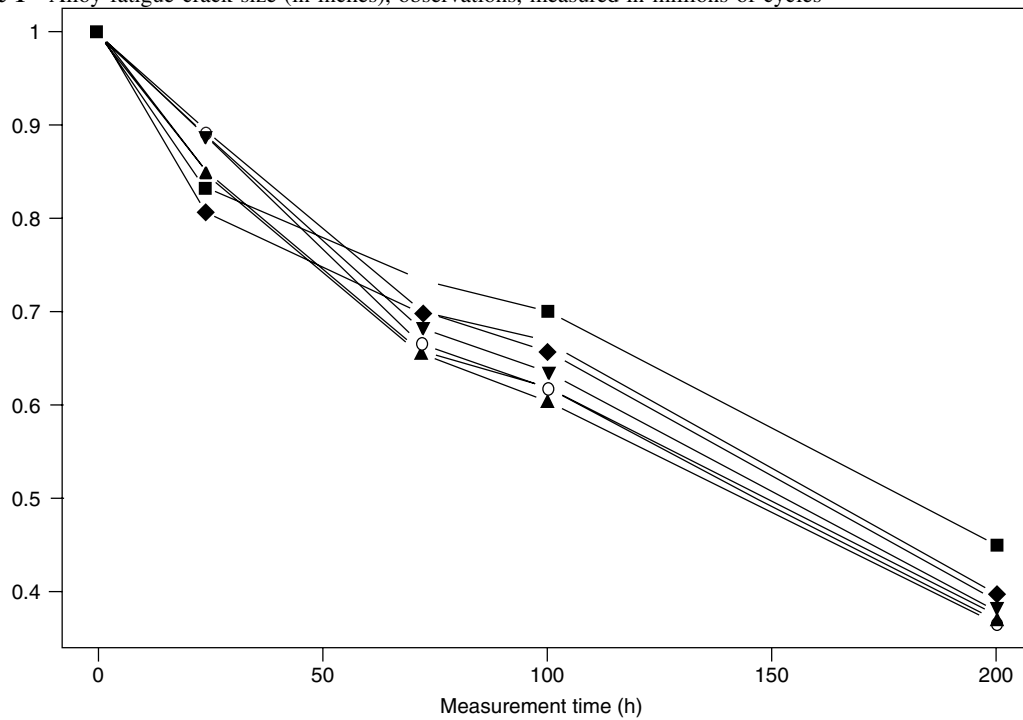


Figure 2 Luminosity degradation in seven VFD devices

plus a measurement error term ϵ . At the m_i time points, the degradation measurements of device i are

$$y_{ij} = \eta(t_{ij}) + \epsilon_{ij}, \quad 1 \leq i \leq n, \quad 1 \leq j \leq m_i \quad (1)$$

The form of η can be chosen to have a strict form or it can be more arbitrary. For example, the degradation of many electronic components is known to be of a log-linear form. That is, if we replace y with $\log(y)$, η can be a linear function of time such as $\eta(t) = \beta_1 + \beta_2 t$. On the other hand, the form for degradation can be unknown and nonparametric regression techniques are required to analyze the degradation data. Shiau and Lin [5] use this approach, but compared to a well chosen parametric model, their method can be quite inefficient. However, they show that the nonparametric technique can be useful for selecting a suitable parametric form of η .

If a specific form is described for η , we will have an unknown set of parameter values $\beta = (\beta_1, \dots, \beta_k)$ that must be estimated in order to fully characterize the degradation. In any realistic application of degradation analysis, the test units will degrade in a similar way but on distinct paths. To model this *unit-to-unit variability*, a distribution is assigned to β , allowing n distinctive paths to describe the degradation of the n test units. Meeker and Escobar [3] present a convincing argument for why the **normal distribution** adequately characterizes this randomness (i.e., $\beta \sim \mathcal{N}(\bar{\beta}, \Sigma)$). Along with β , let λ represent unknown parameters that are common across test units (thus no random effects are necessary) so the degradation path is expressed by

$$y_{ij} = \eta(t_{ij}; \lambda, \beta_i) + \epsilon_{ij}, \quad 1 \leq i \leq n, \quad 1 \leq j \leq m_i \quad (2)$$

Degradation Processes

In many experiments, degradation is a continuous progression of wear and decay, so it makes intuitive sense to model the degradation path with a stochastic process. If we let a stochastic process $W(t)$ to describe the degradation level of an item at time t , the mean degradation $\mu(t) = E[W(t)]$ is typically increasing and sometimes known through physical principles. For example, the wear on an automobile tire might be measured in terms of usage (t = odometer reading) and if the tire wear is constant, $\mu(t) = \lambda t$ for some unknown value λ .

The degradation characteristics of several electronic components can be described through stochastic processes. For example, Mitsuo [6] determined the mean degradation for light emitting diodes as $\mu(t) = \lambda t^k$ for some $\lambda > 0$ and $k \in \mathbb{R}$. Aven and Jensen [7] show how different processes imply different lifetime distributions. For example, in the case where the degradation follows a Wiener process, the time when the degradation level first reaches a fixed failure threshold has an inverse Gaussian distribution.

Stochastic processes are also helpful to infer lifetime distributions from **damage models**, which are a special case of degradation models, for example, see Kahle and Wendt [8]. For example, suppose $M(t)$ is a Poisson shock process with rate λ . Let P_k be the probability that the system survives k shocks, so that $1 = P_0 \leq P_1 \leq P_2 \leq \dots$. Then, the system survival can be expressed as

$$\sum_{k=0}^{\infty} P_k \frac{(\lambda t)^k}{k!} e^{-\lambda t}$$

If $P_{i+j} \leq P_i P_j$, then compared to a new device, the probability of surviving k additional shocks is smaller when the device has already absorbed some shocks. Based on this premise, it is possible to show that the lifetime distribution is *new better than used*, that is, for any $x, y > 0$, $P(X \geq x)P(X \geq y) \geq P(X \geq x + y)$ (see **Stochastic Orders and Aging**).

Relating Degradation to Lifetime

If the **physics of failure** is known through the model of an item's degradation over time, it's lifetime distribution can be inferred from this model as well (see **Degradation and Failure**). For materials, specimens exposed to constant stress cycles in a given stress range, lifetime is measured in number-of-cycles until failure (N). The Whöler curve (or $S - N$ curve) relates stress level (S) to N as $NS^b = k$, with known material parameters b and k . The $S - N$ equation is expressed in log form as $Y = \log N = \log k - b \log S$. If N is log-normally distributed, Y is normally distributed and regular regression models can be applied for predicting cycles-to-failure. The log-normal distribution is considered for modeling the failure time distribution when the corresponding degradation process based on rates that combine multiplicatively, and it is convenient for modeling fatigue

4 Degradation Models

crack growth in metals and composites. Sobczyk and Spencer [9] features numerous settings and examples.

An alternative model based on fatigue, introduced by Birnbaum and Saunders [10], defined B_n to be the measurable damage (*see Cumulative Damage Models Based on Gamma Processes*) to the test item after n cycles with accumulated amount of damage ζ_i in the i th cycle via $B_n = \zeta_1 + \dots + \zeta_n$, $i = 1, \dots, n$. If the ζ_i s are identically and independently distributed as with mean μ and variance σ^2 ,

$$P(N \leq n) = P(B_n > B^*) \approx \Phi\left(\frac{B^* - n\mu}{\sigma\sqrt{n}}\right) \quad (3)$$

where Φ is the standard normal cumulative distribution function (CDF). This results because B_n will be approximately normal if n is large enough. The reliability function for the test unit is

$$R(t) \approx \Phi\left(\frac{B^* - n\mu}{\sigma\sqrt{n}}\right) \quad (4)$$

This is called the *Birnbaum–Saunders* distribution, and it follows that

$$W = \frac{\mu\sqrt{N}}{\sigma} - \frac{B^*}{\sigma\sqrt{N}} \quad (5)$$

has a normal distribution, which leads to accessible implementation in lifetime modeling. Bogdanoff and Kozin [2] overview properties of the Birnbaum–Saunders distribution for materials testing.

Statistical Inference

If we assume the degradation path model in equation (2), with item-to-item variability reflected through random coefficients in η , we have an accumulated set of unknown parameters $\theta = (\lambda, \bar{\beta}, \Sigma)$, which makes for a difficult computation of the lifetime distribution. Numerical methods and simulations are typically employed to generate point estimates and confidence statements.

In selecting a degradation model based on longitudinal measurements of degradation, monotonic models are typically chosen under the assumption that degradation is a one-way process. From the degradation model, the lifetime distribution is defined as the

time at which the degradation first reaches the failure threshold, designated y^* :

$$F(t) = P(y(t) > y^*) = P(\eta(t; \lambda, \beta) + \epsilon > y^*) \quad (6)$$

Least squares (**Least-Squares Estimation**) or maximum likelihood (**Maximum Likelihood**) can be used to estimate the unknown parameters in the degradation model. To estimate $F(t_0)$, one can simulate N degradation curves from the estimated regression by generating N random coefficients $\theta_1, \dots, \theta_N$ from the estimated distribution $G(\theta; \hat{\beta})$. Next compute the estimated degradation curve for y_i based on the model with θ_i and $\hat{\lambda}$: $y_i(t) = \eta_i(t; \hat{\lambda}, \theta_i)$. Then $\hat{F}(t_0)$ is the proportion of the N generated curves that have reached the failure threshold y^* by time t_0 .

A nonparametric **bootstrap** sampling procedure can be used for measuring the uncertainty in the lifetime distribution estimate. The bootstrap procedure *resamples* the sample degradation curves *with replacement* (i.e., so some curves may not be represented in the sample while others may be represented multiple times). Meeker and Escobar [3] summarize a bootstrap procedure for making **confidence intervals** for the lifetime distribution:

1. Compute estimates of parameters $\bar{\beta}$, λ , Σ .
2. Use simulation (above) to construct $\hat{F}(t_0)$.
3. Generate $N \geq 1000$ bootstrap samples, and for each one, compute estimates $\hat{F}^{(1)}(t_0), \dots, \hat{F}^{(N)}(t_0)$. This is done as before except now the M simulated degradation paths are constructed with an error term generated from $H(\eta; \hat{\Sigma})$ to reflect variability in any single degradation path.
4. With the collection of bootstrap estimates in equation (7), compute a $1 - \alpha$ level confidence interval for $F(t_0)$ as $(\hat{F}^l(t_0), \hat{F}^u(t_0))$, where the indexes $1 \leq l \leq u \leq N$ are calculated as $N^{-1} = \Phi(2\Phi^{-0.5}(p_0) + \Phi^{-0.5}(0.5\alpha))$ and $uN^{-1} = \Phi(2\Phi^{-0.5}(p_0) + \Phi^{-0.5}(0.5(1 - \alpha)))$, and p_0 is the proportion of bootstrap estimates of $F(t_0)$ less than $\hat{F}(t_0)$.

Example

The VFDs are tested (see Figure 2) at higher filament voltage level than normal usage condition

(Ef = 4.55 V). It is known that the display luminosity for VFDs decreases exponentially over most of the usage period when the degradation path can be expressed as

$$\Lambda(t) = \beta_0 \exp(-\beta_1 t) \quad (7)$$

where $\Lambda(t)$ denotes the luminosity at time t , β_0 is the initial luminosity, and β_1 is the rate of degradation. The light device is considered to fail at the time its luminosity decreases below 50% of its initial measurement. That is, failure is defined as the first time relative luminosity

$$y(t) = \frac{\Lambda(t)}{\Lambda(0)} = \exp(-\beta_1 t) \quad (8)$$

falls below 0.5. Using least-squares regression with the log transformation $\log y(t) = -\beta_1 t$, and by ignoring variation between individuals, the fitted (no-intercept) model is

$$\log y(t) = -0.0048t \quad (9)$$

with $R^2 = 0.9893$. However, the model fails to reflect individual variation of degradation rate (see Table 1).

To reflect individual variation of degradation, a linear random-coefficients model can be applied by assuming that β_1 is random. The random effects model can be written as

$$\log y(t_{ij}) = -(\beta + u_i)t_{ij} + \epsilon_{ij} \quad (10)$$

where t_{ij} is the covariate for the j th measurement time on the i th individual, $\epsilon_{ij} \sim \mathcal{N}(0, \sigma^2)$, and $u_i \sim \mathcal{N}(0, \sigma_u^2)$. Using the SAS NLMIXED procedure, the fitted linear random-coefficients model is computed numerically as

$$\log y(t_{ij}) = -0.004751t_{ij} \quad (11)$$

Table 1 Estimated degradation rates for VFDs in Figure 2

Sample	Degradation rate
1	4.6×10^{-3}
2	5.1×10^{-3}
3	4.6×10^{-3}
4	5.1×10^{-3}
5	4.8×10^{-3}
6	4.0×10^{-3}
7	5.0×10^{-3}

with $\hat{\sigma}_u^2 = 0.001574$ and $\hat{\sigma}^2 = 1.032 \times 10^{-7}$. Based on the estimated parameters, the fitted lines are given in Figure 3.

Comments

Degradation measurements have great potential to improve lifetime data analysis, but they also introduce new problems to the statistical inference. Lifetime models have been researched and refined for many manufactured products that are put on test. On the other hand, parametric degradation models tend to be based on simple physical properties of the test item and its environment (e.g., the Paris crack law, Arrhenius rule, Power law) which often lead to obscure lifetime models. Meeker and Escobar [3] show that most valid degradation models will not yield lifetime distributions with closed-form solutions. Given the improving computational tools available to researchers, this should be no deterrent to using degradation analysis.

In a setting where the lifetime distribution is known, but the degradation distribution is unknown, degradation information does not necessarily complement the available lifetime data. For example, the lifetime data may be distributed as Weibull, but conventional degradation models will contradict the Weibull assumption (actually, the rarely used *reciprocal Weibull* distribution for degradation with a fixed failure threshold leads to Weibull lifetimes). Bae *et al.* [11] discuss this problem in terms of a simple additive degradation model.

Finally, the monotone relationship between degradation and usage (or time) does not necessarily hold for all applications. Bae and Kvam [12] analyze light display data for a VFD that shows nonmonotonic degradation during its **burn-in** period (*see Burn-In and Maintenance Policies*), when impurities in the vacuum are being burned off. As this happens, luminosity actually increases slightly before beginning a long and steady decrease over its usage period. In this case, a more complicated mixture model that captures the burn-in effect proves to be more efficient.

Compared to ordinary life testing or even ALT, degradation analysis procedures tend to be computationally cumbersome. However, for important applications, the increase in statistical efficiency can be dramatic. In the past, these computations have

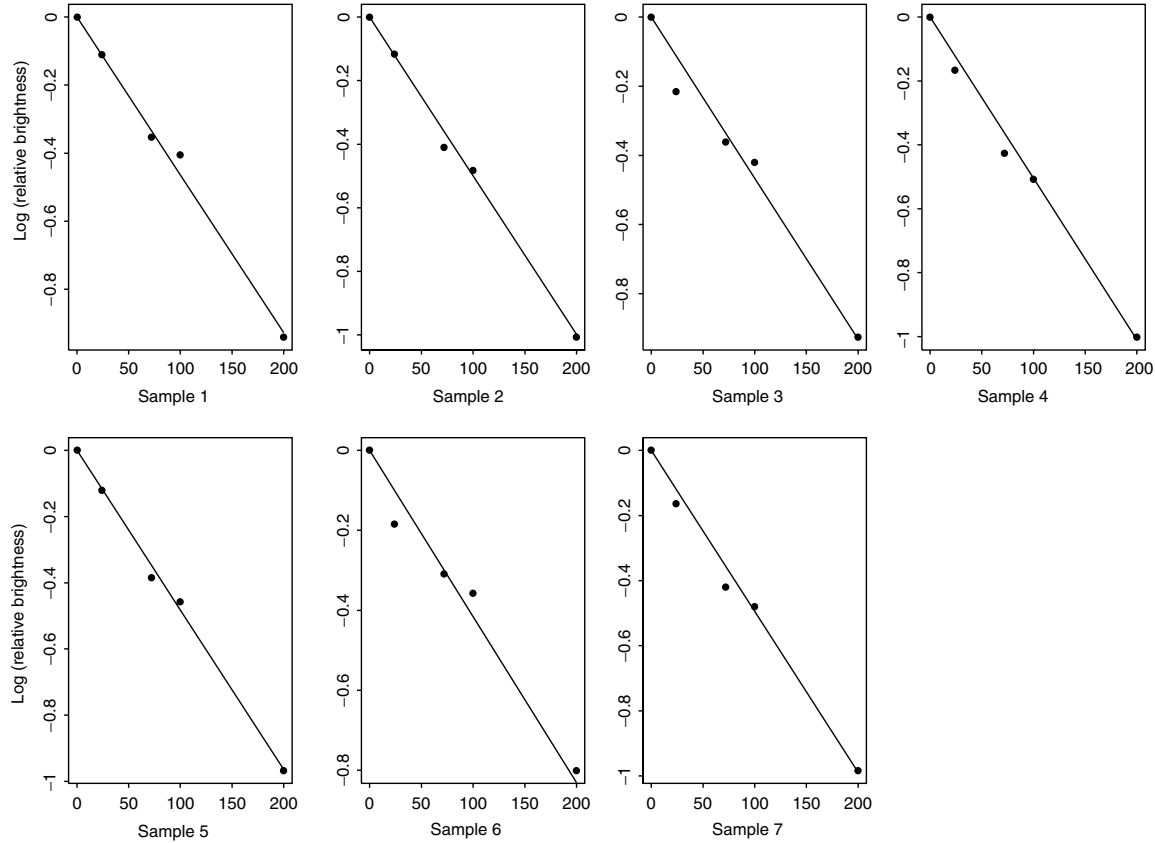


Figure 3 Fitted degradation model

impeded degradation analysis from being a feature of reliability problem solving. Such analyses are easier to implement now, and the reliability analyst need not be coerced into using an overly simplistic model – for instance, a linear model that does not allow for random coefficients. Chen and Zheng [13] construct an **imputation** algorithm to generate a closed-form solution, although its performance is suspect for medium or small samples.

Robinson and Crowder [14] introduce a Bayesian approach and showed if **prior distributions** are chosen from a reasonable class, they have negligible effect on the lifetime **estimation**. Bayesian reliability is summarized in **Bayesian Reliability Analysis**. In most cases the Bayesian procedure is even more computational than the one introduced here, but **Markov chain Monte Carlo** (e.g., **Hierarchical Markov Chain Monte Carlo (MCMC)** for

Bayesian System Reliability) methods have straightforward implementation.

Acknowledgments

This research was supported by NSF grant DMI-0400071.

References

- [1] Ohring, M. (1998). *Reliability and Failure of Electronic Materials and Devices*, Academic Press.
- [2] Bogdanoff, J.L. & Kozin, F. (1985). *Probabilistic Models of Cumulative Damage*, John Wiley & Sons.
- [3] Meeker, W.Q. & Escobar, L.A. (1998). *Statistical Methods for Reliability Data*, John Wiley & Sons.
- [4] Bae, S.J. & Kvam, P.H. (2004). A nonlinear random coefficients model for degradation testing, *Technometrics* **46**, 460–469.
- [5] Shiau, J.-J.H. & Lin, H.H. (1999). Analyzing accelerated degradation data by nonparametric regression, *IEEE Transactions on Reliability* **48**, 149–158.

-
- [6] Mitsuo, F. (1991). *Reliability and Degradation of Semiconductor Lasers and LED*, Artech House, Norwood.
 - [7] Aven, T. & Jensen, U. (1999). *Stochastic Models in Reliability*, Springer-Verlag, New York.
 - [8] Kahle, W. & Wendt, H. (2000). In *Statistical Analysis of Damage Processes, Recent Advances in Reliability Theory: Methodology, Practice and Inference*, N. Limnios & M. Nikulin, eds, Birkhäuser, Boston.
 - [9] Sobczyk, K. & Spencer, B.F. (1992). *Random Fatigue From Data to Theory*, Academic Press.
 - [10] Birnbaum, Z.W. & Saunders, S.C. (1969). A new family of life distributions, *Journal of Applied Probability* **6**, 319–327.
 - [11] Bae, S.J., Kuo, W. & Kvam, P.H. (2007). Degradation models and implied lifetime distribution, *Reliability Engineering and System Safety* **92**, 601–608.
 - [12] Bae, S.J. & Kvam, P.H. (2006). A change-point analysis for modeling incomplete burn-in for light displays, *IEEE Transactions* **38**, 489–498.
 - [13] Chen, Z. & Zheng, S. (2005). Lifetime distribution based degradation analysis, *IEEE Transactions on Reliability* **54**, 3–10.
 - [14] Robinson, M.E. & Crowder, M.J. (2000). Bayesian methods for a growth-curve degradation model with repeated measures, *Lifetime Data Analysis* **6**, 357–374.

Related Articles

Accelerated Life Models; Bayesian Reliability Analysis; Burn-In and Maintenance Policies; Cumulative Damage Models Based on Gamma Processes; Design for Reliability; Degradation Processes; Degradation and Failure; Mean Square Error; Stochastic Orders and Aging; Stochastic Deterioration; Warranty: Usage and Wear Process for.

SUK J. BAE AND PAUL H. KVAM

Reliability Growth Modeling

Introduction

Modeling reliability growth (RG) has received considerable attention in the statistical and engineering literature over the past four decades. At the initial developmental phase of any production involving complex systems, a test-analyze-and-fix (TAAF) process is often undertaken to promote the detection of system faults and incorporate appropriate redesigns or other corrective actions, and the modified system is retested. As this test–redesign–retest sequence contributes to an improvement in the system performance, failure data become increasingly sparse at the later stages of testing making it more difficult to assess the current reliability. An RG model provides a structure through which the failure data from the current as well as previous stages of testing could be analyzed in an integrated way to make efficient inference on system reliability, and other aspects of the underlying failure process.

RG models have been applied extensively in analyzing failure data arising from hardware as well as a piece of software. Further, for systems with single-shot missions that undergo a developmental program, RG models arise quite naturally while studying the effectiveness of design fixes. Below we describe some popular RG models that have been developed in the context of observing either time-to-failure or a sequence of dichotomous/trichotomous outcomes.

Continuous-Time RG Models

The continuous-time RG models are typically classified into two general categories; one class models the *time between successive failures*, whereas the other category builds on a modeling framework directly applied to the cumulative count of failures viewed over the continuum of time. Following Singpurwalla and Wilson [1] we shall term these two classes *Type-I* and *Type-II models*, respectively, noting that a model of each type generates a unique model of the other type. In the discussion that follows, we shall describe some

of the more commonly used models under each classification.

A major thrust to RG modeling in the area of hardware reliability rose from certain empirical findings in Duane [2] from examination of the failure data of a variety of systems such as complex hydromechanical devices, aircraft generators, and jet engines in the course of their development. When plotted on a log–log scale, the cumulative number of failures was typically found to produce a linear pattern of relationship with the cumulative operating time. This phenomenon, later came to be known as the *Duane postulate*, was given a concrete stochastic basis in Crow [3], where failures during the development stage of a new system were assumed to follow a nonhomogeneous **Poisson process** (NHPP) with an **intensity function** $\lambda(t)$ of the Weibull form

$$\lambda(t) = \mu \delta t^{\delta-1} \quad (1)$$

Note that equation (1) corresponds to a Type-II model according to the classification scheme described above. The case $\delta < 1$, corresponds to a decrease in the failure intensity and hence, a growth in reliability. The corresponding cumulative failure rate μt^δ is linear with the accumulated time on a log–log scale. Incidentally, an NHPP formulation was originally proposed by Ascher [4] in modeling the reliability change of a *bad-as-old system*, and was later used by Bassin [5] with the Weibull intensity to obtain optimal overhaul intervals for various machines. A large body of literature has evolved thereafter, focusing on theoretical developments as well as applications of equation (1) in RG analysis. Under this model, a thorough treatment of maximum-likelihood **estimation** and confidence procedures for the parameters was provided by Crow [3] under both time-truncated and failure-truncated testing schemes applied to one or more systems. Testing procedures for equality of the scale parameters for two or more independent systems were studied by Lee [6] under the time-truncated scheme. Exact **confidence intervals** for the scale parameter under the failure-truncated scheme were obtained by Finkelstein [7] using **Monte Carlo** methods. Sequential testing procedures for the growth parameter δ was discussed by Bain and Engelhardt [8]. Inferences on the current system reliability were studied by Lee and Lee [9], Bain and Engelhardt [10], Crow [11], Rigdon and Basu [12], and Sen and Khattree [13].

2 Reliability Growth Modeling

In the continuous-time case, the NHPP with Weibull intensity has gained vast popularity due to its empirical fit to a variety of data *vis-a-vis* its conformity to the Duane learning curve, and availability of elegant distributional results concerning statistical inferences. While empirical fit, nice mathematical properties, and tractability of statistical inferences are important aspects of a stochastic model, a clear physical interpretation is also of paramount importance. Finkelstein [14] pointed out that the Weibull intensity $\lambda(t)$ has the rather unrealistic feature that it tends to ∞ as $t \rightarrow 0$ and to 0 as $t \rightarrow \infty$. Finkelstein and later Littlewood [15] discussed some modifications of the NHPP model to overcome these difficulties. Aside from these limiting values, there is an even more serious criticism concerning the principle of this modeling. Duane [2] indicated that a learning curve is used to monitor developmental progress and plan for reliability improvement. The major point of concern with the NHPP model is its continually changing intensity function irrespective of the failure history, which is in direct conflict with the concept of the Duane learning curve, as well as the conceptual framework of a test–redesign–retest course of development program. If system improvement is assumed to be effected only after a failure is observed, a realistic model should be flexible enough to incorporate a change in the failure rate at the occurrences of the failures. Thompson [16] expressed the same concern by saying that “... some provision needs to be present for altering the process of failures when modifications or corrective actions are applied to the system”.

Incorporating a step change in the failure pattern at the occurrences of successive failures, Sen and Bhattacharyya [17] proposed a Type-I RG model, in which the successive interfailure times are distributed as independent exponential random variables with hazard λ_i at the i th stage of the process. The authors investigated the specific RG parameterization

$$\lambda_i = (\mu/\delta)i^{1-\delta}, \quad \mu > 0, \quad \delta \geq 1 \quad (2)$$

that is intimately connected to the Duane learning curve property described above. The authors showed that under the sampling scheme of observing a pre-specified number of failures, the maximum-likelihood estimators (MLEs) of the parameters are jointly asymptotically *singular* normal with quite nonstandard rates of convergence. More recently, Sen [18] investigated exact as well as large-sample inference

of the estimate of the current intensity, as evaluated at the end of the testing stage.

In the area of software engineering, reliability is viewed synonymously with the probability of failure-free operation of a software code in a prespecified duration. In contrast to the field of hardware reliability, a distinctively large number of RG models have been proposed that have applications in error detection in software programs. One of the earliest, and perhaps the most well-studied software RG model is a Type-I model proposed by Jelinski and Moranda [19], in which the time between the $(i-1)$ st and the i th failure of the software is distributed as an exponential random variable with hazard

$$\lambda_i = \Lambda(N - i + 1), \quad \Lambda > 0, \quad N \text{ integer} \quad (3)$$

independently of other interfailure times. In this parameterization, N denotes an unknown fixed number of bugs present in the code, and Λ is a constant of proportionality. While the importance of the model in equation (3) is immense in terms of initialing the formalization of software RG models, some aspects of the model have drawn criticisms from researchers in this area. Specifically, the assumptions underlying the formulation of equation (3) of (a) equal contribution from each bug to the failure rate and (b) a perfect repair mechanism, are often thought to be overly simplistic, and not true descriptions of the failure process. In view of countering criticism (a), Moranda [20] modified the hazard rate to be of the geometric decay form

$$\lambda_i = D\rho^{i-1}, \quad D > 0, \quad 0 < \rho < 1 \quad (4)$$

Goel and Okumoto [21] imposed the possibility of **imperfect repair** directly within the modeling framework of equation (3) by rewriting the hazard rate as

$$\lambda_i = \Lambda(N - p(i-1)), \quad \Lambda > 0, \quad N \text{ integer}, \quad 0 < p < 1 \quad (5)$$

where p is the probability that a bug is detected and repaired. Schick and Wolverton [22] generalized equation (3) in a direction distinct from the rest by assuming that the conditional failure rate of the i th event is proportional to the number of remaining bugs as well as the accumulated time since last failure. This is tantamount to assuming yet another Type-I model

with an increasing linear rate of failure (starting at zero at each failure), governing the successive interfailure times. The underlying rationale for such a model reflects the belief that a software is highly reliable after a repair, and the reliability decays with time until the next failure (and bug removal) occurs.

A variant of the usual Type-I RG model that directly models the interdependence between the times between failures has been considered by Singpurwalla and Soyer [23]. Specifically, they assume the interfailure times T_i , $i = 1, 2, \dots$ to satisfy the autoregressive relationship

$$\log T_i = \theta_i \log T_{i-1} + \epsilon_i \quad (6)$$

In equation (6) if T_i 's are scaled to have values larger than 1, values of θ_i greater than 1 would imply RG. Under the assumption of normality of $\log T_i$'s and ϵ_i 's, Singpurwalla and Soyer [23] treat θ_i as random effects and offer Bayesian solution to the inference problem associated with equation (6).

Several Type-II models in the context of modeling growth in software reliability have been proposed in the literature. One of the earliest of them is the model proposed by Goel and Okumoto [24], which assumes the counting process of failures to follow an NHPP with intensity function

$$\lambda(t) = ab \exp(-bt), \quad a > 0, b > 0 \quad (7)$$

In equation (7), a symbolizes the average number of failures to be found "eventually", and b is the growth rate in reliability. Musa and Okumoto [25] postulate a relationship

$$\lambda(t) = \lambda \exp(-\theta M(t)), \quad \lambda > 0, \theta > 0 \quad (8)$$

between the intensity function $\lambda(t)$ and the mean number of failures $M(t)$ up to time t of an NHPP. Both equations (7) and (8) entertain the possibility of having an infinite number of bugs in the software with the interfailure times being stochastically dependent, two features that distinguish them from the usual Type-I models. Additionally, the parameterization in equation (7) imposes a positive probability at infinity for the random variable denoting time to first failure. This may be a realistic assumption for a software code that is a newer release (version) of an older software which has undergone several iterations of development.

A number of Bayesian formulations of both types of models described above have been proposed

in the literature. The Bayesian estimation problem for equation (1) has specifically been addressed by Higgins and Tsokos [26], Guida *et al.* [27], Barlev *et al.* [28], and Sen [29]. A general Bayesian framework for the class of NHPP models has been put forth by Kuo and Lang [30].

In the context of Bayesian inference within the class of Type-I models, substantial effort has been invested in studying *piecewise exponential* models, i.e. models with independent and exponentially distributed interfailure times, where the means λ_i are explicit functions of cumulative failure numbers. Various parametric forms for λ_i have been studied in a Bayesian context by several authors (e.g., Littlewood and Verrall [31], Mazzuchi and Soyer [32], and Soyer [33]). In a more recent work, Erkanli *et al.* [34] avoided any parameterization of λ_i 's, and modeled them directly under an *a priori* monotonicity constraint consistent with RG. Basu and Ebrahimi [35] relaxed the monotonicity assumption and modeled the sequence of λ_i 's as a martingale process. Both of these later works relied heavily on the modern computational advances of **Markov chain Monte Carlo** procedures.

Discrete-Time RG Models

For single-mission systems, the test results are binary in nature as opposed to time to failure in a continuous-time framework. Growth in reliability in the discrete case is demonstrated by an increment in the probability of success at successive stages of the development program. The genesis of the discrete RG models has its roots in the early attempt to quantify the increase in system reliability attained after implementing posttest design changes. Numerous proposals for rescoring "observed failures" as "adjusted partial successes" have been put forth in the literature. Simulations showed, however, that most of the proposed methods were very sensitive to the choice of the failure-discounting parameter(s), which was somewhat arbitrary.

A more formal approach is undertaken by considering that the developmental phase comprises a sequence of independent Bernoulli trials that are not identically distributed across stages. Let p_i denote the probability of success at the i th stage of a developmental program. On the basis of the data of x_i successes in n_i trials at stage i , $i = 1, \dots, n$, Barlow *et al.* [36] obtained MLEs for p_i 's under the

order restriction $p_1 \leq \dots \leq p_n$. An extension of the dichotomous characterization of the trial outcome to the case of partitioning of failures in two distinct categories, namely, assignable-cause and inherent, has been an area of major research interest to the reliability practitioners for a long time. The former corresponds to failure modes that can be eliminated by redesigns or by equipment or operational modifications, while the latter reflects the state of the art and remains unchanged as the system configuration advances. Wolman [37] initialized work in this area, although his study was confined to the calculation of probabilistic quantities, such as “chance of eliminating all assignable-cause failure modes by trial n ”. Barlow and Scheuer [38] followed up with a formal trinomial formulation of the process. They assumed that the probability of an inherent failure q_0 remains a constant throughout the developmental program, while the RG feature manifests itself into a decrease of q_i , the probability of the assignable-cause failure, over successive stages. They obtained MLEs for q_0 and for the q_i ’s subject to the monotonicity constraint $q_i < q_{i-1}$.

On the basis of an extensive simulation study, Olsen [39] investigated the pros and cons of the Barlow–Scheuer restricted MLEs. It is concluded in [39] that these estimators should only be used when (a) a reasonably large sample size is available after the last set of corrections, (b) testing is stopped with the occurrence of assignable-cause failures, (c) corrective actions do not fully remove individual failure modes or introduce new failure modes. It was noted that, in general, the Barlow–Scheuer methodology yielded positive biased estimator $\hat{R}_n = 1 - \hat{q}_0 - \hat{q}_n$ for reliability R_n at the n th (final) stage of testing.

Weinrich and Gross [40] and Mazzuchi and Soyer [41] investigated Bayesian formulations of the Barlow–Scheuer piecewise trinomial model. The modification of the Barlow–Scheuer model in [40] entails redistribution of the probability among other causes of inherent and assignable-cause failure, upon removal of an assignable-cause failure mode. Assuming the initial probabilities (success, assignable cause, inherent cause) to possess a Dirichlet **pdf**, the posterior *pdf* of stage- k probabilities is updated as either a Dirichlet or a mixture of Dirichlet depending on the extent of knowledge of the amount of initial assignable-cause failures removed. Instead of specifying a prior for the probabilities only at the initial

stage, Mazzuchi and Soyer [41] imposed a Dirichlet prior on the differences of the assignable-cause probabilities from successive stages. This clever manipulation clearly defines a *pdf* for the assignable-cause probabilities decreasing across stages. The major contributions of Mazzuchi and Soyer [41] formulations were (a) specification of the prior belief in the entire RG pattern, (b) retention of all relevant marginal and posterior *pdf*’s as tractable Dirichlet (or β) forms, and (c) flexible framework for elicitation of prior information from experts.

Under the construct of observing single-shot trials with dichotomous outcomes, a large class of RG models that adheres to a specific parametric form for the probability of success (reliability) p_k at the k th stage of the developmental process has been investigated over the years. Some of the earlier models of this type were extensively studied by Lloyd and Lipow [42] and Gross and Kamins [43, 44]. The “Gompertz” model proposed in Virene [45] is an interesting inclusion in this class in the sense that it was the only model of this form which attempted to characterize S-shaped patterns for RG.

Recent research in discrete reliability models has devoted considerable attention to another class of parametric models intimately linked to the Duane learning curve property described in the previous section. Identifying a single trial to a unit time in the continuous case, Crow [46] developed the following model used extensively by the US Army:

$$p_k = 1 - \lambda \left[(N_k)^\beta - (N_{k-1})^\beta \right] / n_k, \\ 0 < \lambda < 1, \quad 0 < \beta < 1 \quad (9)$$

where k denotes the configuration number, N_k denotes the cumulative number of test trials until the k th configuration, $N_0 \equiv 0$, and $n_k = N_k - N_{k-1}$. The derivation rested on the underlying assumption that for each stage of testing, the number of observations is fixed and the distribution of outcomes in successes and failures is random. Such a situation can arise, for example, when the system producer delivers batches for sequential phases of testing, all of the scheduled test events are executed, and the results for each stage lead to modifications of the system design. Some estimation procedures for this model were suggested by Crow [46] and Finkelstein [47]; asymptotic properties were studied by Bhattacharyya *et al.* [48], Bhattacharyya and Ghosh [49], and Johnson [50].

The development of the models described above is incompatible with another common strategy of stopping testing as soon as any failure occurs and introducing system redesigns or corrective actions as soon as possible, i.e. inverse sampling. This, in particular, is a viable model for destructive testing of very expensive systems, e.g. testing of the launch and landing capabilities of unmanned aerial vehicles. One of the earlier models under this framework was proposed by Dubman and Sherman [51], who considered the formulation

$$p_k = 1 - q\rho^k \quad (10)$$

which amounts to a “geometric decay” in the unreliability. The focus of the detailed study by Dubman and Sherman [51] was on the asymptotic properties of the MLEs. Fries [52] took the learning curve property and directly imposed the inverse sampling strategy to derive the following model expressed as

$$p_k = 1 - \lambda' \left[k^{\beta'} - (k-1)^{\beta'} \right]^{-1}, 0 < \lambda' < 1, \beta' > 1 \quad (11)$$

The inference procedures for equation (11) as well as the effects of using results from the direct-sampling model (9) to data which is generated from equation (11) have been documented in Sen and Fries [53, 54].

Concluding Remarks

RG modeling is often deemed to fall under the larger area of studying **repairable systems**. While some of the models described above can be easily extended to incorporate *decay* in reliability, thereby making it a candidate for modeling failures of a repairable system, others are quite specifically designed to depict a growth in reliability. It is important to realize that the major focus in RG modeling is an assessment of system improvement effected by the design changes. Further, the models described above are mostly appropriate for analyzing a modest chain of events from a single system. The context is relevant for studying very complex, expensive systems, for which only one or a very small number of prototypes may be available. When only a few failures are observed for a substantially large number of systems/units, such as in many medical and manufacturing applications, a nonparametric or semiparametric approach is deemed more reasonable. Substantial

methodological advances have been made in this area, which clearly is focused on issues far removed from the topic of RG modeling, and hence, is not covered in this article.

References

- [1] Singpurwalla, N.D. & Wilson, S.P. (1999). *Statistical Methods in Software Engineering: Reliability and Risk*, Springer-Verlag, New York.
- [2] Duane, J.T. (1964). Learning curve approach to reliability monitoring, *IEEE Transactions on Aerospace* **2**, 563–566.
- [3] Crow, L.H. (1974). Reliability analysis for complex, repairable systems, in *Reliability and Biometry*, F. Proschan & R.J. Serfling, eds, Society for Industrial and Applied Mathematics, Philadelphia, pp. 379–410.
- [4] Ascher, H.E. (1968). Evaluation of repairable system reliability using the bad-as-old concept, *IEEE Transactions on Reliability* **17**, 103–110.
- [5] Bassin, W.M. (1973). A Bayesian optimal overhaul interval model for the Weibull restoration process case, *Journal of the American Statistical Association* **68**, 575–578.
- [6] Lee, L. (1980). Comparing rates of several independent Weibull processes, *Technometrics* **22**, 427–430.
- [7] Finkelstein, J.M. (1976). Confidence bounds on the parameters of the Weibull process, *Technometrics* **18**, 115–117.
- [8] Bain, L.J. & Engelhardt, M. (1982). Sequential probability ratio tests for the shape parameter of a NHPP, *IEEE Transactions on Reliability* **31**, 79–83.
- [9] Lee, L. & Lee, S.K. (1978). Some results on inference for the Weibull process, *Technometrics* **20**, 41–45.
- [10] Bain, L.J. & Engelhardt, M. (1980). Inferences on the parameters and current system reliability for a time truncated Weibull process, *Technometrics* **22**, 421–426.
- [11] Crow, L.H. (1982). Confidence interval procedures for the Weibull process with applications to reliability growth, *Technometrics* **24**, 67–72.
- [12] Rigdon, S.E. & Basu, A.P. (1988). Estimating the intensity function of a Weibull process at the current time: failure truncated case, *Journal of Statistical Computation and Simulation* **30**, 17–38.
- [13] Sen, A. & Khattree, R. (1998). On estimating current intensity of a power law process, *Journal of Statistical Planning and Inference* **74**, 253–272.
- [14] Finkelstein, J.M. (1979). Starting and limiting values for reliability growth, *IEEE Transactions on Reliability* **28**, 111–113.
- [15] Littlewood, B. (1984). Rationale for a modified Duane model, *IEEE Transactions on Reliability* **33**, 157–159.
- [16] Thompson Jr, W. (1988). *Point Process Models with Applications to Safety and Reliability*, Chapman & Hall, New York.

- [17] Sen, A. & Bhattacharyya, G.K. (1993). A piecewise exponential model for reliability growth and associated inferences, *Advances in Reliability*, A.P. Basu, ed, Elsevier/North Holland [Elsevier Science Publishing, New York; North Holland Publishing, Amsterdam], New York, Amsterdam, pp. 331–355.
- [18] Sen, A. (1998). Estimation of current reliability in a Duane-based reliability growth model, *Technometrics* **40**, 334–344.
- [19] Jelinski, Z. & Moranda, P.B. (1972). Software reliability research, *Statistical Computer Performance Evaluation*, W. Freiburger, ed, Academic Press, New York, pp. 465–497.
- [20] Moranda, P.B. (1975). Prediction of software reliability and its applications, in *Proceedings of the Annual Reliability and Maintainability Symposium*, Washington, DC, pp. 327–332.
- [21] Goel, A.L. & Okumoto, K. (1978). An analysis of recurrent software failures on a real-time control system, in *Proceedings of the ACM Annual Technical Conference*, Washington, DC, pp. 496–500.
- [22] Schick, G.J. & Wolverton, R.W. (1978). Assessment of software reliability, *Proceedings in Operations Research*, Physica-Verlag, Vienna, pp. 395–422.
- [23] Singpurwalla, N.D. & Soyer, R. (1985). Assessing (software) reliability growth using a random coefficient autoregressive process and its ramifications, *IEEE Transactions on Software Engineering* **11**, 1456–1464.
- [24] Goel, A.L. & Okumoto, K. (1979). Time-dependent error-detection rate model for software reliability and other performance measures, *IEEE Transactions on Reliability* **28**, 206–211.
- [25] Musa, J.D. & Okumoto, K. (1984). A logarithmic Poisson execution time model for software reliability measurement, in *Proceedings of the Seventh International Conference on Software Engineering*, Orlando, pp. 230–238.
- [26] Higgins, J.J. & Tsokos, C.P. (1981). A quasi-Bayes estimate of the failure intensity of a reliability growth model, *IEEE Transactions on Reliability* **30**, 471–475.
- [27] Guida, M., Calabria, R. & Pulcini, G. (1989). Bayes inference for a non-homogeneous Poisson process with power intensity law, *IEEE Transactions on Reliability* **38**, 603–609.
- [28] Bar-lev, S.K., Lavi, I. & Reiser, B. (1992). Bayesian inference for the power law process, *Annals of Institute of Statistical Mathematics* **44**, 623–639.
- [29] Sen, A. (2002). Bayesian estimation and prediction of the intensity of the power law process, *Journal of Statistical Computation and Simulation* **72**, 613–631.
- [30] Kuo, L. & Yang, T.Y. (1996). Bayesian computation for nonhomogeneous Poisson processes in software reliability, *Journal of the American Statistical Association* **91**, 763–773.
- [31] Littlewood, B. & Verrall, J.L. (1973). A Bayesian reliability growth model for computer software, *The Journal of the Royal Statistical Society. Series C* **22**, 332–346.
- [32] Mazzuchi, T.A. & Soyer, R. (1988). A Bayes empirical-Bayes model for software reliability, *IEEE Transactions on Reliability* **37**, 248–254.
- [33] Soyer, R. (1992). *Monitoring software reliability using non-Gaussian dynamic models*, *Proceedings of the Engineering Systems Design and Analysis Conference*, Istanbul **1** pp. 419–423.
- [34] Erkanli, A., Mazzuchi, T.A. & Soyer, R. (1998). Bayesian computations for a class of reliability growth models, *Technometrics* **40**, 14–23.
- [35] Basu, S. & Ebrahimi, N. (2003). Bayesian software reliability models based on martingale processes, *Technometrics* **45**, 150–158.
- [36] Barlow, R.E., Proschan, F. & Scheuer, E.M. (1966). *Maximum Likelihood Estimation and Conservative Confidence Interval Procedures in Reliability Growth and Debugging Problems*, Memorandum RM-4749-NASA, RAND Corporation, Santa Monica.
- [37] Wolman, W. (1963). Problems in system reliability analysis, in *Statistical Theory of Reliability*, M. Zelen, ed, University of Wisconsin Press, Madison, pp. 149–160.
- [38] Barlow, R.E. & Scheuer, E.M. (1966). Reliability growth during a development testing program, *Technometrics* **8**, 53–60.
- [39] Olsen, D.E. (1977). Estimating reliability growth, *IEEE Transactions on Reliability* **26**, 50–53.
- [40] Weinrich, M.C. & Gross, A.J. (1978). The Barlow-Scheuer reliability growth model from a Bayesian viewpoint, *Technometrics* **20**, 249–254.
- [41] Mazzuchi, T.A. & Soyer, R. (1993). A Bayes method for assessing product-reliability during development testing, *IEEE Transactions on Reliability* **42**, 503–510.
- [42] Lloyd, D.K. & Lipow, M. (1964). *Reliability: Management, Methods, and Mathematics*, Prentice Hall, Englewood Cliffs.
- [43] Gross, A.J. & Kamins, M. (1967). *Reliability Assessment in the Presence of Reliability Growth*, Memorandum RM-5346-PR, RAND Corporation, Santa Monica.
- [44] Gross, A.J. & Kamins, M. (1968). Reliability assessment in the presence of reliability growth, in *Annals of Assurance Sciences: Symposium on Reliability*, San Francisco, pp. 406–416.
- [45] Virene, E.P. (1968). Reliability growth and its upper limit, *Proceedings of the Annual Reliability and Maintainability Symposium*, Institute of Electrical and Electronics Engineers, pp. 265–270.
- [46] Crow, L.H. (1983). *AMSAA Discrete Reliability Growth Model*, Note 1–83, United States Army Material System Analysis Activity Methodology Office, Aberdeen Proving Ground.
- [47] Finkelstein, J.M. (1983). A logarithmic reliability growth model for single-mission systems, *IEEE Transactions on Reliability* **32**, 154–164.
- [48] Bhattacharyya, G.K., Fries, A. & Johnson, R.A. (1989). Properties of continuous analog estimators for a discrete reliability-growth model, *IEEE Transactions on Reliability* **38**, 373–378.

- [49] Bhattacharyya, G.K. & Ghosh, J.K. (1991). Estimation of parameters in a discrete reliability growth model, *Journal of Statistical Planning and Inference* **29**, 43–53.
- [50] Johnson, R.A. (1991). Confidence estimation with discrete reliability growth models, *Journal of Statistical Planning and Inference* **29**, 87–97.
- [51] Dubman, M. & Sherman, B. (1969). Estimation of parameters in a transient Markov chain arising in a reliability growth model, *Annals of Mathematical Statistics* **40**, 1542–1556.
- [52] Fries, A. (1993). Discrete reliability-growth models on learning-curve property, *IEEE Transactions on Reliability* **42**, 303–306.
- [53] Sen, A. & Fries, A. (1997). Estimation in a discrete reliability growth model under an inverse sampling scheme, *Annals of Institute of Statistical Mathematics* **49**, 211–229.
- [54] Sen, A. & Fries, A. (1998). Model misspecification effects within a family of alternative discrete reliability-growth methodologies, *Lifetime Data Analysis* **4**, 65–81.

Related Articles

Analysis of Recurrent Events from Repairable Systems; Bayesian Reliability Demonstration; Imperfect Repair; Intensity Functions for Non-homogeneous Poisson Processes; Nonparametric and Semiparametric Bayesian Reliability Analysis; Nonparametric Methods for Analysis of Repair Data; Repairable Systems Reliability; Software Failure Data Analysis; Software Reliability Modeling and Analysis; Software Reliability: Bayesian Analysis.

ANANDA SEN

Multicomponent Maintenance

Introduction

Multicomponent maintenance is the problem of finding optimal **maintenance policies** for a system consisting of several units of machines or many pieces of equipment, which may or may not depend on each other [1]. Owing to the existence of interactions between components in complex systems, applying single-component maintenance policies is often not optimal. Generally, three types of interactions or dependencies can be distinguished: economic dependence, stochastic dependence, and structural dependence. Economic dependence implies that grouping maintenance actions either saves costs (e.g. because to economies of scale) or yields higher costs (e.g. because of machine downtime) compared to multiple individual maintenance actions. Stochastic dependence occurs if the condition of components influences the lifetime distribution of other components. Structural dependence applies if components structurally form a part, so that the maintenance of a failed component implies maintenance of working components. Before we discuss these dependencies in more detail, we will first give a historical perspective on multicomponent maintenance.

Historical Perspective

Multicomponent maintenance is one of the key problems in maintenance management of complex systems. Intense scientific interest in problems related to maintenance management originated only a few decades ago. The basis for the scientific support of maintenance is found in reliability engineering. The book *Mathematical Theory of Reliability* [2] indicates the start of scientific interest in maintenance problems. In the early stages, maintenance was almost exclusively interpreted as replacement of components (see **Replacement Strategies**). The development of some simple but insightful models describing aging phenomena of technical components, as well as the choice of appropriate actions to cope with these phenomena, showed that a scientific approach can be

useful. In the 1970s and early 1980s, the basic models were extended in several directions. A substantial part of this work was primarily motivated by mathematical curiosity, rather than practical need.

In the 1980s, scientists made a step forward in bridging the gap between analytical models and practice by a systematic analysis of multicomponent systems. Multicomponent systems are much closer to reality than the classical single-component systems. Since maintenance actions in practice quite often come down to the replacement of one or more components in a multicomponent system, the study of these systems was an adequate way to bring reliability studies closer to practical maintenance management. Moreover, similar to many other areas of applied (stochastic) optimization (see **Reliability of Redundant Systems**), the major advances in information technology brought the use of quantitative analytic models within computational reach.

During the last 15 years substantial progress has been made in developing quantitative decision support systems for maintenance management of complex systems. Sophisticated decision support systems are in operation nowadays in the oil industry (both for maintenance of refineries as well as offshore installations), road and railways maintenance, and electric power generation. The number of companies that make use of some kind of quantitative tools to support inspection and replacement of equipment is increasing rapidly.

Dependencies in Multicomponent Systems

Thomas [3] defines three different types of interactions between components: economic, stochastic, and structural. Below, we will discuss these dependencies in more detail.

Economic Dependence

In simple terms, economic dependence implies that the cost of joint maintenance of a group of components does not equal the total cost of individual maintenance of these components. The effect of this dependence comes to the fore in the execution of maintenance activities. The joint execution of maintenance activities can save costs in some cases, but can lead to higher costs in other cases or it may not be allowed. So, we can further subdivide the models

2 Multicomponent Maintenance

with economic dependence into positive and negative economic dependence. In many multicomponent systems both forms of economic dependence between components exist. As an example, we will discuss the different dependencies in *k-out-of-n* systems.

Positive Economic Dependence. We distinguish between the following forms of positive dependence:

- economies of scale
 - general
 - single setup
 - multiple setups
 - hierarchy of setups
- downtime opportunity.

The term *economies of scale* is often used to indicate that joint maintenance activities are cheaper than multiple individual activities. Economies of scale can result from preparatory or setup activities that can be shared when several components are maintained simultaneously. The cost of this setup work is often called *setup cost*. Setup costs can be saved because the execution of a group of activities requires only one setup. In practice, maintenance of multicomponent systems requires different setup activities and in that case there usually is a hierarchy of setups. For instance, consider a system consisting of two components, which both consist of two subcomponents. Maintenance of the subcomponents of the components may require a setup at system level and component level. Firstly, this means that the setup cost at component level is paid only once when the maintenance of two subcomponents of a component is combined. Secondly, the setup cost at system level is paid only once when all subcomponents are maintained at the same time. In multicomponent maintenance models the setup cost structure is incorporated in the objective function of the corresponding optimization problem.

Another form of positive dependence is the downtime opportunity. Component failures can often be regarded as opportunities for **preventive maintenance** of nonfailed components. In a series system (see **Parallel, Series, and Series-Parallel Systems**) a component failure results in a nonoperating system. In that case it may be worthwhile to replace other components preventively at the same time (opportunistic maintenance). This way the **system downtime** results in savings in maintenance costs

since more components can be replaced at the same time. Moreover, by grouping corrective and preventive maintenance, the downtime can be regulated and in some cases it can even be reduced. In some cases the downtime cost can be included in the setup cost. In general, however, it is difficult to assess the cost associated with the downtime.

Negative Economic Dependence. Negative economic dependence between components occurs when maintaining components simultaneously is more expensive than maintaining components individually. There can be several reasons for this such as:

- manpower restrictions
- safety requirements
- redundancy/production loss.

Firstly, grouping maintenance results in a peak in manpower needs. Manpower restrictions may even be violated and additional labor needs to be hired, which is costly. The problem here is to find the balance between workload fluctuation and grouping maintenance.

Secondly, there are often restrictions on the use of equipment, when executing maintenance activities simultaneously. For instance, use of one equipment may hamper the use of other equipment and cause unsafe operations. Safety requirements often prohibit joint operation.

Thirdly, joint (corrective) maintenance of components in systems in which some kind of redundancy (see **Reliability of Redundant Systems**) is available may not be beneficial. Although there may exist economies of scale through simultaneous repair of a number of (identical) components, leaving components in a failed condition for some time increases the risk of costly production losses. Production loss may increase more than linearly with the number of components out of operation. For an example of this type of economic dependence, we refer to Stengos and Thomas [4].

Economic Dependencies in *k-out-of-n* Systems. The *k-out-of-n* system (see ***k-out-of-n* Systems**) is a typical example of a system with both positive and negative economic dependence between components. A *k-out-of-n* system functions if at least *k* components function. If $k = 1$, then it is a parallel system; if $k = n$, then it is a series system. Let us for

the moment distinguish between the cases $k = n$ and $k < n$.

In the series system ($k = n$, see **Parallel, Series, and Series-Parallel Systems**), positive economic dependence between components is present due to downtime opportunities. The failure of one component results in an expensive downtime of the system and this time can be used to group both preventive and corrective maintenance, i.e. opportunistic grouping. Negative economic dependence is not explicitly present in the series system.

If $k < n$, then there is redundancy in the system and it fails less often than its individual components. In this way a certain reliability can be guaranteed. Typically, the components of this system are identical which allows for economies of scale in the execution of maintenance activities. It is not only possible to obtain savings by grouping preventive maintenance, but also by grouping corrective maintenance. In other words, the redundant components introduce additional positive dependence in the system. Whereas positive economic dependence is present upon failure of a component, negative economic dependence plays a role as long as the system operates. A single failure of a component may not always be an opportunity to combine maintenance activities. Firstly, grouping corrective and preventive maintenance upon the failure of the component increases the probability of system failure and costly production losses. Secondly, leaving components in a failed condition for some time, with the intention to group corrective maintenance at a later stage, has the same effect. So, there is a trade-off between the potential loss resulting from a system failure and the benefit of joint maintenance.

Stochastic Dependence

Stochastic dependence (see **Aging and Positive Dependence**) is also referred to as *failure interaction* and *probabilistic/statistical dependence*. It defines a relationship between components upon failure of a component. For example, it may be the case that the failure of one component induces the failure of other components or causes a shock to other components. Generally, the state of components can influence the state of the other components in the system. Here, the state can be given by the age, the failure rate, state of failure, or another condition measure. In the seminal work on stochastic dependence, Murthy and Nguyen [5, 6] introduce two types of failures: natural

and induced. Natural failures are modeled by random variables and represent “normal” failures because of, e.g. aging or heavy usage. Induced failures are the result of other failures. Typically, there are two ways of modeling the induced failures. Firstly, a (natural) failure of one component induces a failure of other components with probability p and has no effect with probability $1 - p$. Scarf and O’Deara [7, 8] have compared several age- and block-replacement policies (see **Multivariate Age and Multivariate Renewal Replacement; Block Replacement**) for a two-component system with both this type of failure dependence and economic dependence. Secondly, a natural failure of one component may act as a shock to another component, resulting in a higher failure rate or condition of this component. This type of failure interaction is often modeled by a nonhomogeneous Poisson process (see **Poisson Processes**). Özekici [9] proposes a quite general Markov (see **Markov Processes; Maintenance and Markov Decision Models**) model for failure interaction. The state of a component at time t , which is given by a continuous random variable, depends on the state of the system up to time t . The state of the system is again given by the vector of component state values.

Finally, we recognize that the maintenance policies for systems with failure interaction are mainly of an opportunistic nature, since the failure of one component is potentially harmful for other components.

Structural Dependence

Structural dependence (see **Group Maintenance Policies; Total Productive Maintenance**) means that some operating components have to be replaced, or at least dismantled, before failed components can be replaced or repaired. In other words, structural dependence between components indicates that they cannot be maintained independently, but only together. It is not about failure dependence, but about maintenance dependence. Since the failure of a component offers an opportunity to replace other components, opportunistic policies are expected to perform well on systems with structural dependence between components. Obviously, preventive maintenance may also be advantageous, since maintenance of structural dependent components can be grouped.

There may be several reasons for structural dependence. For example, a bicycle chain and a cassette form a union, which should always be replaced

together, rather than individually. Another example is from Dekker *et al.* [10], which considers road maintenance. Several deterioration patterns affect roads, e.g., longitudinal and transversal unevenness, cracking, and raveling. For each mechanism one may define a virtual component, but if one applies a maintenance action to such a component it also affects the state with respect to the other failure mechanisms.

The seminal paper in this category is from Sasieni [11]. He considers the production of rubber tires. The machine that produces the tires consists of two “bladders”; one tire is produced on each bladder simultaneously. Upon failure of a bladder, the machine must be stripped down before replacement can be done. This means that the other bladder can be replaced at the same time. Note that immediate replacement is not mandatory, but a failed bladder will produce faulty tires.

Modeling Multicomponent Maintenance

Several models for the maintenance of multicomponent systems have been proposed. Firstly, these models can be classified based on the planning aspect of the model: stationary (long term) and dynamic (short term). Secondly, one can distinguish three types of optimization methods that can be used to select the model: exact, heuristic, or policy optimization. The typical features of the models in these categories are specified below.

Planning Horizon

In stationary models, a long-term stable situation is assumed and mostly these models assume an infinite planning horizon. This assumption facilitates mathematical analysis. Models of this kind provide static rules for maintenance, which do not change over the planning horizon. They generate for example, long-term maintenance frequencies for groups of related activities or control limits for carrying out maintenance, depending on the state of the components. Three main classes of models in this category can be recognized, namely grouping corrective maintenance, grouping preventive maintenance, and opportunistic maintenance. In the first case, components are only correctively maintained and a failed component can be left in the failed state until its corrective maintenance is carried out jointly with that of other failed

components. In the second case, preventive maintenance is carried out to prevent failures or to decrease operating costs, and this is planned in advance and in such a way that setup costs can be saved by simultaneous execution. In the third case, maintenance is not necessarily planned in advance. However, setup savings can be obtained since (corrective or preventive) maintenance of a component yields an opportunity for maintenance of other components.

In dynamic models, short-term information such as a varying deterioration of components or unexpected opportunities can be taken into account. Situations that are not stationary, such as a varying use of components, and unexpected events that may create an opportunity for doing maintenance at lower costs, can now be incorporated. These models generate dynamic decisions that may change over the planning horizon.

The dynamic-grouping models can be further subdivided into two categories: those with a finite horizon and those with a rolling horizon. Finite-horizon models consider the system in this horizon only, and hence assume implicitly that the system is not used afterward, unless a so-called residual value is incorporated to estimate the industrial value of the system at the end of the horizon. Rolling-horizon models also use a finite horizon, but they do so repeatedly and on the basis of a long-term (infinite-horizon) plan. That is, once decisions of the finite horizon are implemented or when new information becomes available, a new horizon is considered, and a tentative plan based on the long-term one is adapted according to short-term circumstances.

Optimization Method

Exact optimization methods are designed to find the global optimal solution to a problem. However, if the computing time of the optimization method increases exponentially with the number of components, then exact methods are desirable only to a certain extent. In that case, solving problems with many components is impossible and heuristics should be used. Heuristics are local optimization methods that do not pretend to find the global optimum, but can be applied to find a solution to the problem in a reasonable time. The quality of such a solution depends on the problem instance. In some cases, it is possible to give an upper bound on the gap between the optimal solution and the solution found by the heuristic.

In many papers, maintenance planning is done by optimizing a certain type of policy. Well-known maintenance policies are the age- and block-replacement policies (*see Multivariate Age and Multivariate Renewal Replacement; Block Replacement*) and their extensions. The advantage of policy optimization over other optimization methods is that it gives more insight into the solution of the problem. Policy optimization will not always result in the global optimal solution, since there may be another policy that results in a better solution.

Heuristic and exact methods are often used in finite-horizon models, especially when time is discretized and the decision parameters only take integer values. Policy optimization is mainly used in infinite-horizon models. In these models, it is often possible to derive analytical expressions for optimal control parameters and the corresponding optimal costs.

Example

We will now illustrate the preceding review of multicomponent maintenance models with an example. It is based on two studies by Scarf and O'Deara [7, 8], who consider the maintenance of a clutch system in a bus fleet. The system is a series system consisting of the clutch assembly (component 2) and the clutch controller (component 1). In this particular system a (natural) failure of component 1 causes a failure of component 2 with probability 1, whereas a failure of component 2 *never* induces a failure of component 1. Moreover, the system exhibits (positive) economic dependence in that combining (preventive or failure) replacements saves costs. Replacement restores the respective component to a new condition. Below, we highlight two replacement policies for this system: combined failure-based replacement (CFR) and combined age-based replacement (CAR; *see Multivariate Age and Multivariate Renewal Replacement*).

Note that because we deal with a series system (*see Parallel, Series, and Series-Parallel Systems*), a system failure implies that either both components have failed or just one of the components has failed. It is assumed that failed components are replaced immediately. CFR now prescribes that either (a) both failed components are replaced or (b) the failed component is replaced correctively and the component that has not failed is replaced preventively. CAR dictates that

both components are replaced at age T or upon failure, whichever occurs first. Observe that both policies prescribe that components that have not failed are replaced preventively. Scarf and O'Deara [7] also consider age- and failure-based replacement policies that do not replace components that have not failed preventively. Since these policies do not take into account the positive economic dependence between components, they are outperformed by their extensions (CAR and CFR, respectively).

Let us introduce some notation. The cost of a preventive replacement of both components is given by C_{PP} and is less than or equal to the independent preventive replacement of both components. The cost of failure replacement of both components C_{FF} , is assumed to be equal to the cost of combined failure and preventive replacements of the components. The lifetime distribution $F_i(x)$ of component i is the Weibull distribution with parameters α_i and β_i , $i = 1, 2$. That is, $F_i(x) = 1 - \exp(-(x/\beta_i)^{\alpha_i})$, $i = 1, 2$. The lifetime is measured in months. Now, it can be shown that the mean cost rates for CFR and CAR are given by

$$\begin{aligned} C_{\infty}^{\text{CFR}} &= \frac{C_{FF}}{\int_0^{\infty} (1 - F_1(x))(1 - F_2(x)) dx} \quad \text{and} \\ C_{\infty}^{\text{CAR}}(T) &= \frac{C_{FF} + (C_{PP} - C_{FF})(1 - F_1(T))(1 - F_2(T))}{\int_0^T (1 - F_1(x))(1 - F_2(x)) dx} \end{aligned} \quad (1)$$

respectively. Minimizing $C_{\infty}^{\text{CAR}}(T)$ with respect to T yields the optimum replacement age.

Figure 1 shows the (optimal) mean cost rate for CFR and CAR for the lifetime parameters as estimated in [7, 8]. Cost parameter C_{PP} has been set to 2740 and C_{FF} has been varied between 4420 and 6420 (this corresponds with the case studies by Scarf and O'Deara [7], where the unavailability cost is 1600 and the penalty cost of a failure is varied between 0 and 2000). As expected CAR outperforms CFR and the difference becomes more pronounced as the cost of failure increases. Scarf and O'Deara [7] show that the performance (with respect to both mean cost rate and implementation) of CAR is comparable with that of well-performing

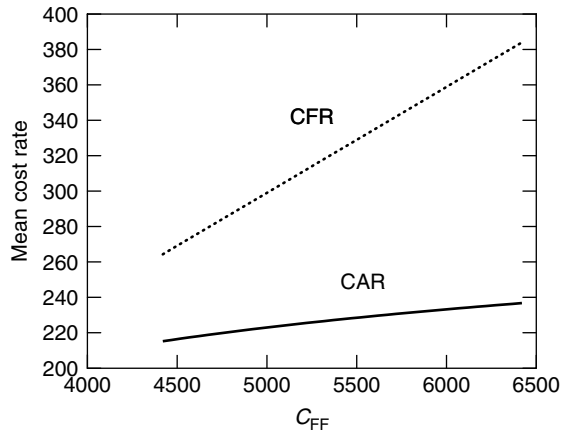


Figure 1 Mean cost rate for replacement policies CFR and CAR versus C_{FF} (system failure replacement cost). The Weibull parameters are given by $\alpha_1 = 8$, $\beta_1 = 34.8$, $\alpha_2 = 7.7$, and $\beta_2 = 17.8$. The cost of preventive replacement of both components is $C_{PP} = 2740$. The cost of combined failure and preventive replacements C_{FF} is varied between 4420 and 6420

opportunistic age-replacement policies. Notice that CAR treats the system as a single component and prescribes to replace it at age T . As such it is easy to implement in a decision support system. Moreover, a further study [8] shows that CAR competes with the best (modified) block-replacement policies, which are also relatively easy to manage.

Further Reading

Several overview articles have appeared on multicomponent maintenance models. Cho and Parlar [1] review articles from 1976 up to 1991. It divides the literature into five topical categories: machine-interference/repair models, group/block/cannibalization/opportunistic models, inventory/maintenance models, other maintenance/replacement models, and inspection/maintenance models. Dekker *et al.* [12] exclusively deal with multicomponent maintenance models that are based on economic dependence. The models are classified on the basis of their planning aspect: stationary (long term) or dynamic (short term). Wang [13] gives an overview of maintenance policies of deteriorating systems. The emphasis is on policies for single-component systems. One

section is devoted to opportunistic maintenance policies for multicomponent systems. The author primarily considers models with economic dependence. In a recent article, Nicolai and Dekker [14] review articles in which multicomponent maintenance of systems with economic (both positive and negative), stochastic, and structural dependence is modeled (and optimized).

References

- [1] Cho, D. & Parlar, M. (1991). A survey of maintenance models for multi-unit systems, *European Journal of Operational Research* **51**, 1–23.
- [2] Barlow, R.E. & Proschan, F. (1965). *Mathematical Theory of Reliability*, John Wiley & Sons, New York.
- [3] Thomas, L. (1986). A survey of maintenance and replacement models for maintainability and reliability of multi-item systems, *Reliability Engineering* **16**, 297–309.
- [4] Stengos, D. & Thomas, L. (1980). The blast furnaces problem, *European Journal of Operational Research* **4**, 330–336.
- [5] Murthy, D.N.P. & Nguyen, D. (1985a). Study of a multicomponent system with failure interaction, *European Journal of Operational Research* **21**, 330–338.
- [6] Murthy, D.N.P. & Nguyen, D. (1985b). Study of two-component system with failure interaction, *Naval Research Logistics Quarterly* **32**, 239–247.
- [7] Scarf, P.A. & O'Deara, M. (1998). On the development and application of maintenance policies for a two-component system with failure dependence, *IMA Journal of Mathematics Applied in Business and Industry* **9**, 91–107.
- [8] Scarf, P.A. & O'Deara, M. (2003). Block replacement policies for a two-component system with failure dependence, *Naval Research Logistics* **50**(1), 70–87.
- [9] Özekici, S. (1988). Optimal periodic replacement of multicomponent reliability systems, *Operations Research* **36**, 542–552.
- [10] Dekker, R., Plasmeijer, R. & Swart, J. (1998). Evaluation of a new maintenance concept for the preservation of highways, *IMA Journal of Mathematics Applied in Business and Industry* **9**, 109–156.
- [11] Sasieni, M. (1956). A Markov chain process in industrial replacement, *Operational Research Quarterly* **7**, 148–155.
- [12] Dekker, R., van der Duyn Schouten, F. & Wildeman, R. (1996). A review of multi-component maintenance models with economic dependence, *Mathematical Methods of Operations Research* **45**, 411–435.
- [13] Wang, H. (2002). A survey of maintenance policies of deteriorating systems, *European Journal of Operational Research* **139**, 469–489.
- [14] Nicolai, R. & Dekker, R. (2007). Optimal maintenance of multi-component systems: a review, in *Maintenance*

of Complex Systems: Blending Theory with Practice, K.A.H. Kobbacy & D.N.P. Murthy, eds, Springer-Verlag, Berlin.

Related Article

Aging and Positive Dependence; Block Replacement; Group Maintenance Policies; Inspection

Policies for Reliability; k -out-of- n Systems; Maintenance and Markov Decision Models; Maintenance Optimization; Maintenance Optimization in Random Environments; Markov Processes; Multivariate Age and Multivariate Renewal Replacement; Poisson Processes; Reliability of Redundant Systems; Replacement Strategies; Total Productive Maintenance.

ROBIN P. NICOLAI AND ROMMERT DEKKER

Markov Renewal Processes in Reliability Modeling

Introduction and Mission Process

In this article, we consider a reliability system that is designed to perform a random sequence of phases or stages with random durations. These are referred to as *phased-mission* or *mission-based systems* where the stochastic structure of the mission process plays a critical role. It is assumed that the mission follows a Markov **renewal process**, and we will analyze various issues related to the reliability of such systems. This line of stochastic modeling and research is often classified under stochastic models in random environments. The idea is used not only in a reliability setting, but in other applications as well. For example, Özekici [1] discusses other applications in inventory and queuing. Asmussen [2] and Rolski *et al.* [3] also give further applications in queuing, insurance, and finance.

In reliability modeling, a device generally consists of a large number of components with stochastically dependent lifetimes. **Random environments** are used to provide a tractable model of dependence since this is taken as an external process that affects the deterioration, aging, and failure of all of the components. Since all components are subjected to the same environmental conditions, their lifetimes are dependent *via* their common environmental process. Thus, the environmental process is actually a factor of variation in the failure structure of the components. An interesting model was introduced by Çınlar and Özekici [4] where stochastic dependence is introduced by a randomly changing common environment that all components of the system are subjected to. This model is based on the simple observation that the aging or deterioration process of any component depends very much on the environment that the component is operating in. They propose to construct an intrinsic clock that ticks differently in different environments to measure the intrinsic age of the device. The environment is modeled by a semi-Markov jump process and the intrinsic age is represented by the cumulative hazard accumulated in time during the operation of the device in the randomly varying

environment. This is a rather stylish choice, which envisions that the intrinsic lifetime of any device has an exponential distribution with parameter 1. The concept of random hazard functions is also used in Gaver [5] and Arjas [6]. The intrinsic aging model of Çınlar and Özekici [4] is studied further in Çınlar *et al.* [7] to determine the conditions that lead to associated component lifetimes, as well as multivariate increasing failure rate (IFR) and new better than used (NBU) life distribution characterizations. It was also extended in Shaked and Shanthikumar [8] by discussions on several different models with multicomponent replacement policies. Lindley and Singpurwalla [9] discuss the effects of the random environment on the reliability of a system consisting of components that share the same environment. Although the initial state of the environment is random, they assume that it remains constant in time and components have exponential life distributions in each possible environment. This model is also studied by Lefèvre and Malice [10] to determine partial orderings on the number of functioning components and the reliability of k -out-of- n systems, for different partial orderings of the probability distribution on the environmental state. The association of the lifetimes of components subjected to a randomly varying environment is discussed in Lefèvre and Milhaud [11]. Singpurwalla and Youngren [12] also discuss multivariate distributions that arise in models where a dynamic environment affects the failure rates of the components.

Reliability models in random environments also provide a perfect opportunity to analyze mission-based or so-called phased-mission reliability systems. Our aim in this article is to present an example along this direction. These systems involve devices or machines that are designed to perform or assigned to missions consisting of a number of phases. The sequence as well as the durations of the phases may be random and, as in all random environment models, all stochastic and deterministic failure properties depend on the phases of the mission that is performed at a given time. Therefore, the random environmental process is the mission process in such models. A phased-mission consists of a number of stages or phases that must be successfully completed. These systems were introduced by Esary and Ziehms [13] and a vast literature has accumulated since then. The phases may be deterministic

2 Markov Renewal Processes in Reliability Modeling

or stochastic as in Kim and Park [14]. The components are either repairable or nonrepairable. Alam and Al-Saggaf [15] discuss a phased-mission system with **repairable components** where the repair activity begins as soon as a component fails. A case study involving spacecraft is provided by Mura and Bondavalli [16] where the mission consists of four basic phases: launch, hibernation, planet, and scientific observation. The reader is referred to these papers and the many references cited in them for various models, techniques, and applications on phased-mission systems.

Let X_n denote the n th phase of the mission and T_n denote the time at which the n th phase starts with $T_0 \equiv 0$. The main assumption is that the process $(X, T) = \{(X_n, T_n); n \geq 0\}$ is a Markov renewal process on the countable state space E with some semi-Markov kernel Q . The state space E is actually that of the process X and it is implicitly understood that the process T always takes values in $R_+ = [0, +\infty)$ since they denote times at which certain events occur. We refer the reader to Çinlar [17] for a more rigorous and detailed treatment of Markov renewal processes and theory. The Markov renewal property states that

$$\begin{aligned} P\{X_{n+1} = j, T_{n+1} - T_n \leq t | X_0, \dots, X_n; T_0, \dots, T_n\} \\ = P\{X_{n+1} = j, T_{n+1} - T_n \leq t | X_n\} \end{aligned} \quad (1)$$

where we suppose that the process is time homogeneous with the semi-Markov kernel

$$Q(i, j, t) = P\{X_{n+1} = j, T_{n+1} - T_n \leq t | X_n = i\} \quad (2)$$

for any $i, j \in E$ and $t \geq 0$. It is well known that X is a Markov chain on E with transition matrix $P(i, j) = P\{X_{n+1} = j | X_n = i\} = Q(i, j, +\infty)$. We further assume that the Markov renewal process has infinite lifetime so that $\sup_n T_n = +\infty$. The probabilistic structure of T is given by the conditional distributions

$$\begin{aligned} G(i, j, t) &= P\{T_{n+1} - T_n \leq t | X_n = i, X_{n+1} = j\} \\ &= \frac{Q(i, j, t)}{P(i, j)} \end{aligned} \quad (3)$$

for $P(i, j) > 0$. Moreover, it is easy to see that

$$F_i(t) = P\{T_1 \leq t | X_0 = i\} = \sum_{j \in E} Q(i, j, t) \quad (4)$$

denotes the distribution of the duration of phase i . Finally, the mission process $Y = \{Y_t; t \geq 0\}$ is the minimal semi-Markov process associated with (X, T) so that Y_t is the stage or phase of the mission at time t . More precisely, $Y_t = X_n$ whenever $T_n \leq t < T_{n+1}$.

To simplify the notation, for any event A we will let $P_i\{A\} = P\{A | X_0 = i\}$ denote the conditional probability of A given $X_0 = i$. The reliability of a mission-based system will be analyzed with two extreme repair policies: maximal repair and no repair. The system is replaced by a brand-new one at the beginning of each phase in the first case whereas no repair or replacement is done in the second case until the system fails. We discuss system and mission reliability with maximal repair in the second section. Mission reliability can be defined in different ways. We consider two alternatives where we focus on the probability of completing the first n phases successfully in the first one, and the probability of completing a given critical phase in the second alternative. These issues are discussed for the no repair case in the third section using intrinsic aging concepts.

Models with Maximal Repair

Suppose that the system performs the mission such that at the beginning of each phase it is replaced by a brand-new one. This simplifying assumption allows us to obtain various reliability measures quite easily by using the renewal property. It is clear that the probability of survival during a given phase depends on it, and we let

$$\bar{P}_i(t) = 1 - P_i(t) = P_i\{L > t | Y_s = i; s \in [0, t]\} \quad (5)$$

denote the survival probability of the system in phase i for t units of time where L is the lifetime of the system.

We can obtain another Markov renewal processes using (X, T) , which will be used during the analysis. Define a new Markov renewal process $(\tilde{X}, \tilde{T})$

through its minimal semi-Markov process $\tilde{Y}$ so that

$$\tilde{Y}_t = \begin{cases} Y_t & \text{if } t < L \\ \Delta & \text{if } t \geq L \end{cases} \quad (6)$$

where Δ is an absorbing state. This also implies that

$$\tilde{T}_n = \begin{cases} 0 & n = 0 \\ \inf\{t > T_{n-1} : \tilde{Y}_t \neq \tilde{Y}_{T_{n-1}}\} & n \geq 1 \end{cases} \quad (7)$$

and $\tilde{X}_n = \tilde{Y}_{\tilde{T}_n}$ for $n \geq 0$. Clearly, the state space of $(\tilde{X}, \tilde{T})$ is $\tilde{E} = E \cup \{\Delta\}$ and its semi-Markov kernel is

$$\begin{aligned} \tilde{Q}(i, k, ds) &= \begin{cases} \bar{P}_i(s) \tilde{Q}(i, k, ds) & \text{if } i, k \in E \\ \bar{F}_i(s) P_i(ds) & \text{if } i \in E, k = \Delta \\ 0 & \text{if } i, k = \Delta \end{cases} \end{aligned} \quad (8)$$

Note that $\tilde{Q}(\Delta, \Delta, t) = 0$ for every $t \geq 0$ and $\tilde{P}(\Delta, \Delta) = 1$. In plain words, the new process $\tilde{Y}$ is obtained by “stopping” or “killing” the process Y whenever the system fails at time L . At the time of failure, the process is dumped to the absorbing state Δ that denotes system failure. The state space is therefore extended by including this new state Δ . We can find the transition matrix of the Markov chain $\tilde{X}$ as

$$\begin{aligned} \tilde{P}(i, j) &= \tilde{Q}(i, j, \infty) = \int_0^\infty \tilde{Q}(i, j, ds) \\ &= \int_0^\infty \bar{P}_i(s) \tilde{Q}(i, j, ds) \end{aligned} \quad (9)$$

for $i, j \in E$, and

$$\tilde{P}(i, \Delta) = \int_0^\infty \bar{F}_i(s) P_i(ds) \quad (10)$$

for $i \in E$.

Suppose that one is interested in the successful completion of a given critical phase j of the mission. Letting

$$U_j = \inf\{t \geq 0; Y_t \neq Y_{t-} = j\} \quad (11)$$

denote the first time that the mission process leaves state j , we can now define another new Markov renewal process $(\bar{X}, \bar{T})$ as appropriate through its

minimal semi-Markov process $\bar{Y}$. The new process is defined by

$$\bar{Y}_t = \begin{cases} Y_t & \text{if } t < \min\{L, U_j\} \\ \Delta & \text{if } L \leq \min\{t, U_j\} \\ S & \text{if } U_j \leq \min\{t, L\} \end{cases} \quad (12)$$

so that we now extend E such that $\bar{E} = E \cup \{\Delta, S\}$. We let the semi-Markov mission process Y go to the absorbing state Δ as soon as the system fails before completing phase j or to the absorbing state S as soon as phase j is completed without failure. Thus, if $\bar{X}_n = \Delta$ or $\bar{X}_n = S$, then $\bar{T}_{n+1} = \infty$. If the minimal semi-Markov process associated with this Markov renewal process is $\bar{Y}$, then the probability of completing the critical phase j until time t is $P_i\{\bar{Y}_t = S\}$ if the initial state is i . The Markov renewal kernel of this process can be obtained as

$$\bar{Q}(i, k, ds) = \begin{cases} \tilde{Q}(i, k, ds) & i \in E - \{j\}, k \in E \\ \tilde{Q}(i, \Delta, ds) & i \in E, k = \Delta \\ F_j(ds) \bar{P}_j(s) & i = j, k = S \\ 0 & \text{otherwise} \end{cases} \quad (13)$$

Note that $\bar{Q}(S, S, t) = \bar{Q}(\Delta, \Delta, t) = 0$ for every $t \geq 0$, and $\bar{P}(S, S) = \bar{P}(\Delta, \Delta) = 1$. If $i \in E - \{j\}, k \in E$, then phase i is completed successfully before system failure and hence $\bar{Q}(i, k, ds) = \tilde{Q}(i, k, ds)$. If $i \in E$ and $k = \Delta$, then $\bar{Q}(i, k, ds)$ represents the probability of failure during phase i in the vicinity of time s , which means that the duration of phase i is longer than s units of time and a failure will occur in the vicinity of time s . Therefore, this probability is equal to $\tilde{Q}(i, \Delta, ds) = \bar{F}_i(s) P_i(ds)$. Moreover, $\bar{Q}(j, S, ds)$ represents the probability that the system will complete phase j in the vicinity of time s given that the system is in working condition at the beginning of phase j . This situation occurs if the duration of mission j is in the vicinity of time s and the system survives more than s units of time in these conditions; hence, $\bar{Q}(j, S, ds) = F_j(ds) \bar{P}_j(s)$. It is clear from the definition of the process $(\bar{X}, \bar{T})$ that the states S and Δ are absorbing states. Therefore, if the process gets into states S or Δ , it remains there forever where S now represents the “success” state and Δ represents the “failure” state.

4 Markov Renewal Processes in Reliability Modeling

System Reliability

System reliability is given by the probability that the system will function until time t . Let $f(i, t) = P_i \{L > t\}$ denote the desired probability. Then, conditioning on T_1

$$\begin{aligned} f(i, t) &= P_i \{L > t, T_1 > t\} + P_i \{L > t, T_1 \leq t\} \\ &= P_i \{L > t | T_1 > t\} P_i \{T_1 > t\} \\ &\quad + P_i \{L > t, T_1 \leq t\} = \bar{P}_i(t) \bar{F}_i(t) \\ &\quad + \sum_{j \in E} \int_0^t Q(i, j, ds) \bar{P}_i(s) f(j, t-s) \\ &= g(i, t) + \tilde{Q} * f(i, t) \end{aligned} \quad (14)$$

where $g(i, t) = \bar{P}_i(t) \bar{F}_i(t)$ and $\tilde{Q}(i, j, ds) = Q(i, j, ds) \bar{P}_i(s)$ for $i \in E$. Clearly, $f(\Delta, t) = g(\Delta, t) = 0$. Then, equation (14) is a Markov renewal equation and it has the unique solution $f = \tilde{R} * g$ so that

$$\begin{aligned} P_i \{L > t\} &= \sum_{j \in E} \int_0^t \tilde{R}(i, j, ds) g(j, t-s) \\ &= \sum_{j \in E} \int_0^t \tilde{R}(i, j, ds) \bar{P}_j(t-s) \bar{F}_j(t-s) \end{aligned} \quad (15)$$

where $\tilde{R} = \sum_{n=0}^{\infty} \tilde{Q}^n$ is the Markov renewal kernel corresponding to $\tilde{Q}$.

Mission Reliability

In a given application, it may be important to determine the probability that the system will complete the first n phases successfully. We will now focus on this issue and show that this probability can be calculated using a recursive formula and then obtain an explicit solution. We first find the probability that the first phase will be completed without failure. This can be done by conditioning on the next phase after the first one so that

$$P_i \{L > T_1\} = \sum_{j \in E} P_i \{L > T_1, X_1 = j\}$$

$$\begin{aligned} &= \sum_{j \in E} \int_0^{\infty} \bar{P}_i(s) Q(i, j, ds) \\ &= \int_0^{\infty} \bar{P}_i(s) F_i(ds) \end{aligned} \quad (16)$$

Now, we can write,

$$\begin{aligned} P_i \{L > T_{n+1}\} &= \sum_{j \in E} P_i \{L > T_{n+1}, X_1 = j\} \\ &= \sum_{j \in E} P_j \{L > T_n\} P_i \{X_1 = j, L > T_1\} \\ &= \sum_{j \in E} P_j \{L > T_n\} \int_0^{\infty} \bar{P}_i(s) Q(i, j, ds) \end{aligned} \quad (17)$$

which is a recursive relationship with the boundary condition $P_i \{L > T_0\} = 1$ for $n \geq 0$. Using equations (9) and (17), we have

$$P_i \{L > T_{n+1}\} = \sum_{j \in E} \tilde{P}(i, j) P_j \{L > T_n\} \quad (18)$$

Letting $f_n(i) = P_i \{L > T_n\}$, we can rewrite equation (18) as

$$f_{n+1} = \tilde{P} f_n \quad (19)$$

for $n \geq 0$ with the boundary condition $f_0 = 1$. Then, it is clear that $f_1 = \tilde{P} f_0$, $f_2 = \tilde{P} f_1 = \tilde{P}^2 f_0$, and, more generally, $f_n = \tilde{P} f_{n-1} = \tilde{P}^n f_0$ so that the solution is

$$f_n(i) = \tilde{P}^n f_0(i) = \sum_{j \in E} \tilde{P}^n(i, j) \quad (20)$$

It is also possible to obtain the same solution by noting that

$$P_i \{L > T_n\} = P_i \{\tilde{X}_n \in E\} = \sum_{j \in E} \tilde{P}^n(i, j) \quad (21)$$

since the n th phase of the mission is successfully completed if $\tilde{X}_n \neq \Delta$.

Critical Phase Reliability

For a complex system, a given phase can be more important than the others because of the overall objective of the mission. Therefore, the probability that this phase will be completed in a fixed period is

an important measure to represent the reliability of the system. To analyze this case, define

$$W_j(t) = \begin{cases} 1 & \text{if phase } j \text{ is completed before time } t \\ 0 & \text{otherwise} \end{cases} \quad (22)$$

where j is the critical phase. Then, using the definition of $(\bar{X}, \bar{T})$, the desired probability is

$$P_i \{W_j(t) = 1\} = P_i \{\bar{Y}_t = S\} \quad (23)$$

and, using Proposition (10.5.4) in [18], we have

$$P_i \{W_j(t) = 1\} = P_i \{\bar{Y}_t = S\} = \bar{R}(i, S, t) \quad (24)$$

where $\bar{R} = \sum_n \bar{Q}^n$ is the Markov renewal kernel corresponding to $\bar{Q}$.

Models with No Repair

In the previous section, we considered systems with maximal repair so that we had a brand-new system at the beginning of each phase. In this section, we remove this assumption and, hence, the system will deteriorate or get older in time. Therefore, the concept of “aging” comes into consideration. In this study, we will use the “intrinsic aging” model introduced by Çinlar and Özekici [4]. The analysis will be presented for a series system with m components, which requires that all components must function for the system to be functional.

Before the analysis, some new notation should be introduced. Let $H_k(i, t)$ be the cumulative hazard of component k at time t in phase i . Then, the intrinsic age of component k at time t is defined as $H_k(i, t)$ provided that the system performs phase i throughout $[0, t]$. The intrinsic aging rate of component k during phase i is defined as

$$r_k(i, a) = \frac{d}{dt} H_k(i, t) \big|_{t=\tau(x, a)} \quad (25)$$

at any age $a \geq 0$, where $\tau_k(i, a)$ is the time at which the intrinsic age of component k becomes a if the system performs phase i ; or

$$\tau_k(i, a) = \inf \{t \geq 0 : H_k(i, t) > a\} \quad (26)$$

Let $A_t(k)$ denote the intrinsic age process of component k . We assume that the intrinsic age process satisfies

$$\frac{dA_t(k)}{dt} = r_k(Y_t, A_t(k)) \quad (27)$$

Therefore, if the intrinsic age of component k at time s when phase i starts is $A_s(k) = a$, then after t units of time its age becomes

$$A_{s+t}(k) = h_k(i, a, t) = H_k(i, \tau_k(i, a) + t) \quad (28)$$

Let $B_0(k) = A_0(k)$ and define an embedded process $\{B_n(k)\}$ recursively through

$$B_{n+1}(k) = h(X_n, B_n(k), T_{n+1} - T_n) \quad (29)$$

for $n \geq 0$. Then, the intrinsic age of component k at time t is given by

$$A_t(k) = h(X_n, B_n(k), t - T_n) \quad (30)$$

whenever $T_n \leq t \leq T_{n+1}$. Let $L(k)$ and $\hat{L}(k)$ denote the lifetime and the intrinsic lifetime of component k respectively. Then,

$$\hat{L}(k) = A_{L(k)}(k) \quad (31)$$

for all k . Following Çinlar and Özekici [4], $\hat{L}(k)$ is exponentially distributed with rate 1 for all k and

$$\begin{aligned} P\{L(k) > t | A_0(k) = a(k)\} \\ &= P\{\hat{L}(k) > A_t(k) | A_0(k) = a(k)\} \\ &= E[e^{-(A_t(k) - a(k))}] \end{aligned} \quad (32)$$

Let A_t , a , $h(i, a, t)$, L , and $\hat{L}$ denote the vectors with elements $A_t(k)$, $a(k)$, $h_k(i, a, t)$, $L(k)$, and $\hat{L}(k)$ for all k , respectively. We let $F = R_+^m$ denote the collection of all nonnegative vectors where m is the number of components in the system. Each a in F represents a vector of intrinsic ages of all components. Assume that if b is a vector, then $b > x$ for a scalar x means that $b(k) > x$ for all k ; and if a, b are both vectors in F , then $b > a$ means that $b(k) > a(k)$ for all k . We let σ be a row vector with m entries which are all equal to 1, and

$$H(i, a, t; db) = \begin{cases} 1 & \text{if } h(i, a, t) = b \\ 0 & \text{otherwise} \end{cases} \quad (33)$$

To simplify our notation, for any event A we will let $P_{ia}\{A\} = P\{A | X_0 = i, A_0 = a, \hat{L} > a\}$ denote the

6 Markov Renewal Processes in Reliability Modeling

conditional probability of A given $X_0 = i$, $A_0 = a$, $\hat{L} > a$. Our analysis in the previous section used the fact that (X, T) is a Markov renewal process to obtain Markov renewal characterizations on various reliability measures. This analysis was possible because of the maximal repair assumption so that the system was brand new at the beginning of each phase. However, when this is not the case, we need to base our analysis on the new Markov renewal process $((X, B), T)$ that also includes information on the intrinsic ages of the components. The state space is now extended to $E \times F$.

System Reliability of a Series System

We first consider the system reliability of a series system that is performing a mission. In other words, we want to determine the probability that the system will work without failure until time t . Using a Markov renewal characterization, we define $f(i, a, t) = P_{ia} \{L > t\}$ and obtain

$$\begin{aligned} P_{ia} \{L > t\} &= P_{ia} \{L > t, T_1 > t\} \\ &+ P_{ia} \{L > t, T_1 \leq t\} = \bar{F}_i(t) e^{-\sigma(h(i,a,t)-a)} \\ &+ \sum_{j \in E} \int_{F \times [0,t]} Q(i, j, ds) H(i, a, s; db) \\ &\times e^{-\sigma(b-a)} f(j, b, t-s) \end{aligned} \quad (34)$$

which is a Markov renewal equation of the form $f = g + \hat{Q} * f$, where

$$g(i, a, t) = \bar{F}_i(t) e^{-\sigma(h(i,a,t)-a)} \quad (35)$$

and

$$\begin{aligned} \hat{Q}(i, a, j, db, ds) &= Q(i, j, ds) H(i, a, s; db) \\ &\times e^{-\sigma(b-a)} \end{aligned} \quad (36)$$

It is clear that $\hat{Q}$ is a semi-Markov kernel on the extended state space $\tilde{E} \times F$ and the Markov renewal equation has the unique solution $f = \hat{R} * g$ such that

$$\begin{aligned} P_{ia} \{L > t\} &= \sum_{j \in \tilde{E}} \int_{F \times [0,t]} \hat{R}(i, a; j, db; ds) \\ &\times \bar{F}_j(t-s) e^{-\sigma(h(j,b,t-s)-b)} \end{aligned} \quad (37)$$

where $\hat{R} = \sum_n \hat{Q}^n$ is the Markov renewal kernel corresponding to $\hat{Q}$.

Mission Reliability of a Series System

We now focus on computing mission reliability involving the first n phases of the mission. We first find the probability that the first phase will be completed by conditioning on the next phase so that

$$\begin{aligned} P_{ia} \{L > T_1\} &= \sum_{j \in E} P_{ia} \{L > T_1, X_1 = j\} \\ &= \sum_{j \in E} P_{ia} \{L > T_1 | X_1 = j\} P(i, j) \end{aligned} \quad (38)$$

Note that

$$\begin{aligned} P_{ia} \{L > T_1 | X_1 = j\} &= \int_{F \times [0,\infty)} P_{ia} \{L > T_1, T_1 \in ds, B_1 \in db | X_1 = j\} \\ &= \int_{F \times [0,\infty)} e^{-\sigma(b-a)} G(i, j, ds) H(i, a, s; db) \end{aligned} \quad (39)$$

Combining equations (38) and (39), we obtain

$$\begin{aligned} P_{ia} \{L > T_1\} &= \sum_{j \in E} \int_{F \times [0,\infty)} Q(i, j, ds) \\ &\times e^{-\sigma(b-a)} H(i, a, s; db) \end{aligned} \quad (40)$$

Similarly, we can now obtain a recursive relationship

$$\begin{aligned} P_{ia} \{L > T_{n+1}\} &= \sum_{j \in E} P_{ia} \{L > T_{n+1} | X_1 = j\} P(i, j) \\ &= \sum_{j \in E} \int_{F \times [0,\infty)} P_{ia} \{L > T_{n+1}, T_1 \in ds, \\ &\quad B_1 \in db | X_1 = j\} P(i, j) \\ &= \sum_{j \in E} \int_{F \times [0,\infty)} e^{-\sigma(b-a)} Q(i, j, ds) \\ &\quad \times H(i, a, s; db) P_{jb} \{L > T_n\} \end{aligned} \quad (41)$$

for $n \geq 0$. Let $\hat{Q}(i, a; j, db; ds) = Q(i, j, ds) H(i, a, s; db) e^{-\sigma(b-a)}$, and define

$$\hat{P}(i, a; j, db) = \hat{Q}(i, a; j, db; \infty)$$

$$= P_{ia} \{X_1 = j, B_1 \in db\} \quad (42)$$

and

$$\begin{aligned} \hat{P}^n(i, a; j, db) &= \sum_{j \in E} \int_F \hat{P}^{n-1}(i, a; k, dc) \hat{P}(k, c; j, db) \\ &= P_{ia} \{X_n = j, B_n \in db\} \end{aligned} \quad (43)$$

for $n \geq 1$.

Then, it follows from equation (40) that

$$\begin{aligned} P_{ia} \{L > T_1\} &= \sum_{j \in E} \int_{F \times [0, \infty)} \hat{Q}(i, a; j, db; ds) \\ &= \sum_{j \in E} \int_F \hat{Q}(i, a; j, db; \infty) \\ &= \sum_{j \in E} \int_F \hat{P}(i, a; j, db) \end{aligned} \quad (44)$$

Now, using induction, we will show that

$$P_{ia} \{L > T_n\} = \sum_{j \in E} \int_F \hat{P}^n(i, a; j, db) \quad (45)$$

for every $n \geq 1$. We have already shown that equation (45) holds for $n = 1$. Suppose that

$$P_{ia} \{L > T_k\} = \sum_{j \in E} \int_F \hat{P}^k(i, a; j, db) \quad (46)$$

for all $k \leq n - 1$. Then, using equations (41) and (46),

$$\begin{aligned} P_{ia} \{L > T_n\} &= \sum_{j \in E} \int_{F \times [0, \infty)} \hat{Q}(i, a; j, db; ds) P_{jb} \{L > T_{n-1}\} \\ &= \sum_{j \in E} \int_F \hat{P}(i, a; j, db) \sum_{k \in E} \int_F \hat{P}^{n-1}(j, b; k, dc) \\ &= \sum_{j \in E} \int_F \hat{P}^n(i, a; j, db) \end{aligned} \quad (47)$$

Critical Phase Reliability of a Series System

Suppose that there is a critical phase j and we are interested in computing the probability

$P_{ia} \{W_j(t) = 1\}$ that phase j will be completed successfully by time t . Using a similar approach as in the critical reliability part of the previous section, we note that the process $\bar{Y}$ defined by equation (12) is a semiregenerative process. In this case, the semi-Markov kernel corresponding to the Markov renewal process $((\bar{X}, B), \bar{T})$ is given by

$$\bar{Q}(i, a; k, db; ds) = \begin{cases} Q(i, k, ds) H(i, a, s; db) \times e^{-\sigma(b-a)} & \text{if } i \in E - \{j\}, \\ & k \in E \\ \bar{F}_i(s) H(i, a, s; db) U_{ia}(ds) & \text{if } i \in E, k = \Delta \\ F_j(ds) H(i, a, s; db) e^{-\sigma(b-a)} & \text{if } i = j, k = S \\ 0 & \text{otherwise} \end{cases} \quad (48)$$

where

$$U_{ia}(s) = 1 - e^{-\sigma(h(i, a, s) - a)}. \quad (49)$$

Note that $\bar{Q}(S, a; S, db; t) = \bar{Q}(\Delta, a; \Delta, db; t) = 0$ trivially for every $t \geq 0$, $a, b \in F$ and $\bar{P}(S, a; S, da) = \bar{P}(\Delta, a; \Delta, da) = 1$ for every $a \in F$. We now find the probability $f(i, a, t) = P_{ia} \{\bar{Y}_t = j, A_t \in db\}$ for some $j \in \bar{E}$ and $b \in F$. Using Markov renewal theory,

$$\begin{aligned} f(i, a, t) &= P_{ia} \{\bar{Y}_t = j, A_t \in db, \bar{T}_1 > t\} \\ &\quad + P_{ia} \{\bar{Y}_t = j, A_t \in db, \bar{T}_1 \leq t\} \\ &= I(i, j) I(a, b) q(i, a, t) \\ &\quad + \sum_{l \in \bar{E}} \int_{F \times [0, t]} \bar{Q}(i, a; l, dc; ds) \\ &\quad \times P_{lc} \{\bar{Y}_{t-s} = j, A_{t-s} \in db\} \\ &= I(i, j) I(a, b) q(i, a, t) \\ &\quad + \sum_{l \in \bar{E}} \int_{F \times [0, t]} \bar{Q}(i, a; l, dc; ds) \\ &\quad \times f(l, c, t-s) \end{aligned} \quad (50)$$

where

$$\begin{aligned} q(i, a, t) &= P_{ia} \{\bar{T}_1 > t\} \\ &= 1 - \sum_{l \in \bar{E}} \int_F \bar{Q}(i, a; l, db; t) \end{aligned} \quad (51)$$

We have a Markov renewal equation $f = g + \bar{Q} * f$ with

$$g(i, a, t) = I(i, j) I(a, b) q(i, a, t) \quad (52)$$

Following Proposition A.2 in [4], equation (50) has a unique solution

$$\begin{aligned} f(i, a, t) &= P_{ia} \{ \bar{Y}_t = j, A_t \in db \} \\ &= \sum_{l \in \bar{E}} \int_{F \times [0, t]} \bar{R}(i, a; l, dc; ds) g(l, c, t - s) \\ &= \int_{[0, t]} \bar{R}(i, a; j, db; ds) q(j, b, t - s) \end{aligned} \quad (53)$$

where $\bar{R} = \sum_n \bar{Q}^n$ is the Markov renewal kernel corresponding to $\bar{Q}$.

Then, using equation (53)

$$\begin{aligned} P_{ia} \{ W_j(t) = 1 \} &= \int_F P_{ia} \{ \bar{Y}_t = j, A_t \in db \} \\ &= \int_{F \times [0, t]} \bar{R}(i, a; j, db; ds) \\ &\quad \times q(j, b, t - s) \end{aligned} \quad (54)$$

Acknowledgment

This research is supported by the Scientific and Technological Research Council of Turkey through grant 106M044.

References

- [1] Özekici, S. (1996). Complex systems in random environments, In *Reliability and Maintenance of Complex Systems, NATO ASI Series*, S. Ozekici, ed, Springer-Verlag, Berlin, Vol. F154, pp. 137–157.
- [2] Asmussen, S. (2000). *Ruin Probabilities*, World Scientific, Singapore.
- [3] Rolski, T., Schmidli, H., Schmidt, V. & Teugels, J. (1999). *Stochastic Processes for Insurance and Finance*, John Wiley & Sons, Chichester.
- [4] Çinlar, E. & Özekici, S. (1987). Reliability of complex devices in random environments, *Probability in the Engineering and Informational Sciences* **1**, 97–115.
- [5] Gaver, D.P. (1963). Random hazard in reliability problems, *Technometrics* **5**, 211–226.
- [6] Arjas, E. (1981). The failure and hazard process in multivariate reliability systems, *Mathematics of Operations Research* **6**, 551–562.
- [7] Çinlar, E., Shaked, M. & Shanthikumar, J.G. (1989). On lifetimes influenced by a common environment, *Stochastic Processes and their Applications* **33**, 347–359.
- [8] Shaked, M. & Shanthikumar, J.G. (1989). Some replacement policies in a random environment, *Probability in the Engineering and Informational Sciences* **3**, 117–134.
- [9] Lindley, D.V. & Singpurwalla, N.D. (1986). Multivariate distributions for the lifelengths of components of a system sharing a common environment, *Journal of Applied Probability* **23**, 418–431.
- [10] Lefèvre, C. & Malice, M.-P. (1989). On a system of components with joint lifetimes distributed as a mixture of exponential laws, *Journal of Applied Probability* **26**, 202–208.
- [11] Lefèvre, C. & Milhaud, X. (1990). On the association of the lifelengths of components subjected to a stochastic environment, *Advances in Applied Probability* **22**, 961–964.
- [12] Singpurwalla, N.D. & Youngren, M.A. (1993). Multivariate distributions induced by dynamic environments, *Scandinavian Journal of Statistics* **20**, 251–261.
- [13] Esary, J.D. & Ziehms, H. (1975). Reliability analysis of phased missions, *Proceedings of the Conference on Reliability and Fault Tree Analysis*, SIAM, pp. 213–236.
- [14] Kim, K. & Park, K.S. (1994). Phased-mission system reliability under Markov environment, *IEEE Transactions on Reliability* **43**, 301–309.
- [15] Alam, M. & Al-Saggaf, U.M. (1986). Quantitative reliability evaluation of repairable phased-mission systems using Markov approach, *IEEE Transactions on Reliability* **35**, 498–503.
- [16] Mura, I. & Bondavalli, A. (1999). Hierarchical modelling and evaluation of phased-mission systems, *IEEE Transactions on Reliability* **48**, 360–368.
- [17] Çinlar, E. (1969). Markov renewal theory, *Advances in Applied Probability* **1**, 123–187.
- [18] Çinlar, E. (1975). *Introduction to Stochastic Processes*, Prentice Hall, Englewood Cliffs.

Related Articles

Aging and Positive Dependence; General Minimal Repair Models Imperfect Repair; Maintenance Optimization in Random Environments; Markov Processes; Parallel, Series, and Series-Parallel Systems; Renewal Theory; Repairable Systems Reliability.

BORA ÇEKYAY AND SÜLEYMAN ÖZEKICI

Warranty Cost Analysis

Introduction

New products have been appearing in the market at an ever-increasing rate due to the ever-increasing pace of technological innovation. Each generation of products typically incorporates new features and is more complex than the one it replaces. With the increase in complexity and the use of new materials and design methodologies, the reliability of products is of concern to buyers. One way for the manufacturer to assure the buyer that the product will perform satisfactorily is to offer a warranty with the sale of the product. In very simple terms, the warranty assures the buyer that the manufacturer will either repair or replace items that do not perform satisfactorily or refund a fraction or the whole of the sale price. Offering warranty results in additional costs to the manufacturer. The manufacturer must assess this cost very early in the product development stage in order to make decisions with regard to desired product reliability and appropriate product pricing.

This article deals with **warranty cost** analysis for some simple **warranty policies**. This involves building models for claims under warranty, taking into account the reliability of the product. The outline of the article is as follows. The section titled “Product Warranty” looks at the basic concept of a warranty and provides a few examples of commonly used consumer warranties. The cost of warranty is impacted by product reliability in an obvious fashion. Some key reliability results needed in assessing this impact are summarized in the section titled “Product Reliability”. Probabilistic results used in modeling warranty costs are given in the section titled “Warranty Cost Analysis” and applied to analysis of cost per unit sold and **life cycle cost** of warranty in the sections titled “Modeling Warranty Cost per Unit Sold” and “Warranty Cost over Life Cycle” respectively.

Product Warranty

Warranty Concept

A warranty is a contractual agreement between the buyer and the manufacturer (or seller) that is entered

into upon sale of the product or service. The purpose of a warranty is to establish liability of the manufacturer in the event that an item fails or is unable to perform its intended function when properly used. The warranty (often referred to as the *base warranty*) is an integral part of the sale (a service bundled with the product) and is factored into the price. In contrast, extended warranty is a service contract that a buyer might purchase separately to extend the warranty coverage once the base warranty expires.

Warranties are an integral part of nearly all consumer and commercial purchases and also of many government transactions that involve product purchases. In such transactions, warranties serve a somewhat different purpose for the buyer and the manufacturer. For a more detailed discussion of the role of warranty from the perspectives of buyers and manufacturers see [1]. There are many aspects to warranty, and these have been studied by researchers from different disciplines. A discussion of these can be found in [2].

Classification of Warranties

A taxonomy for warranty classification was first proposed by Blischke and Murthy [3]. The first criterion for classification is whether the manufacturer is required to carry out further product development (for example, to improve product reliability) subsequent to the sale of the product. Policies that do not involve such further product development can be further divided into two groups, the first consisting of policies applicable for single-item sales, and the second, policies used in the sale of groups of items (called *lot* or *batch sales*).

Policies in the first group can be subdivided into two subgroups, based on whether the policy is renewing or nonrenewing. For renewing policies, the warranty period begins anew with each replacement, while for nonrenewing policies the replacement (or repaired) item assumes the remaining warranty time of the item it replaced. A further subdivision comes about in that warranties may be classified as “simple” or “combination”. The commonly used consumer warranties featuring free-replacement warranty (FRW) and replacement at *pro rata* cost (PRW), discussed in the next section, are simple policies. A combination policy is a simple policy combined with some additional features or a policy that combines the terms of two or more simple policies. Each of these

2 Warranty Cost Analysis

four groupings can be further subdivided into two subgroups based on whether the policy is one- or two- (or higher)-dimensional. A one-dimensional policy is one that is most often based on either time or age of the item, but could instead be based on usage. An example is a one-year warranty on a television with the option for buyers to purchase extended warranty coverage. In contrast, a two-dimensional policy is based on time or age as well as usage. An example is an automobile with a two-dimensional warranty with a time limit of 3 years and usage limit of 30 000 miles. The warranty expires when one of the limits is exceeded.

Some Simple Warranty Policies

Three simple warranty policies are defined in this section. The cost analysis for each of these is given in the later sections.

Policy 1: One-Dimensional (1-D) Nonrenewing Free-Replacement Warranty (FRW). The manufacturer agrees to repair or provide replacements for failed items free of charge up to a time W from the time of the initial purchase. The warranty expires at time W after purchase.

Typical applications of this warranty are consumer products, ranging from inexpensive items such as photographic film to relatively expensive repairable items such as automobiles, and on expensive nonrepairable items such as microchips and other electronic components.

Policy 2: One-Dimensional Nonrenewing Pro rata Warranty (PRW). The manufacturer agrees to refund a fraction of the purchase price if the item fails before time W from the time of the initial purchase. The buyer is not constrained to buy a replacement item.

Under the *pro rata* warranty (PRW), the refund given depends on the age of the item at failure. In the most common form of PRW, the refund is determined by the linear relationship $q(x) = [(W - x)/W]P$, where W is the warranty period, x is the age at failure, P is the sale price, and $q(\cdot)$, called the *rebate function*, is the amount of refund. Typically, these policies are offered on relatively inexpensive nonrepairable products such as batteries, tires, ceramics, and so on.

Policy 3: Two-Dimensional (2-D) Nonrenewing FRW. The manufacturer agrees to repair or provide a replacement for failed items free of charge up to a time W from the time of initial purchase (limit on time) or up to a usage U (limit on usage), whichever occurs first.

Product Reliability

The modeling of first failures is different from that of subsequent failures. In the latter case, the result depends on whether the failed item is repaired or replaced by a new item.

One-Dimensional Failure Modeling

First Failure. The time to first failure, T , is a random variable that can assume values in the interval $[0, \infty)$, with distribution function

$$F(t) = P\{T \leq t\}, \quad 0 \leq t < \infty \quad (1)$$

The *density function* (if $F(t)$ is differentiable) is given by $f(t) = dF(t)/dt$. The *reliability* (of the item) is given by the survivor function

$$S(t) = P\{T > t\} = 1 - P\{T \leq t\} = 1 - F(t) \quad (2)$$

The failure rate function (or hazard function) $h(t)$ associated with $F(t)$ is defined as

$$h(t) = \frac{f(t)}{1 - F(t)} \quad (3)$$

Subsequent Failures. When a repairable component fails, it can either be repaired or replaced by a new component. In the case of a nonrepairable component, the only option is to replace the failed component with a new item. Since failures occur in an uncertain manner, the number of components needed over a given time interval is a nonnegative, integer-valued random variable. The distribution of this variable depends on the failure distribution of the component, the actions (repair or replace) taken after each failure, and, if repaired, the type of repair. The number of failures $N(t)$ over the interval $[0, t]$, starting with a new component at $t = 0$, can be modeled as a *counting process*.

Nonrepairable Item. In the case of a nonrepairable component, every failure results in the replacement

of the failed component by a new item. If the time to replace a failed component by a new one is small relative to the mean time to failure, then it can be ignored. In this case, the number of failures (replacements) over time is given by a *renewal process* and the mean number of failures over $[0, t]$, is given by the *renewal integral equation*:

$$M(t) = F(t) + \int_0^t M(t-t') dF(t') \quad (4)$$

Repairable Item. There are different types of repairs. With *minimal repair* (see **General Minimal Repair Models**), the failure rate after repair is the same as the failure rate of the item immediately before it failed [4]. With *imperfect repair* (see **Imperfect Repair**), the failure rate after repair is less than that just before failure but not as good as that for a new component.

With minimal repair and the repair time negligible in relation to the mean time between failures, failures over time occur according to a *nonhomogeneous Poisson process* [5]. The mean number of failures over $[0, t]$ is given by

$$E[N(t)] = \Lambda(t) = \int_0^t h(x) dx \quad (5)$$

In the case of imperfect repair, the analysis is more complicated. See [6] and references therein for additional details.

Two-Dimensional Failure Modeling

There are two approaches to modeling failures in the 2-D case.

Approach 1: 1-D Approach. The two-dimensional problem is effectively reduced to a one-dimensional problem by treating usage as a function of age.

First Failure. Let $Z(t)$ denote the usage of the item at age t and assume that it is a linear function of t of the form

$$Z(t) = Rt \quad (6)$$

where $R, 0 \leq R < \infty$, represents the usage rate and is a nonnegative random variable with a distribution function $G(r)$ and a density function $g(r)$.

The hazard function, conditional on $R = r$, is denoted $h(t|r)$. Various forms of $h(t|r)$ have been proposed, including the linear case given by

$$h(t|r) = \theta_0 + \theta_1 r + \theta_2 t + \theta_3 Z(t) \quad (7)$$

The conditional distribution function of time to first failure is given by

$$F(t|r) = 1 - \exp \left\{ - \int_0^t h(u|r) du \right\} \quad (8)$$

On removing the conditioning, we have the distribution function of time to first failure given by

$$F(t) = \int_0^\infty \left(1 - \exp \left\{ - \int_0^t h(u|r) du \right\} \right) g(r) dr \quad (9)$$

Subsequent Failures. Under minimal repair, the failures over time occur according to a nonhomogeneous Poisson process [7] with **intensity function** $\lambda(t|r) = h(t|r)$. Under perfect repair, the failures occur according to the renewal process associated with $F(t|r)$.

The bulk of the literature deals with a linear relationship between usage and age. See, for example, [1, 8, and 9]. For applications, see [10], which deals with motorcycle failures under warranty, and [8] and [11], which deal with automobile warranty data analysis.

Approach 2: 2-D Approach. *First Failure.* Let T and X denote the item age and usage at first failure. Both of these are random variables. T and X are modeled jointly by the bivariate distribution function:

$$F(t, x) = P\{T \leq t, X \leq x\}; \quad t \geq 0, x \geq 0 \quad (10)$$

The survivor function is given by

$$\begin{aligned} \bar{F}(t, x) &= \Pr\{T > t, X > x\} \\ &= \int_t^\infty \int_x^\infty f(u, v) dv du \end{aligned} \quad (11)$$

If $F(t, x)$ is differentiable in both arguments, the bivariate failure density function is given by

$$f(t, x) = \frac{\partial^2 F(t, x)}{\partial t \partial x} \quad (12)$$

Subsequent Failures. In the case of nonrepairable items, failed items are replaced by new ones. In

4 Warranty Cost Analysis

this case, failures occur according to a 2-D renewal function [4] and the expected number of failures over $[0, t) \times [0, x)$ is given by the solution of the two-dimensional integral equation:

$$M(t, x) = F(t, x) + \int_0^t \int_0^x M(t - u, x - v) \times f(u, v) dv du \quad (13)$$

Warranty Cost Analysis

Two types of costs are of interest to the manufacturer: (a) expected cost per unit and (b) expected cost over the product life cycle.

Expected Cost per Unit Sold

Whenever a failed item is returned for rectification action under warranty, the manufacturer incurs various costs – handling, material, labor, facilities, and so on. These costs are random variables. The total cost of servicing all **warranty claims** for an item over the warranty period is the sum of a random number of such individual costs, since the number of claims over the warranty period is also random. The expected warranty service cost per unit is needed in deciding the selling price of the item.

Expected Cost over the Life Cycle of a Product

From a marketing perspective, the product life cycle, L , is the period from the instant a new product is launched to the instant it is withdrawn from the market. Over the product life cycle, product sales (first and repeat purchases) occur over time in a dynamic manner. The manufacturer must service the warranty claims with each such sale. In the case of products sold with one-dimensional nonrenewing warranty, this period is $L + W$, where W is the warranty period.

The expected cost incurred over $[t, t + \delta t)$ depends on the sales over the interval $[t - \psi, t)$, where

$$\psi = \max\{0, t - W\} \quad (14)$$

Sales over time can be modeled either by a continuous function (appropriate for large-volume sales) or a point process (appropriate for small-volume sales).

Assume that the sales rate is given by $s(t)$, $0 \leq t \leq L$, so that total sales S over the life cycle is

$$S = \int_0^L s(t) dt \quad (15)$$

Assumptions

In developing tractable and reasonable warranty cost models, various assumptions regarding the process are made. These typically include the following:

1. All consumers are alike in their usage. This can be relaxed by dividing consumers into two or more groups of similar usage patterns.
2. All items are statistically similar. This can be relaxed by including two types of items (conforming and nonconforming) to take into account quality variations in manufacturing.
3. Whenever a failure occurs, it results in an immediate claim. Relaxing this assumption involves modeling the delay time between failure and claim.
4. All claims are valid. This can be relaxed by assuming that a fraction of the claims is invalid, either because the item was used in a mode not covered by the warranty or because it was a bogus claim.
5. The time to rectify a failed item (either through repair or replacement) is sufficiently small in relation to the mean time between failures so that it can be approximated as being zero.
6. The manufacturer has the logistic support (spares and facilities) needed to carry out rectification actions without delays.
7. The cost of each rectification under warranty (repair or replacement) is modeled by the following cost parameters:

C_m : Cost of replacing a failed unit by a new unit (manufacturing cost per unit)

C_h : Handling cost per claim

C_r : Average cost of each minimal repair

8. The product life cycle, L , is modeled as a deterministic variable. One can easily relax this assumption and treat it as a random variable with a specified distribution.
9. All model parameters (costs and the various distributions involved) are known.

Modeling Warranty Cost per Unit Sold

Cost Analysis of Policy 1

Replace by New. Since failures under warranty are rectified instantaneously through replacement by new items, the number of failures over the period $[0, t)$, $N(t)$, is characterized by a renewal process with time between renewals distributed according to $F(t)$. As a result, the cost $\tilde{C}(W)$ to the manufacturer over the warranty period W is a random variable given by

$$\tilde{C}(W) = [C_m + C_h][N(W)] \quad (16)$$

The expected number of failures over the warranty period is

$$E[N(W)] = M(W) \quad (17)$$

where $M(t)$ is the renewal function given by equation (4). As a result, the expected warranty cost per unit to the manufacturer is

$$C(W) = [C_m + C_h]M(W) \quad (18)$$

Minimal Repair. Failures over the warranty period occur according to a nonhomogeneous Poisson process with intensity function equal to the failure rate $r(t)$ and the expected warranty cost to the manufacturer is

$$C(W) = [C_r + C_h] \int_0^W r(x) dx \quad (19)$$

Cost Analysis of Policy 2

In the case of linear rebate function

$$q(x) = \begin{cases} (1 - x/W)P & 0 \leq x < W \\ 0 & \text{otherwise} \end{cases} \quad (20)$$

the warranty cost to the manufacturer is

$$\tilde{C}(W) = C_h + q(T) \quad (21)$$

where T is the lifetime of the item supplied. As a result, the expected cost per unit to the manufacturer is

$$C(W) = E[\tilde{C}(W)] = C_h F(W) + \int_0^W q(x) f(x) dx \quad (22)$$

Using equation (20) in equation (22) and carrying out the integration, we have

$$C(W) = C_h F(W) + P \left[F(W) - \frac{\mu_W}{W} \right] \quad (23)$$

where

$$\mu_W = \int_0^W t f(t) dt \quad (24)$$

is the *partial expectation* of T .

Cost Analysis of Policy 3

As mentioned in the section titled “Approach 1: 1-D Approach” one can use two different approaches to modeling failures. Approach 1 is more appropriate for repairable products (with minimal repair) and Approach 2, for nonrepairable products.

Approach 1. The product is repairable, and the manufacturer rectifies all failures under warranty through minimal repair. In this case, the distribution of failures over the warranty period can be characterized by an intensity function $\lambda(t|r)$. Conditional on $R = r$, we assume a linear intensity function given by equation (7).

Note that conditional on $R = r$, the warranty expires at time W if the usage rate r is less than γ , where

$$\gamma = \frac{U}{W} \quad (25)$$

and at time t_r given by

$$t_r = \frac{U}{r} \quad (26)$$

if the usage rate $r > \gamma$. The expected numbers of failures over the warranty period for these two cases are given by

$$E[N(W, U)|r] = \int_0^W \lambda(t|r) dt \quad \text{and} \quad (27)$$

$$E[N(W, U)|r] = \int_0^{t_r} \lambda(t|r) dt$$

respectively. On removing the conditioning, we have the expected number of failures during warranty and, using this, we find the expected warranty cost to be

$$C(W, U) = [C_r + C_h] \left\{ \int_0^\gamma \left[\int_0^W \lambda(t|r) dt \right] g(r) dr + \int_\gamma^\infty \left[\int_0^{t_r} \lambda(t|r) dt \right] g(r) dr \right\} \quad (28)$$

Approach 2. The product is nonrepairable and the manufacturer rectifies all failures under warranty through replacement by a new product. In this case, the number of failures over the warranty region is given by a 2-D renewal process and the expected warranty cost is given by

$$C(W, U) = [C_m + C_h]M(W, U) \quad (29)$$

where $M(W, U)$ is obtained from equation (13).

Warranty Cost over Life Cycle

Consider the case where the product is nonrepairable and the life cycle L exceeds the warranty period W . Since the manufacturer must provide a refund or replacements for items that fail before reaching age W , and since the last sale occurs at or before time L , the manufacturer has an obligation to service warranty claims over the interval $[0, L + W]$. Life cycle costs for Policies 1 and 2 will be considered. Some results for Policy 3 can be found in Blischke and Murthy [1, Sections 8.3.1 and 8.4.1].

Cost Analysis for Policy 1

Since failed items are replaced by new ones over the warranty period, the demand for spares in the interval $[t, t + \delta t)$ is due to failure of items sold in the period $[\psi, t)$ where ψ is given by equation (14). It can be shown (details of derivation can be found in Chapter 9 of [1]) that the expected demand rate for spares at time t , $\rho(t)$, is given by

$$\rho(t) = \int_{\psi}^t s(x)m(t-x) dx \quad (30)$$

where $m(t)$ is the renewal density function given by

$$m(t) = \frac{dM(t)}{dt} = f(t) + \int_0^t m(t-t')f(t') dt' \quad (31)$$

The expected total number of spares required to service the warranty over the product life cycle, $ETNS$, is given by

$$ETNS = \int_0^{L+W} \rho(t) dt \quad (32)$$

and the total expected warranty cost over the product life cycle is given by

$$LC(W, L) = [C_m + C_h] \int_0^{W+L} \rho(t) dt \quad (33)$$

Cost Analysis for Policy 2

Under Policy 2 the manufacturer refunds a fraction of the sale price on failure of an item within the warranty period. In order to carry this out, the manufacturer must set aside a fraction of the sale price. This is called *warranty reserving*. We assume a linear rebate function. The rebate over the interval $[t, t + \delta t)$ is due to failure of items that are sold in the interval $[t - \psi, t)$, where ψ is given by equation (14). Let $v(t)$ denote the expected refund rate (i.e. the amount refunded per unit time) at time t . Then it is easily shown that

$$v(t) = P \int_{\psi}^t s(x) \left(\frac{t-t'}{W} \right) f(t-t') dt' \quad (34)$$

for $0 \leq t \leq (W + L)$. (For details, see Chapter 9 of [1].) The expected total reserve (ETR) needed to service the warranty over the product life cycle is given by

$$LC(W, L) = \int_0^{W+L} v(t) dt \quad (35)$$

Concluding Remarks

The cost analysis of many other types of warranty policies can be found in [1] and in more recent papers; references to some of these can be found in [12]. This type of analysis is carried out during the front-end (feasibility) stage of the new product development process. The model selection and parameter values are based on similar products, expert judgment (see **Expert Opinion in Reliability**), and so on. For more on this, see [13], which looks at warranty management from a product life cycle perspective. Once the product is launched, warranty claims provide data for assessing the models and updating the warranty cost estimates. For more on this, relevant references can be found in [12] and [13].

References

- [1] Blischke, W.R. & Murthy, D.N.P. (1994). *Warranty Cost Analysis*, Marcel Dekker, New York.

- [2] Blischke, W.R. & Murthy, D.N.P. (1996). *Product Warranty Handbook*, Marcel Dekker, New York.
- [3] Blischke, W.R. & Murthy, D.N.P. (1992). Product warranty management – I: a taxonomy for warranty policies, *European Journal of Operational Research* **62**, 127–148.
- [4] Blischke, W.R. & Murthy, D.N.P. (2000). *Reliability*, John Wiley & Sons, New York.
- [5] Murthy, D.N.P. (1991). A note on minimal repair, *IEEE Transactions on Reliability* **40**, 245–246.
- [6] Doyen, L. & Gaudoin, O. (2004). Classes of imperfect repair models based on reduction of failure intensity or virtual age, *Reliability Engineering and System Safety* **84**, 45–56.
- [7] Ross, S.M. (1983). *Stochastic Processes*, John Wiley & Sons, New York.
- [8] Lawless, J., Hu, J. & Cao, J. (1995). Methods for estimation of failure distributions and rates from automobile warranty data, *Lifetime Data Analysis* **1**, 227–240.
- [9] Gertsbakh, I.B. & Kordonsky, K.B. (1998). Parallel time scales and two-dimensional manufacturer and individual customer warranties, *IIE Transactions* **30**, 1181–1189.
- [10] Iskandar, B.P. & Blischke, W.R. (2002). Reliability and warranty analysis of a motorcycle based on claims data, in *Case Studies in Reliability and Maintenance*, W.R. Blischke & D.N.P. Murthy, eds, John Wiley & Sons, New York, Chapter 27.
- [11] Kalbfleisch, J.D., Lawless, J.F. & Robinson, J.A. (1991). Methods for the analysis and prediction of warranty claims, *Technometrics* **33**, 273–285.
- [12] Murthy, D.N.P. & Djameludin, I. (2002). Product warranty – a review, *International Journal of Production Economics* **79**, 231–260.
- [13] Murthy, D.N.P. & Blischke, W.R. (2005). *Warranty Management and Product Manufacture*, Springer-Verlag, London.

Related Articles

Life Cycle Costs and Reliability Engineering; Multiattribute Warranty Policies; Warranty Claims and Costs: Statistical Analysis of; Warranty: Usage and Wear Process for; Warranty Modeling; Warranty Cost Prediction Based on Warranty Data; Warranty Servicing.

D. N. PRABHAKAR MURTHY AND
WALLACE R. BLISCHKE

Age-Dependent Minimal Repair and Maintenance

Introduction

It is of great importance to avoid the failure of complex systems during actual operation when such an event is costly and/or dangerous. In such situations, one important area of interest in reliability theory is the study of various **maintenance policies** in order to reduce the operating cost and the risk of a catastrophic breakdown.

Maintenance involves planned and unplanned actions carried out to retain a system in or restore it to an acceptable operating condition. One important research area in reliability is the study of various maintenance policies in order to prevent the occurrence of system failure and improve **system availability**. Maintenance can be classified into two major categories: corrective or preventive. Corrective maintenance (CM) is any maintenance that occurs when the system is failed. Preventive maintenance (PM) is any maintenance that occurs when the system is operating. In this article, we think that maintenance can be further classified according to the degree to which the operating conditions of an item is restored by maintenance in the following two ways (Pham and Wang [1]):

1. *Perfect repair or perfect maintenance* – a maintenance action that restores the system operating condition to as good as new. Generally, replacement of a failed system by a new one is a perfect repair.
2. *Minimal repair or minimal maintenance* – a maintenance action that restores the system to the failure rate it had when it failed; that is, after maintenance the system operating state is often called *as bad as old*.

Optimal maintenance policies aim to provide optimum system reliability/availability and safety performance at lowest possible maintenance costs.

Two policies, of particular note, are the age-replacement maintenance policy and the block-replacement maintenance policy. In the age-replacement maintenance policy (see Barlow and

Hunter [2] for example), an operating system is replaced at age T or at failure, whichever occurs first. This model has been generalized by Glasser [3], Beichelt [4], Tahara and Nishida [5], Muth [6], Cleroux *et al.* [7], Nakagawa and Kowada [8], Bai and Yun [9], Berg *et al.* [10], Block *et al.* [11], Sheu *et al.* [12], and Sheu [13]. In the block-replacement maintenance policy (see Barlow and Proschan [14]), an operating system is replaced by a new one at times kT , $k = 1, 2, 3, \dots$ and at failures. Such a policy is rather wasteful, since sometimes almost-new systems are replaced. To overcome this undesirable feature, various modifications have been advocated to reduce the wastage, see Barlow and Hunter [2], Beichelt [15], Nakagawa [16, 17], Boland and Proschan [18], Boland [19], Nguyen and Murthy [20], Berg and Cleroux [21], Abdel-Hameed [22], Ait Kadi and Cleroux [23], Sheu [13, 24, 25], Ait Kadi *et al.* [26], and Chien and Sheu [27]. Moreover, Makabe and Morimura [28–30] also proposed the new replacement maintenance model in which a system is replaced at the n th failure and discussed the optimum policy. This model has been generalized by Morimura [31], Nakagawa [16, 32], Nakagawa and Kowada [8], and Park [33, 34], Sheu *et al.* [12], Sheu and Griffith [35], Sheu [36], and Chien *et al.* [37].

In this article, we present two PM policies for a complex system with age-dependent minimal repair and general random repair costs. Introducing costs due to minimal repair as well as planned and unplanned replacements, we develop the expected cost per unit time in the long run as a criterion of optimality and seek the optimum replacement policies by minimizing that cost. Under certain conditions, the existence and uniqueness of the minimum-cost policy is shown and discussed. Various special cases are presented. Finally, numerical examples are given.

Mathematical Formulation

Preliminaries

A system has two types of failures when it fails at age z . Type I failure (minor failure) occurs with probability $q(z)$ and is corrected with minimal repair, and where type II failure (catastrophic failure) occurs with probability $p(z) = 1 - q(z)$ the system has to be replaced. Minimal repair means that the repaired system is returned to the same condition as it was,

2 Age-Dependent Minimal Repair and Maintenance

i.e., the failure rate of the repaired system remains the same as it was just prior to failure. Assume that the system has a failure time distribution $F(x)$ with finite mean μ and density $f(x)$. Then, the failure rate (or the hazard rate) is $r(x) = f(x)/\bar{F}(x)$ and the cumulative hazard is $R(x) = \int_0^x r(y) dy$, which has a relation $\bar{F}(x) = \exp\{-R(x)\}$, where $\bar{F}(x) = 1 - F(x)$. It is further assumed that the failure rate $r(x)$ is continuous, monotone increasing, and remains undisturbed by minimal repair.

In this article, two general maintenance replacement policies are separately considered in which minimal repair or replacement takes place according to the following scheme.

Policy 1 ((n, T) policy) A system is replaced at age T or at the n th type I failure or at the first type II failure, whichever occurs first.

Policy 2 ((t, T) policy) A system is replaced at the first type II failure before age t or the first failure after age t or at age T , whichever occurs first.

For operating a system in an infinite time span, let U_i^* denote the length of the i th successive replacement cycle for $i = 1, 2, 3, \dots$. Let V_i^* denote the operational cost over the renewal interval U_i^* . Thus, $\{(U_i^*, V_i^*)\}$ constitutes a renewal reward process. The pairs (U_i^*, V_i^*) , $i = 1, 2, 3, \dots$ are independent and identically distributed. If $D(z)$ denotes the expected cost of the operating system over the time interval $[0, z]$, then it is well known that

$$\lim_{z \rightarrow \infty} \frac{D(z)}{z} = \frac{E(V_1^*)}{E(U_1^*)} \quad (1)$$

(see e.g., Ross [38, p. 52]).

In order to derive the expressions for $E(U_1^*)$ and $E(V_1^*)$, we must describe in more detail the failure process which governs the cost over the interval $[0, U_1^*]$. Consider a nonhomogeneous **Poisson process** $\{N(t), t \geq 0\}$ with intensity $r(t)$ and successive arrival times $S_1, S_2, \dots$. At time S_n a coin is flipped, the outcome is designated by δ_n which takes the value one (head) with probability $p(S_n)$ and the value zero (tail) with probability $q(S_n)$. Let $L(t) = \sum_{n=1}^{N(t)} \delta_n$ and $M(t) = N(t) - L(t)$. Then it can be shown that the process $\{L(t), t \geq 0\}$ and $\{M(t), t \geq 0\}$ are independent nonhomogeneous Poisson process with respective intensities $p(t)r(t)$ and $q(t)r(t)$ (see e.g., Savits [39]). This is similar to the classical decomposition of a Poisson process for constant p . Let Y_1

denote the waiting time until the first type II failure, then $Y_1 = \inf\{t \geq 0 : L(t) = 1\}$. Note that Y_1 is independent of $\{M(t), t \geq 0\}$. Thus, the survival distribution of the time until the first type II failure is given by

$$\begin{aligned} \bar{F}_p(t) &= P(Y_1 > t) = P(L(t) = 0) \\ &= \exp \left\{ - \int_0^t p(x)r(x) dx \right\} \end{aligned} \quad (2)$$

In order to develop the cost model, the following lemma from Sheu [40] is required.

Lemma 1 Let $\{M(t), t \geq 0\}$ be a nonhomogeneous Poisson process with intensity $\lambda(t)$, $t \geq 0$ and $\Lambda(t) = E[M(t)] = \int_0^t \lambda(x) dx$. Denote the successive arrival times by $S_1, S_2, \dots$. Assume that at time S_i ($i = 1, 2, \dots$) a cost of $\phi(C(S_i), c_i(S_i))$ is incurred. Suppose that $C(t)$ at age t is a random variable with finite mean $E[C(t)]$. If $\pi(t)$ is the total cost incurred over $[0, t]$, then

$$\begin{aligned} E[\pi(t)] &= E \left[\sum_{i=1}^{M(t)} \phi(C(S_i), c_i(S_i)) \right] \\ &= \int_0^t \vartheta(x) \lambda(x) dx \end{aligned} \quad (3)$$

where $\vartheta(t) = E_{M(t)}[E_{C(t)}[\phi(C(t), c_{M(t)+1}(t))]]$.

Model Development for Policy 1

For the maintenance policy 1, a system is replaced at age T or at the n th type I failure or at the first type II failure, whichever occurs first. Let c_2 denote the cost of replacement at the n th type I failure or at the first type II failure. Let c_1 denote the cost of replacement at age T . The cost of the minimal repair at age z is $\phi(C(z), c(z))$, where $C(z)$ is the age-dependent random part, $c(z)$ is the age-dependent deterministic part, and ϕ is a positive nondecreasing continuous function. Suppose that the random part $C(z)$ at age z has distribution $W_z(x)$, density function $w_z(x)$, and finite mean $E[C(z)]$. After a replacement the procedure is repeated. We assumed all failures are instantly detected and repaired.

Given that a type II failure does not occur in z , the conditional probability of k type I failures

during $[0, z]$ is

$$p_k(z) = P(M(z) = k | L(z) = 0) = P(M(z) = 0) \frac{e^{-\int_0^z q(x)r(x) dx} \left(\int_0^z q(x)r(x) dx \right)^k}{k!} \quad (4)$$

Let $\hat{Y}_1$ denote the time until the n th type I failure or at the first type II failure for our model with $T = \infty$. The survival distribution of $\hat{Y}_1$ is given by

$$\begin{aligned} \bar{G}(z) &= P(\hat{Y}_1 > z) \\ &= \sum_{k=0}^{n-1} P(L(z) = 0, M(z) = k) \\ &= \sum_{k=0}^{n-1} \bar{F}_p(z) p_k(z) \end{aligned} \quad (5)$$

For our model we have

$$U_1^* = \begin{cases} \hat{Y}_1 & \text{if } \hat{Y}_1 \leq T \\ T & \text{if } \hat{Y}_1 > T \end{cases} \quad (6)$$

Thus, the expected length of a replacement cycle is given by

$$\begin{aligned} E(U_1^*) &= \int_0^T z dG(z) + T \cdot \bar{G}(T) = \int_0^T \bar{G}(z) dz \\ &= \sum_{k=0}^{n-1} \int_0^T \bar{F}_p(z) p_k(z) dz \end{aligned} \quad (7)$$

On the other hand, the V_1^* can be expressed as follows:

$$\begin{aligned} V_1^* &= c_2 \cdot I_{[0, T]}(\hat{Y}_1) + c_1 \cdot I_{(T, \infty)}(\hat{Y}_1) \\ &+ \sum_{i=1}^{n-1} \phi(C(S_i), c(S_i)) \bar{F}_p(S_i) \cdot I_{[0, T]}(S_i) \end{aligned} \quad (8)$$

where S_i is the time of the i th minimal repair and

$$I_{[0, T]}(S_i) = \begin{cases} 1 & \text{if } S_i \in [0, T] \\ 0 & \text{otherwise} \end{cases} \quad (9)$$

Therefore, the expected total cost in a replacement cycle is

$$\begin{aligned} E(V_1^*) &= c_2 \left(1 - \sum_{i=0}^{n-1} \bar{F}_p(T) p_i(T) \right) \\ &+ c_1 \sum_{i=0}^{n-1} \bar{F}_p(T) p_i(T) + \sum_{i=1}^{n-1} \int_0^T \theta(z) \bar{F}_p(z) \\ &\times p_{i-1}(z) q(z) r(z) dz \end{aligned} \quad (10)$$

where $\theta(t) = E_{C(t)}[\phi(C(t), c(t))]$ is expectation with respect to the random variable $C(z)$. For the detailed derivation of equation (10), please refer to Sheu *et al.* [12].

Under policy 1, for the infinite-horizon case, we shall denote the right-hand side of equation (1) by $CR_1(n, T)$, the total expected long-run cost per unit time. And we want to obtain optimal n^* and optimal T^* which minimize $CR_1(n, T)$. Recall that

$$\begin{aligned} CR_1(n, T) &= \\ &= \frac{c_2 \left(1 - \sum_{i=0}^{n-1} \bar{F}_p(T) p_i(T) \right) + c_1 \sum_{i=0}^{n-1} \bar{F}_p(T) p_i(T) \\ &+ \sum_{i=0}^{n-2} \int_0^T \theta(z) \bar{F}_p(z) p_i(z) q(z) r(z) dz}{\sum_{i=0}^{n-1} \int_0^T \bar{F}_p(z) p_i(z) dz} \end{aligned} \quad (11)$$

Model Development for Policy 2

For the maintenance policy 2, if an operating system fails at age $z \leq t$, it is either replaced by a new system (type II failure) with probability $p(z)$ at a cost c_u , or it undergoes minimal repair (type I failure) with probability $q(z) = 1 - p(z)$. Otherwise, an operating system is replaced by a new system at a cost c_r when the first failure after t occurs, or is preventively replaced by a new system at a cost c_p when the total operating time reaches age T ($0 \leq t \leq T$), whichever occurs first. The cost of i th minimal repair at age z is $\phi(C(z), c_i(z))$, where $C(z)$ is the age-dependent random part, $c_i(z)$ is the deterministic part which depends on the age and the number of minimal repairs, and ϕ is a

4 Age-Dependent Minimal Repair and Maintenance

positive, nondecreasing, and continuous function. Suppose that the random part $C(z)$ at age z has distribution $W_z(x)$, density function $w_z(x)$, and finite mean $E[C(z)]$. After a replacement the procedure is repeated. We assumed all failures are instantly detected and repaired.

Let X_t denote the residual life of a system of age $t \geq 0$. Thus, X_0 means the life of a new system. Let $F_t(x)$ be the distribution function of X_t . If at time t a minimal repair is done, then

$$F_t(x) = \frac{F(t+x) - F(t)}{\bar{F}(t)} \quad (12)$$

and the corresponding survival function is

$$\bar{F}_t(x) = \frac{\bar{F}(t+x)}{\bar{F}(t)} \quad (13)$$

For our model we have

$$U_1^* = \begin{cases} Y_1, & \text{if } Y_1 \leq t \\ t + X_t, & \text{if } Y_1 > t, \quad 0 \leq X_t \leq T - t \\ T, & \text{if } Y_1 > t, T - t < X_t \end{cases} \quad (14)$$

and

$$V_1^* = \begin{cases} c_u + \sum_{i=1}^{M(Y_1)} \phi(C(S_i), c_i(S_i)), & \text{if } Y_1 \leq t \\ c_r + \sum_{i=1}^{M(t)} \phi(C(S_i), c_i(S_i)), & \text{if } Y_1 > t, \\ & 0 \leq X_t \leq T - t \\ c_p + \sum_{i=1}^{M(t)} \phi(C(S_i), c_i(S_i)), & \text{if } Y_1 > t, \\ & T - t < X_t \end{cases} \quad (15)$$

By equations (14) and (15), the expressions for $E(U_1^*)$ and $E(V_1^*)$ are given as follows:

$$E(U_1^*) = \int_0^t \bar{F}_p(x) dx + \bar{F}_p(t) \cdot U(t, T) \quad (16)$$

and

$$\begin{aligned} E(V_1^*) = & c_u \cdot F_p(t) + \int_0^t \vartheta(y) \bar{F}_p(y) \\ & \times q(y) r(y) dy + (c_r F_t(T - t) \\ & + c_p \bar{F}_t(T - t)) \bar{F}_p(t) \end{aligned} \quad (17)$$

where

$$\begin{aligned} U(t, T) &= \int_0^{T-t} \bar{F}_t(x) dx \\ &= \frac{\int_0^{T-t} \bar{F}(t+x) dx}{\bar{F}(t)} \end{aligned} \quad (18)$$

and

$$\vartheta(t) = E_{M(t)}[E_{C(t)}[\phi(C(t), c_{M(t)+1}(t))]] \quad (19)$$

For the detailed derivation of equations (16) and (17), please refer to Sheu *et al.* [41].

Under policy 2, for the infinite-horizon case, we shall denote the right-hand side of equation (1) by $CR_2(t, T)$, the total expected long-run cost per unit time, and we want to obtain an optimal pair (t^*, T^*) which minimizes $CR_2(t, T)$. Recall that

$$\begin{aligned} CR_2(t, T) = & \frac{c_u \cdot F_p(t) + \int_0^t \vartheta(y) \bar{F}_p(y) q(y) r(y) dy \\ & + (c_r F_t(T - t) + c_p \bar{F}_t(T - t)) \bar{F}_p(t)}{\int_0^t \bar{F}_p(y) dy + \bar{F}_p(t) U(t, T)} \end{aligned} \quad (20)$$

Optimization

Optimization Results for Policy 1

For Policy 1, we shall attempt to minimize $CR_1(n, T)$ with respect to (n, T) under the following assumptions:

- A1. The failure rate $r(z)$ is continuous, increasing and $r(z) \rightarrow \infty$ as $z \rightarrow \infty$.
- A2. The occurrence probability for type II failure, $p(z)$, is continuous and non-increasing.
- A3. The expected minimal repair cost at age z , $\theta(z)$, is continuous, nonincreasing, differentiable and $\theta(z)r(z) - \theta'(z)$ is continuous and increasing in z , where $\theta'(z) = (d\theta(z)/dz)$.

We see that the inequalities $CR_1(n+1, T) \geq CR_1(n, T)$ and $CR_1(n, T) < CR_1(n-1, T)$ are satisfied if and only if:

$$L(n, T) \geq c_1 \text{ and } L(n-1, T) < c_1 \quad (21)$$

where

$$L(n, T) = \begin{cases} \left\{ \frac{\sum_{i=0}^{n-1} \int_0^T \bar{F}_p(z) p_i(z) dz}{\int_0^T \bar{F}_p(z) p_n(z) dz} \left[\int_0^T (\theta(z)r(z) - \theta'(z)) \bar{F}_p(z) p_n(z) dz \right. \right. \\ \left. \left. + (\theta(T) + c_1 - c_2) \bar{F}_p(T) p_n(T) \right] \right\} \\ - \left\{ \sum_{i=0}^{n-2} \int_0^T \theta(z) \bar{F}_p(z) p_i(z) q(z) r(z) dz + (c_2 - c_1) \left(1 - \sum_{i=0}^{n-1} \bar{F}_p(T) p_i(T) \right) \right\}, \\ n = 1, 2, 3, \dots; \\ 0, \quad n = 0 \end{cases} \quad (22)$$

Differentiating $CR_1(n, T)$ in equation (11) with respect to T , we also see that $(dCR_1(n, T)/dT) = 0$ if and only if

$$Q(T, n) = c_1 \quad (23)$$

where

$$\begin{aligned} Q(T, n) = & r(T) \sum_{i=0}^{n-1} \int_0^T \bar{F}_p(z) p_i(z) dz \\ & \times \left[(\theta(T)q(T) + (c_2 - c_1)p(T)) \right. \\ & \left. - \frac{(\theta(T) + c_1 - c_2)p_{n-1}(T)q(T)}{\sum_{i=0}^{n-1} p_i(T)} \right] \\ & - \left[(c_2 - c_1) \left(1 - \sum_{i=0}^{n-1} \bar{F}_p(T) p_i(T) \right) \right. \\ & \left. + \sum_{i=0}^{n-2} \int_0^T \theta(z) \bar{F}_p(z) p_i(z) q(z) r(z) dz \right] \end{aligned} \quad (24)$$

Then, the main properties of the optimum pair (n^*, T^*) which minimizes $CR_1(n, T)$ are summarized below.

Theorem 1 *Under the assumptions A1, A2, and A3, we have the following results:*

1. If $c_2 \geq \theta(z) + c_1$ for $z \geq 0$, then the system is replaced at age T^* or at the first type II failure, whichever occurs first, and T^* satisfies $Q(T^*, \infty) = c_1$.
2. If $c_1 < c_2 < \theta(z) + c_1$ for $z \geq 0$, then the system is replaced at the n^* th type I failure or at age T^* or at the first type II failure, whichever occurs first, and n^* is a unique solution of $L(n^*, T) \geq c_1$ and $L(n^* - 1, T) < c_1$ and T^* satisfies $Q(T^*, n^*) = c_1$.
3. If $c_2 \leq c_1$, then there exists a finite and unique value of n^* that satisfies equation (21) for all $T > 0$, i.e., the system should be replaced at the n^* th type I failure or at age T or at the first type II failure, whichever occurs first.

For proofs of Theorem 1, see Sheu *et al.* [12, pp. 640–644].

Optimization Results for Policy 2

For Policy 2, we shall attempt to minimize $CR_2(t, T)$ with respect to (t, T) under the following assumptions:

B1. The failure rate $r(z)$ is continuous, increasing and $r(z) \rightarrow \infty$ as $z \rightarrow \infty$.

B2. The occurrence probability for type II failure, $p(z)$, and the expected minimal repair cost at age

$z, \vartheta(z)$, are continuous and nondecreasing functions, and $p(0) = 0$.

B3. $c_u \geq c_r > c_p > 0$, $(c_u - c_r)p(z) + \vartheta(z)q(z) > (c_r - c_p)q(z)$, and $c_r > \vartheta(z) > 0$ for $z \geq 0$.

The main results of the optimum pair (t^*, T^*) which minimizes $CR_2(t, T)$ are summarized below.

Theorem 2 Under the assumptions B1, B2, and B3, we have the following results:

1. If (t^*, T^*) is an optimal pair minimizing $CR_2(t, T)$, then $0 < t^* < T^* < \infty$.
2. The optimal pair (t^*, T^*) is uniquely determined, t^* is the solution of

$$r(T)U(t, T) + \bar{F}_t(T - t) - \frac{(c_u - c_r)p(t) + \vartheta(t)q(t)}{(c_r - c_p)q(t)} = 0 \quad (25)$$

and T^* is given by

$$T^* = r^{-1} \left(\frac{c_u - \frac{(c_u - c_r)\bar{F}_p(t^*)}{q(t^*)} - \vartheta(t^*)\bar{F}_p(t^*) + \int_0^{t^*} \vartheta(y)\bar{F}_p(y)q(y)r(y)dy}{(c_r - c_p) \int_0^{t^*} \bar{F}_p(y)dy} \right) \quad (26)$$

For proofs of Theorem 2, see Sheu *et al.* [41, pp. 342–346].

Remark 1 The assumption $c_u \geq c_r > c_p > 0$ is reasonable because

1. $c_r > c_p$ means that the corrective (unplanned) replacement cost is greater than that of the preventive (planned) replacement.
2. $c_u \geq c_r$ means that the replacement due to a type II failure before age t is more urgent (as well as unexpected) than if the first failure occurred after t .

Special Cases

Several well-known replacement policies are the special cases of policies 1 and 2.

Special Cases for Policy 1

- Case 1–1 ($n = \infty, p(z) = 0$, and $\phi(C(z), c(z)) = c$): this is the classical periodic replacement with minimal repair at failure. Barlow and Hunter [2] considered this case.
- Case 1–2 ($n = \infty, p(z) = 1$, and $\phi(C(z), c(z)) = 0$): this is the classical age replacement considered by Barlow and Hunter [2].
- Case 1–3 ($n = \infty, p(z) = 0$, and $\phi(C(z), c(z)) = c(z)$): this is the case considered by Boland [19].
- Case 1–4 ($n = \infty, p(z) = p, 0 < p < 1$, and $\phi(C(z), c(z)) = C$): this is the case considered by Cleroux *et al.* [7].
- Case 1–5 ($n = \infty$ and $\phi(C(z), c(z)) = C(z)$): this is the case considered by Berg *et al.* [10].
- Case 1–6 ($n = \infty, p(z) = 0$, and $\phi(C(z), c(z)) = c + c(z)$): this is the case considered by Tilquin and Cleroux [42].
- Case 1–7 ($n = \infty$ and $\phi(C(z), c(z)) = c$): this is the case considered by Beichelt [4].
- Case 1–8 ($T = \infty, p(z) = 0$, and $\phi(C(z), c(z)) = c_1$): Makabe and Morimura [28–30] considered this case. Park [34] also considered this case when $F(x)$ is a Weibull failure distribution.
- Case 1–9 ($T = \infty, p(z) = p$, and $\phi(C(z), c(z)) = c_1$): this is the case considered by Nakagawa [16].
- Case 1–10 ($T = \infty, p(z) = p$, and $\phi(C(z), c(z)) = C$): this is the case considered by Park [34].
- Case 1–11 ($p(z) = 0$ and $\phi(C(z), c(z)) = c_1$): this is the case considered by Morimura [31].
- Case 1–12 ($\phi(C(z), c(z)) = C(z) + c(z)$): here the cost of minimal repair is the age-dependent random part $C(z)$ plus the

deterministic part $c(z)$ and so $\theta(z) = E[C(z)] + c(z)$.

Special Cases for Policy 2

- Case 2-1 ($p(z) = 0$ and $\phi(C(z), c_i(z)) = c$): this is the policy considered by Tahara and Nishida [5].
- Case 2-2 ($t = 0$): this is the classical age replacement considered by Barlow and Hunter [2].
- Case 2-3 ($t = T$, $p(z) = 0$, and $\phi(C(z), c_i(z)) = c$): this is the policy 2 considered by Barlow and Hunter [2]. The problem reduces to the classical periodic replacement problem with minimal repair at failure.
- Case 2-4 ($t = T$, $p(z) = 0$, and $\phi(C(z), c_i(z)) = c(z)$): this is the case considered by Boland [19].
- Case 2-5 ($t = T$, $p(z) = 0$, and $\phi(C(z), c_i(z)) = c_i$): Boland and Proschan [18] investigated this case. In particular, they considered the cost structure $c_i = a + ic$.
- Case 2-6 ($t = T$, $p(z) = p$, $0 \leq p \leq 1$, and $\phi(C(z), c_i(z)) = C$): this is the case considered by Cleroux *et al.* [7].
- Case 2-7 ($t = T$ and $\phi(C(z), c_i(z)) = C(z)$): this is the case considered by Berg *et al.* [10].
- Case 2-8 ($t = T$ and $\phi(C(z), c_i(z)) = c_i(z)$): this is the case considered by Block *et al.* [11].
- Case 2-9 ($t = T$): this is the case considered by Sheu [43].
- Case 2-10 ($T = \infty$, $p(z) = 0$, and $\phi(C(z), c_i(z)) = c$): this is the case considered by Muth [6].
- Case 2-11 ($T = \infty$, $p(z) = p$, $q(z) = q = 1 - p$, $c_u = c_r$, and $\phi(C(z), c_i(z)) = C + cq$): this is the case considered by Bai and Yun [9].

Numerical Examples

Barlow and Proschan [14, p. 90] discussed an example of the classical age-replacement problem related to electron tubes used in airline communication and

give the solution when the underlying life distribution is truncated normal. Glasser [3] also give the solution in the cases of Gamma and Weibull distributions. Indeed, the very same case was also used by Cleroux *et al.* [7], Berg *et al.* [10], and Sheu [43] to exemplify the age-replacement policy and the block-replacement policy, respectively. In this example, we have $\mu = 9080$ and $\sigma = 3027$ for the underlying life distribution where the time unit is the hour.

In the numerical analysis here, we shall consider the system with Weibull distribution, which is one of the most common in reliability studies. The pdf of the Weibull distribution with shape parameter β and scale parameter α is given by

$$f(x) = \frac{\beta}{\alpha} \left(\frac{x}{\alpha}\right)^{\beta-1} e^{-\left(\frac{x}{\alpha}\right)^\beta}, \quad x > 0, \beta, \alpha > 0 \quad (27)$$

and the parameters of the distribution will be chosen to be $\beta = 2$ and $\alpha = 1012.2$.

Suppose that $\phi(C(x), c_i(x)) = C(x) + c(x)$. Here, we discuss a model in which, at failure, one replaces the system or repairs it depending on the random variable C . Let c_∞ be the constant. A replacement (type II failure) upon failure at age x takes place if $C > \delta(x)c_\infty$; if $C \leq \delta(x)c_\infty$, then one proceeds with a minimal repair (type I failure). The parameter $\delta(x)$ can be interpreted as a fraction of the constant c_∞ at age x , and $0 \leq \delta(x) \leq 1$. Here, we consider the following parametric form of the repair cost limit function $\delta(x) = \delta e^{-ax}$ with $0 \leq \delta \leq 1$ and $a \geq 0$, which is chosen for ease of computation. Suppose that the random variable C has a **normal distribution** $W(\cdot)$ and density function $w(\cdot)$ with mean μ_C and standard deviation σ_C .

If an operating system fails at age x , it is either replaced with a new system with probability

$$p(x) = 1 - \int_0^{\delta(x)c_\infty} w(z) dz \quad (28)$$

or it undergoes minimal repair with probability

$$q(x) = \int_0^{\delta(x)c_\infty} w(z) dz \quad (29)$$

The random part $C(x)$ of minimal repair cost at age x has density $l_x(z) = \frac{l(z)}{q(z)}$ for $0 \leq z \leq \delta(x)c_\infty$

8 Age-Dependent Minimal Repair and Maintenance

and

$$\begin{aligned}\vartheta(x) &= E[\phi(C(x), c(x))] = E[C(x) + c(x)] \\ &= \int_0^{\delta(x)c_\infty} z \cdot w_x(z) \, dz + c(x) \\ &= \frac{1}{q(x)} \int_0^{\delta(x)c_\infty} z \cdot w(z) \, dz + c(x) \quad (30)\end{aligned}$$

Numerical Results for Policy 1

In our computations, we consider the following three different cost cases:

Case 1a : $c_1 = 600, c_2 = 1200, c_\infty = 1000$,
 $\phi(C(x), c(x)) = 500, C \sim N(300, 60^2)$;

Case 1b : $c_1 = 1000, c_2 = 1200, c_\infty = 1000$,
 $\phi(C(x), c(x)) = C(x) + c(x), c(x) = 0.3x$,
 $C \sim N(300, 60^2)$;

Case 1c : $c_1 = 1200, c_2 = 1000, c_\infty = 1000$,
 $\phi(C(x), c(x)) = C(x) + c(x), c(x) = 0.3x$,
 $C \sim N(300, 60^2)$.

The parameters δ and a were varied to take the different values in order to see their effect on the optimal solution. When a is not equal to zero, the probability of minimal repairing is age dependent. The results are given in Tables 1, 2, and 3.

On the basis of the numerical results, we have the following remarks:

1. Some improvement can be made in the minimum cost per unit time if one allows for minimal repair at failure.

Table 1 Optimal replacement policy under cost case 1a

$q(y)$	δ	a	n^*	T^*	$CR_1(n^*, T^*)$
1.0	1	0	∞	1108.8	1.0822
0.9	0.377	0	∞	1108.8	1.1038
0.8	0.3505	0	∞	1108.6	1.1253
0.7	0.3315	0	∞	1108.4	1.1467
0.6	0.3125	0	∞	1108.1	1.1680
0.5	0.3	0	∞	1107.6	1.1892
0.3	0.2685	0	∞	1106.5	1.2313
0.1	0.223	0	∞	1105.0	1.2727
0.0	0	0	∞	1104.1	1.2932

2. From Table 1, the minimum cost per unit time will be reduced when the probability of minimal repairing is greater.
3. From Tables 2 and 3, the minimum cost per unit time will be reduced when the probability of minimal repairing is age dependent.
4. It can be seen that the present model of policy 1 is a generalization on previously known policies.

Table 2 Optimal replacement policy under cost case 1b

δ	a	n^*	T^*	$CR_1(n^*, T^*)$
1	0	2	1944.1	1.3120
1	0.0004	2	1944.1	1.3120
0.377	0	2	1977.4	1.3059
0.377	0.0004	4	3162.6	1.2379
0.3505	0	2	2015.7	1.3028
0.3505	0.00035	3	3190.4	1.2504
0.3315	0	2	2060.1	1.3012
0.3315	0.0003	3	3205.0	1.2608
0.3	0	2	2176.2	1.3018
0.3	0.0002	3	3161.9	1.2804

Table 3 Optimal replacement policy under cost case 1c

δ	a	$T = 2000$		$T = 4000$		$T = 6000$	
		n^*	$CR_1(n^*, T)$	n^*	$CR_1(n^*, T)$	n^*	$CR_1(n^*, T)$
1	0.0004	1	1.1206	1	1.1148	1	1.1148
1	0.0004	1	1.1206	1	1.1148	1	1.1148
0.377	0	1	1.1206	1	1.1148	1	1.1148
0.377	0.0004	2	1.0664	3	1.0587	3	1.0587
0.3505	0	1	1.1206	1	1.1148	1	1.1148
0.3505	0.0004	3	1.0697	3	1.0625	3	1.0625
0.3315	0	1	1.1206	1	1.1148	1	1.1148
0.3315	0.0004	3	1.0755	3	1.0688	3	1.0688
0.3	0	1	1.1206	1	1.1148	1	1.1148
0.3	0.0003	2	1.0880	3	1.0808	3	1.0808

Numerical Results for Policy 2

In our computations, we consider the following two different cost cases:

Case 2a : $c_u = 1200, c_r = 1200, c_p = 1000,$

$c_\infty = 1100, \phi(C(x), c_i(x)) = C(x),$

$C \sim N(700, 200^2);$

Case 2b : $c_u = 1200, c_r = 1200, c_p = 1000,$

$c_\infty = 1100, \phi(C(x), c_i(x)) = C(x) + 0.1x,$

$C \sim N(700, 200^2).$

The parameters δ and a were varied to take the different values in order to see their influence on the optimal solution. When a is not equal to zero, the probability of minimal repair is age dependent. The results are given in Tables 4 and 5.

On the basis of the numerical results, we have the following remarks:

1. Some improvement can be made in the minimum cost per unit time if one allows for minimal repair at failure.
2. From Tables 4 and 5, the minimum cost per unit time will be reduced when the probability of minimal repairing is age dependent.
3. It can be seen that the present model of policy 2 is a generalization on previously known policies.

Table 4 Optimal replacement policy under cost case 2a

$q(y)$	δ	a	t^*	T^*	$CR_2(t^*, T^*)$
0.977	1	0	554	3322	1.2968
^a	1	0.0005	820	3289	1.2841
0.933	10/11	0	582	3314	1.2938
^a	10/11	0.0005	1285	3266	1.2752
0.841	9/11	0	636	3303	1.2897
^a	9/11	0.0003	880	3277	1.2796
0.691	8/11	0	727	3292	1.2853
^a	8/11	0.0003	1170	3269	1.2764
0.500	7/11	0	869	3289	1.2839
^a	7/11	0.0003	1882	3282	1.2815
0.158	5/11	0	1416	3333	1.3012
0.0023	3/11	0	2822	3402	1.3284

^a#: age dependent

Table 5 Optimal replacement policy under cost case 2b

$q(y)$	δ	a	t^*	T^*	$CR_2(t^*, T^*)$
0.977	1	0	479	3343	1.3050
^a	1	0.0008	1072	3299	1.2879
0.933	10/11	0	501	3337	1.3029
^a	10/11	0.0008	1439	3292	1.2853
0.841	9/11	0	542	3329	1.2997
^a	9/11	0.0008	1714	3300	1.2883
0.691	8/11	0	610	3322	1.2968
^a	8/11	0.0005	1071	3306	1.2908
0.500	7/11	0	709	3321	1.2965
^a	7/11	0.0003	968	3319	1.2958
0.158	5/11	0	1026	3357	1.3106
0.0023	3/11	0	1531	3408	1.3305

^a#: age dependent

Conclusions

In this article, two extensions of age-replacement policies were presented and analyzed for a complex system that has two types of failures. Type I (minor) failure is removed by a minimal repair, whereas type II (catastrophic) failure is removed only by a replacement. The failure types are age dependent, and the repair costs are also generalized. Under policy 1 (i.e., the (n, T) policy), a system is replaced at the n th type I failure, or the first type II failure, or at age T , whichever occurs first. Under policy 2 (i.e., the (t, T) policy), if an operating system fails before age t , it is either replaced by a new system (type II failure), or it undergoes minimal repair (type I failure); otherwise, the system is correctively replaced by a new system when the first failure after t arrives, or preventively replaced by a new system when the total operating time reaches age T ($0 \leq t \leq T$), whichever occurs first. For operating the system in the long run, the expected cost per unit time was developed for the two maintenance policies, and the optimal pairs (n^*, T^*) and (t^*, T^*) that minimize the corresponding cost rate were determined, respectively. We have shown that under some mild conditions, the optimal solutions exist and they are unique. Various special cases were included and numerical examples were presented. Our results have extended many of the well-known results for the age-replacement policies.

However, we must point out the limitations of our models: Firstly, we have assumed that repairs are minimal. That means the repaired system is "as bad as old". In some practical situations, the repairs

are not minimal. Therefore, it is more reasonable to consider so-called imperfect repairs which bring the repaired system to a “better than the before-repair” state. Secondly, it is assumed that the repairs and replacements occur instantaneously. In fact, these actions may take time. Taking these realistic factors into consideration in the two maintenance policies can be one direction for future research. From the analysis of this study, we feel that our approach may not be feasible if more realistic factors are included in the model. Therefore, we need to adopt other modeling techniques to solve more general optimal replacement problems such as stochastic dynamic programming and simulation.

References

- [1] Pham, H. & Wang, H. (1996). Invited review: imperfect maintenance, *European Journal of Operational Research* **94**, 425–438.
- [2] Barlow, R.E. & Hunter, L.C. (1960). Optimum preventive maintenance policies, *Operations Research* **8**, 90–100.
- [3] Glasser, G.J. (1967). The age replacement problem, *Technometrics* **9**, 83–91.
- [4] Beichelt, F. (1976). A general preventive maintenance policy, *Mathematische Operationsforschung und Statistik* **7**, 927–932.
- [5] Tahara, A. & Nishida, T. (1975). Optimal replacement policy for minimal repair model, *Journal of the Operations Research Society of Japan* **18**, 113–124.
- [6] Muth, E.J. (1977). An optimal decision rule for repairs vs replacement, *IEEE Transactions on Reliability* **26**, 179–181.
- [7] Cleroux, R., Dubuc, S. & Tilquin, C. (1979). The age replacement problem with minimal repair and random repair cost, *Operations Research* **27**, 1158–1167.
- [8] Nakagawa, T. & Kowada, M. (1983). Analysis of a system with minimal repair and its applications to replacement policy, *European Journal of Operational Research* **12**, 176–182.
- [9] Bai, D.S. & Yun, W.Y. (1986). An age-replacement policy with minimal repair cost limit, *IEEE Transactions on Reliability* **35**, 452–454.
- [10] Berg, M., Bienvenu, M. & Cleroux, R. (1986). Age-replacement policy with age-dependent minimal repair, *INFOR* **24**, 26–32.
- [11] Block, H.W., Borges, W.S. & Savits, T.H. (1985). Age-dependent minimal repair, *Journal of Applied Probability* **22**, 370–385.
- [12] Sheu, S.H., Griffith, W.S. & Nakagawa, T. (1995). Extended optimal replacement model with random minimal repair costs, *European Journal of Operational Research* **85**, 636–649.
- [13] Sheu, S.H. (1998). A generalized age and block replacement of a system subject to shocks, *European Journal of Operational Research* **108**, 345–362.
- [14] Barlow, R.E. & Proschan, F. (1965). *Mathematical Theory of Reliability*, John Wiley & Sons, New York.
- [15] Beichelt, F. (1981). A general block-replacement policy, *IEEE Transactions on Reliability* **2**, 171–172.
- [16] Nakagawa, T. (1981). Generalized models for determining optimal number of minimal repairs before replacement, *Journal of the Operations Research Society of Japan* **24**, 325–337.
- [17] Nakagawa, T. (1982). A modified block replacement with two variables, *IEEE Transactions on Reliability* **31**, 398–400.
- [18] Boland, P.J. & Proschan, F. (1982). Periodic replacement with increasing minimal repair costs at failure, *Operations Research* **30**, 1183–1189.
- [19] Boland, P.J. (1982). Periodic replacement when minimal repair costs vary with time, *Naval Research Logistics Quarterly* **29**, 541–546.
- [20] Nguyen, D.G. & Murthy, D.N.P. (1984). A combined block and repair limit replacement policy, *Journal of the Operational Research Society* **35**, 653–658.
- [21] Berg, M. & Cleroux, R. (1982). The block replacement problem with minimal repair and random repair costs, *Journal of Statistical Computation and Simulation* **15**, 1–7.
- [22] Abdel-Hameed, M.S. (1987). An imperfect maintenance model with block replacement, *Applied Stochastic Models and Data Analysis* **3**, 63–72.
- [23] Ait Kadi, D. & Cleroux, R. (1988). Optimal block replacement policies with multiple choice at failure, *Naval Research Logistics* **35**, 99–110.
- [24] Sheu, S.H. (1992). Optimal block replacement policies with multiple choice at failure, *Journal of Applied Probability* **29**, 129–141.
- [25] Sheu, S.H. (1996). A modified block replacement policy with two variable and general random minimal repair cost, *Journal of Applied Probability* **33**, 557–572.
- [26] Ait Kadi, D., Beaucaire, C. & Cleroux, R. (1990). A periodic maintenance model with used equipment and random minimal repair, *Naval Research Logistics* **37**, 855–865.
- [27] Chien, Y.H. & Sheu, S.H. (2006). Extended optimal age-replacement policy with minimal repair of a system subject to shocks, *European Journal of Operational Research* **174**, 169–181.
- [28] Makabe, H. & Morimura, H. (1963a). On some preventive maintenance, *Journal of the Operations Research Society of Japan* **5**, 110–124.
- [29] Makabe, H. & Morimura, H. (1963b). On some preventive maintenance Policies, *Journal of the Operations Research Society of Japan* **6**, 17–47.
- [30] Makabe, H. & Morimura, H. (1965). Some considerations on preventive maintenance policies with numerical analysis, *Journal of the Operations Research Society of Japan* **7**, 154–171.

- [31] Morimura, H. (1970). On some preventive maintenance policies for IFR, *Journal of the Operations Research Society of Japan* **12**, 94–124.
- [32] Nakagawa, T. (1983). Optimal number of failures before replacement time, *IEEE Transactions on Reliability* **32**, 115–116.
- [33] Park, K.S. (1979). Optimal number of minimal repairs before replacement, *IEEE Transactions on Reliability* **28**, 137–140.
- [34] Park, K.S. (1987). Optimal number of minimal repairs before replacement, *International Journal of Systems Science* **18**, 333–337.
- [35] Sheu, S.H. & Griffith, W.S. (1996). Optimal number of minimal repairs before replacement of a system subject to shocks, *Naval Research Logistics* **43**, 319–333.
- [36] Sheu, S.H. (1999). Extended optimal replacement model for deteriorating systems, *European Journal of Operational Research* **112**, 503–516.
- [37] Chien, Y.H., Sheu, S.H., Zhang, Z.G. & Love, E. (2006). An extended optimal replacement model of systems subject to shocks, *European Journal of Operational Research* **175**, 399–412.
- [38] Ross, S.M. (1970). *Applied Probability Models with Optimization Applications*, Holden-Day, San Francisco.
- [39] Savits, T.H. (1988). Some multivariate distributions derived from a non-fatal shock model, *Journal of Applied Probability* **25**, 383–390.
- [40] Sheu, S.H. (1991a). A general replacement model with minimal repair and general random repair costs, *Microelectronics Reliability* **31**, 1009–1017.
- [41] Sheu, S.H., Kuo, C.M. & Nakagawa, T. (1993). Extended optimal age-replacement policy with minimal repair, *Recherche operationnelle/Operations Research* **27**(3), 337–351.
- [42] Tilquin, C. & Cleroux, R. (1975). Periodic replacement with minimal repair at failure and general cost function, *Journal of Statistical Computation and Simulation* **4**, 63–67.
- [43] Sheu, S.H. (1991b). A generalized block replacement policy with minimal repair and general random repair costs for a multi-unit system, *Journal of the Operational Research Society* **42**, 331–341.

Further Reading

Sheu, S.H. & Griffith, W.S. (1991). Multivariate age-dependent imperfect repair, *Naval Research Logistics* **38**, 839–850.

Related Articles

Block Replacement; General Minimal Repair Models; Group Maintenance Policies; Imperfect Repair; Maintenance Optimization; Maintenance Optimization in Random Environments; Multivariate Age and Multivariate Renewal Replacement; Multivariate Imperfect Repair Models; Replacement Strategies.

SHEY-HUEI SHEU AND YU-HUNG CHIEN

Burn-In and Maintenance Policies

Introduction

In this article we will use product, component, and system alternatively unless stated otherwise.

Burn-in is a widely used method in engineering. It can eliminate weak products from a population before they are released to customers. Thus it can greatly improve reliability performance of products. This method is particularly useful when the failure rate function of the products has a bathtub shape. By bathtub shape it is meant that the failure rate function strictly decreases at first until certain time point, say t_1 , then keeps as a constant until another time point, say t_2 , and it becomes strictly increasing afterward. The two time points t_1 and t_2 are called the *change points* of the failure rate function. The class of lifetime distributions with bathtub-shaped failure rate functions include many other lifetime classes such IFR (increasing failure rate), DFR (decreasing failure rate), and exponential distributions (constant failure rate) as special cases.

To perform a burn-in procedure, the products are operated for a certain time, say b , usually subject to more severe conditions encountered in field application. For more explanation about burn-in we refer to [1–3]. Insufficient burn-in time will not be able to improve the quality of a product significantly, and on the other hand excessive burn-in will damage the products' quality and will be very costly. A major concern of burn-in is to decide the optimal burn-in time, denoted as b^* , so that either the desired reliability measure of the burned-in products is optimized or the associated cost function is minimized.

Another important method applied in engineering is maintenance. Maintenance can raise the reliability of products, prolong mean residual life of products, reduce number of failure of products in field application, and increase availability of products. The most common maintenance actions include the following: (a) performing preventive repair/replacement; (b) allocating redundancies; and (c) conducting inspection (see [4]).

There are two basic repair models. One is complete repair, which repairs the failed product as good

as new. Sometimes the term *replacement* is an alternative to complete repair unless otherwise specified. The other is the so-called **minimal repair**, which repairs the failed product in such a way that the failure rate of the repaired product is exactly the same as the rate prior to the failure of the product.

Suppose that a system consists of several components, say $C_1, \dots, C_n$. One way to “bolster” this system is to provide spare components as redundancies of some of the components $C_i, 1 \leq i \leq n$. A major concern is which spare component should be used as **redundancy** of which $C_i, 1 \leq i \leq n$. In many cases the purpose of providing spare components is to achieve higher reliability subject to one or more constraints on expense, weight, volume, and so on. There are two different types of redundancy. One is standby (or cold) redundancy, which means that the redundant component will be operated at the failure of the original component. The other is active (or hot, or warm) redundancy, which means that both the redundant component and the original component are operated simultaneously and is thus equivalent to forming a parallel subsystem using the redundancy and the original component.

Inspection is necessary for detecting failure of a product if its failure is not self-announcing. Sometimes an age policy could be specified, which requires the product to be replaced at a specific age regardless of the status of the product at that time. It is certainly too costly to conduct too frequent inspections. However, on the contrary, too few inspections will significantly affect the availability of the product since the failed product will not be repaired for a long time. A trade-off between the frequency of inspection and the **system availability** must be made.

Optimal Burn-In Time

Let $F(x)$ be the distribution function of the lifetime X of a product. Assuming that the product exhibits a bathtub-shaped failure rate function, it has been shown that for any of the following reliability measures and cost functions the optimal burn-in time b^* never exceeds the first change point t_1 of the failure rate function.

Reliability Measures

Conditional Survival Probability. A product surviving the burn-in procedure with time b has a conditional survival function $\bar{F}_b(x) = \bar{F}(b+x)/\bar{F}(b)$,

where $\bar{F}(t) = 1 - F(t)$ is the survival function of X . Let $\tau > 0$ be a given mission time. Thus, the probability that this product operates for τ time units without failure is $\bar{F}_b(\tau)$. Often we may want to find out the optimal time b^* maximizing $\bar{F}_b(\tau)$ (see [5]).

Mean Number of Failures. Suppose that in field operation we use products surviving the same burn-in time b and that at failure, the failed product will be replaced by an independent and identically distributed (i.i.d.) one. Let $\tau > 0$ be a given mission time, and $N_b(\tau)$ be the number of failures in the interval $[0, \tau]$. Then the mean of $N_b(\tau)$ is given as $E(N_b(\tau)) = \sum_{i=1}^{\infty} F_b^{(i)}(\tau)$, where $F_b^{(i)}(x)$ is the i th convolution of $F_b(x)$. In practice we often want to find the optimal burn-in time b^* minimizing $E(N_b(\tau))$ (see [5]).

Mean Residual Life. After surviving the burn-in time b , the mean residual life of the product is given by $\mu(b) = E(X - b | X > b) = \int_b^{\infty} \bar{F}(t) / \bar{F}(b) dt$. It is natural that we want to find the optimal burn-in time b^* to maximize $\mu(b)$ (see [6]).

Cost Functions

During the burn-in procedure a product may fail before time b . Assume that the failed product will be replaced by a product with the same distribution and then the procedure will be continued. Let each replacement cost c_s . Further we assume that the cost of the burn-in per unit time is c_0 . It is shown in [7] that the mean cost incurred by this burn-in procedure is given by $k(b) = c_0 \int_0^b \bar{F}(t) / \bar{F}(b) dt + c_s F(b) / \bar{F}(b)$. Suppose that a product surviving the burn-in procedure has been put into field operation. If it fails before mission time τ , then a cost C will be incurred; and if it does not fail during the mission, then a gain G is obtained. Combining all the expenses and gain, it follows that the mean total cost function is given by $c_1(b) = k(b) + C(F(b + \tau) - F(b)) / \bar{F}(b) - G\bar{F}(b + \tau) / \bar{F}(b)$.

Suppose that the operation of a product surviving the burn-in procedure will yield a gain proportional to its mean life. Combining the burn-in cost and the gain in field operation, we have another total mean cost function expressed as $c_2(b) = k(b) - K \int_b^{\infty} \bar{F}(t) / \bar{F}(b) dt$, where K is the proportion constant.

Maintenance Models

Age-replacement policy is a common replacement model. Let $X_i, i \geq 1$ be the lifetimes of a sequence of products and $\tau > 0$ is a specified age. Suppose that at time zero a product is operated. This product will be replaced (or repaired) by another one at its failure or at age τ , whichever comes first. That is, $T_i = \min\{X_i, \tau\}$ expresses the operating time of the i th product. Usually the lifetimes $X_i, i \geq 1$ are i.i.d. Block *et al.* [8] extended this to the following general model by considering both complete repair and minimal repair. An operating product is completely replaced whenever it reached age τ . If it fails at age $x < \tau$, then it is either replaced by an i.i.d. one with probability $p(x)$ or undergoes minimal repair with probability $q(x) = 1 - p(x)$.

Block-replacement policy is another commonly used replacement model. Under this policy the product is replaced (or repaired) by another one at its failure or at scheduled (or planned) times $\tau, 2\tau, \dots$, where $\tau > 0$ is a predetermined number.

Cost Functions Resulting from Maintenance

In either the age-replacement policy or the block-replacement policy, both unscheduled (or unplanned) replacement at failures and scheduled (planned) replacement at age τ or at times $k\tau, k \geq 1$ are performed. In the literature there are thousands of papers that discussed various cost functions $C(\tau)$ associated with different combinations of preventive replacement policies and repair models. For more information we refer to [9, 10]. The optimal replacement policy, denoted as τ^* , is desired for minimizing the function $C(\tau)$. If the same τ is used for both age- and block-replacement policies and let $A(\tau)$ and $B(\tau)$ be the corresponding expected long-run cost per unit time respectively, then reference [11] gives an interesting result showing the connection between $A(\tau)$ and $B(\tau)$ by $B(\tau) = \int_{[0, \tau)} A(\tau - x) dU(x)$ where $U(x)$ is an appropriate renewal function. This equation means that $A(\tau)$ and $B(\tau)$ can be expressed in terms of each other.

Cost Resulting from Both Burn-In and Age Replacement

Suppose that an age-replacement policy τ is applied and i.i.d. products surviving burn-in time b will be

used for replacement. Combining the cost with both burn-in and age-replacement policy, we obtain the long-run average cost function $c(b, \tau) = c_1 F_b(\tau) + c_2 \bar{F}_b(\tau) + k(b) / \int_0^\tau \bar{F}_b(t) dt$, where $k(b)$ is the burn-in cost defined before, cost c_1 is suffered for each unscheduled replacement, and c_2 is a cost suffered from each scheduled replacement.

Cost Resulting from Both Burn-In and Block Replacement

We assume that at each failure the products will undergo a minimal repair and at each scheduled replacement a product surviving the burn-in time b will be used. Let c_m be the cost of each minimal repair, and c_r be the cost of a replacement. Combining the costs with both burn-in and age-replacement policy, we obtain the long-run average cost per unit time as a function of b and τ

$$c(b, \tau) = \frac{1}{\tau} \left(-c_s + \frac{c_s + c_0 \int_0^b \bar{F}(t) dt}{\bar{F}(b)} \right) + \frac{c_r + c_m \int_b^{b+\tau} r(t) dt}{\tau} \quad (1)$$

The problem of simultaneously searching the optimal burn-in time b^* and the optimal age/block-replacement policy τ^* minimizing the function $c(b, \tau)$ is considered in [12]. It is shown that $0 \leq b^* \leq t_1$ and $t_2 < b^* + \tau^* < \infty$ if the failure rate function of the products has a bathtub shape with change points $t_1 \leq t_2$.

System Availability without Inspection

Suppose that at time zero a system with lifetime U_1 is operated and will be repaired at its failure with time D_1 . Upon the completion of the repair the system is operated again with lifetime U_2 and will be repaired at its failure with time D_2 . The process will be continued in the same manner. At the failure of the system either complete or minimal repair is to be performed. The time period from the completion of one complete repair to the completion of the following complete repair is called a *cycle*. The point (or instantaneous) system availability at time t , denoted as $A(t)$, is defined as the probability

that the system is functioning at time t . The limit of $A(t)$ as $t \rightarrow \infty$, denoted as A , is defined as the stable (or steady) system availability if this limit exists. The average system availability on interval $[0, t]$, denoted as $\bar{A}(t)$, is the ratio of the total functioning time in $[0, t]$ to the length t of the interval $[0, t]$. The limit of $\bar{A}(t)$ as t goes to infinity, denoted as $\bar{A}$, is named the *long-run stable system availability* if the limit exists.

Mi [13, 14] generalized the above models and allows nonidentical (U_i, D_i) , $i \geq 1$. In particular, Mi [13] proved that under certain dominance conditions the steady system availability A exists and derived its expression.

Cha and Kim [15] and Iyer [16] considered the following model. Let $\{U_i, i \geq 1\}$ be the i.i.d. lifetimes of the systems. U_i has cumulative distribution function (CDF) $F(t)$, failure rate function $r(t)$, and mean μ . A new system starts working at time zero. It is assumed that in each cycle the probability of performing complete repair is given by $p(t)$, a predetermined function satisfying $0 < p(t) \leq 1$. Here, t is the age of the system in each cycle, i.e., prior to the start of the complete repair in each cycle the elapsed time from the initial instant of the cycle minus any accumulated minimal repair times within this cycle. Let $\{X_i, i \geq 1\}$ and $\{Y_i, i \geq 1\}$ be the i.i.d. minimal repair times and the i.i.d. complete repair times, respectively. Assume that $\{X_i, i \geq 1\}$ and $\{Y_i, i \geq 1\}$ are independent of each other as well as of the failure process. Denote the means of X_i and Y_i as v_m and v_c . It has been shown that under the above assumptions the long-run system availability is given as

$$\bar{A} = \frac{\int_0^\infty \bar{F}_p(t) dt}{\int_0^\infty \bar{F}_p(t) dt + v_m \int_0^\infty q(t) r(t) \bar{F}_p(t) dt + v_c} \quad (2)$$

where $q(t) = 1 - p(t)$ is the probability of performing minimal repair when the system has age t and $\bar{F}_p(t) = \exp\{\int_0^t p(u) r(u) du\}$.

System Availability with Inspection

Consider the case when the failure of a system can only be detected by inspection. The failed system found at inspection will be replaced by an i.i.d. one immediately and the times for implementing replacement are i.i.d. random variables Y_i , $i \geq 1$. The

4 Burn-In and Maintenance Policies

time needed for inspection is assumed negligible. The system will operate without any intervention if it is found to be functioning at an inspection. The inspection is characterized by a constant $\tau > 0$, and there is an age policy $\eta > 0$, where η is an integer or $\eta = \infty$, such that at time $\eta\tau$ the system is replaced regardless of the status of the system at that time.

Mi [17] showed that if $Y_i = 0, i \geq 1$, i.e. the time for implementing replacement is also negligible, then the long-run system availability is given by

$$\bar{A} = \frac{\int_0^{\eta\tau} \bar{F}(t) dt}{\tau \sum_{i=0}^{\eta-1} \bar{F}(i\tau)} \quad (3)$$

where F is the CDF of the lifetime of the system. Cui and Xie [18] assumed that $\eta = \infty$, Y_i has CDF G and showed that the long-run system availability is given by

$$\bar{A} = \frac{\int_0^{\infty} \bar{F}(t) dt}{\tau \sum_{i=0}^{\infty} \bar{F}(i\tau) + \int_0^{\infty} \bar{G}(t) dt} \quad (4)$$

Optimal Allocation of Components in Parallel-Series System

Let $P_i, 1 \leq i \leq k$ be the disjoint min path sets of a parallel-series system with sizes $n_1 \leq \dots \leq n_k$. Suppose that there are $n = \sum_{i=1}^k n_i$ independent components with reliabilities $p_i, 1 \leq i \leq k$. It is shown in [19] that the reliability of the system is maximized when the n_1 most reliable components are allotted to P_1 , the n_2 next most reliable components are allotted to P_2 , ..., and the n_k least reliable components are allotted to P_k .

Optimal Allocation of One Active Redundancy to k-out-of-n System

Suppose that $C, C_i, 1 \leq i \leq n$ are $n+1$ components with independent lifetimes $T, T_1 \leq^{\text{st}} \dots \leq^{\text{st}} T_n$. Let the lifetime of a k -out-of- n system composed of components $C_i, 1 \leq i \leq n$ be $\tau_k = \tau(T_1, \dots, T_i, T_{i+1}, \dots, T_n)$, and let $\tau_k^{(i)} = \tau(T_1, \dots, T_i \vee T,$

$T_{i+1}, \dots, T_n)$, where $T_i \vee T$ means that the component C is an active redundancy with component C_i .

It is shown in [20] that $\tau_k^{(1)} \leq^{\text{st}} \dots \leq^{\text{st}} \tau_k^{(n)}$. Mi [21] derived sufficient conditions under which one can find the optimal way of bolstering one component so that the reliability of the resulting k -out-of- n system is maximized.

Optimal Allocation of One Redundancy to Coherent System

Let $C, C_i, 1 \leq i \leq n$ be $n+1$ components with independent lifetimes $T, T_1 \leq^{\text{st}} \dots \leq^{\text{st}} T_n$. Denote the lifetime of coherent system (S, ϕ) consisting of components $C_i, 1 \leq i \leq n$ by $\tau_k = \tau(\mathbf{T})$ and let $\tau^{(i)} = \tau(T_i \vee T, \mathbf{T}^{(i)})$ where $\mathbf{T} = (T_1, \dots, T_n)$ and $\mathbf{T}^{(i)}$ is obtained by removing component T_i from $\mathbf{T}$. Using the concepts of “more critical” and “permutation equivalent” introduced in [22] the following results are shown in [23]: (a) If $i \stackrel{c}{=} j$ and $T_i \leq^{\text{st}} T_j$, then $\tau^{(j)} \leq^{\text{st}} \tau^{(i)}$; (b) if $i \stackrel{c}{>} j$ and $T_i \leq^{\text{st}} T_j$, then $\tau^{(j)} \leq^{\text{st}} \tau^{(i)}$; and (c) under the same conditions as in (b) it holds that $\tau^{(j)}(T_j + T, \mathbf{T}^{(j)}) \leq^{\text{st}} \tau^{(i)}(T_i + T, \mathbf{T}^{(i)})$.

Optimal Allocation of More Than One Redundancies to k-out-of-n System

Suppose that there are $n+r$ ($1 \leq r \leq n$) components available, of which r will be used for active redundancy. From the given $n+r$ components, r components are chosen to be used as active redundancies, and another r components receive active redundancies. Consider a k -out-of- n system consisting of the r bolstered and the other $n-r$ original components. Mi [24] studied the problem of which r components should be used for active redundancy, and where to allocate them to maximize the lifetime of the resulting k -out-of- n system. It is shown there that under the usual stochastic ordering $\leq^{\text{st}}$ the first r weakest components should be used for active redundancy and allocated in reverse order to the next r weakest components.

References

- [1] Jensen, F. & Petersen, N. (1982). *Burn-in*, John Wiley & Sons, New York.

- [2] Kuo, W. (1984). Reliability enhancement through optimal burn-in, *IEEE Transactions on Reliability* **33**, 145–156.
- [3] Block, H. & Savits, T. (1997). Burn-in, *Statistical Science* **12**, 1–13.
- [4] Aven, T. & Jensen, U. (1999). *Stochastic Models in Reliability*, Springer.
- [5] Mi, J. (1994). Maximization of a survival probability and its application, *Journal of Applied Probability* **31**, 1026–1033.
- [6] Mi, J. (1995). Bathtub failure rate and upside-down bathtub mean residual life, *IEEE Transactions on Reliability* **44**, 388–391.
- [7] Mi, J. (1996). Minimizing some cost functions related to both burn-in and field use, *Operations Research* **44**, 497–500.
- [8] Block, H., Borges, W. & Savits, T. (1988). A general age replacement model with minimal repair, *Naval Research Logistics* **35**, 365–372.
- [9] Barlow, R. & Proschan, F. (1965). *Mathematical Theory of Reliability*, John Wiley & Sons.
- [10] Beichelt, F. (1993). A unifying treatment of replacement policies with minimal repair, *Naval Research Logistics* **40**, 51–67.
- [11] Savits, T. (1988). A cost relationship between age and block replacement policies, *Journal of Applied Probability* **25**, 789–796.
- [12] Mi, J. (1994). Burn-in and maintenance policies, *Advances in Applied Probability* **26**, 207–221.
- [13] Mi, J. (1995). Limiting behavior of some measures of system availability, *Journal of Applied Probability* **32**, 482–493.
- [14] Mi, J. (2006). Limiting availability of system with non-identical lifetime distribution and non-identical repair time distributions, *Statistics and Probability Letters* **76**, 729–736.
- [15] Cha, J.H. & Kim, J.J. (2002). On the existence of the steady state availability of imperfect repair model, *Sankhyā* **64**(Series B), 76–81.
- [16] Iyer, S. (1992). Availability results for imperfect repair, *Sankhyā* **54**(Series B), 249–256.
- [17] Mi J. (2001). Inspection-age-replacement policy and system availability, in *System & Bayesian Reliability*, Hayakawa Y., Irony T. & Xie M., eds, World Scientific, pp. 73–94.
- [18] Cui, L. & Xie, M. (2005). Availability of a periodically inspected system with random repair or replacement times, *Journal of Statistical Planning and Inference* **131**, 89–100.
- [19] El-Newehi, E., Proschan, F. & Sethuraman, J. (1986). Optimal allocation of components in parallel-series and series-parallel systems, *Journal of Applied Probability* **23**, 770–777.
- [20] Boland, P., El-Newehi, E. & Proschan, F. (1992). Stochastic order for redundancy allocations in series and parallel systems, *Advances in Applied Probability* **24**, 161–171.
- [21] Mi, J. (2003). A unified way of comparing the reliability of coherent systems, *IEEE Transactions on Reliability* **52**, 38–43.
- [22] Boland, P., El-Newehi, E. & Proschan, F. (1988). Active redundancy allocation in coherent systems, *Probability in Engineering and Information Science* **2**, 343–353.
- [23] Meng, F. (1996). More on optimal allocation of components in coherent systems, *Journal of Applied Probability* **33**, 548–556.
- [24] Mi, J. (1999). Optimal active redundancy allocation in k -out-of- n system, *Journal of Applied Probability* **36**, 927–933.

Related Articles

Age-Dependent Minimal Repair and Maintenance; Block Replacement; General Minimal Repair Models; Imperfect Repair; Inspection Policies for Reliability; Maintenance and Markov Decision Models; Maintenance Optimization; Maintenance Optimization in Random Environments; Multivariate Age and Multivariate Renewal Replacement; Renewal Theory; Replacement Strategies; System Availability.

JIE MI

Warranty Cost Prediction Based on Warranty Data

Introduction

Generally, an extended warranty signals higher quality and reliability of a product and provides greater assurance and satisfaction to customers. However, extended warranties impose additional costs to the manufacturer. When a legitimate claim is made under warranty, the manufacturer incurs a cost to rectify the claim, which is known as *warranty cost* (see **Warranty Cost Analysis**). Warranty cost depends on the **warranty policy** and on the lifetime distribution (or reliability) of the product. Prediction of warranty cost means an assertion about the future warranty costs for a specified calendar period or about the amount of usage based on observed costs of claims or incorporating costs data from previous model years. For example, consider an automobile manufacturer that wants to predict the warranty costs up to an age of 5 years or a mileage of 50 000 km, whichever comes first, based on the initial claims data observed during the first year or first 10 000 km. The purposes of estimating and predicting warranty cost are to assess and improve product reliability by fitting the distribution of claims and to better manage the warranty program associated with the given warranty policy. Many approaches may be used to predict warranty costs depending on the warranty policy and on the nature of the available data and information [1]. In general, to predict future claims and costs, we need a statistical model (see **Warranty Modeling**) that depends on a vector of parameters, θ , to describe the claim population under the assumption that the failure modes and failure mechanisms of a product do not vary substantially during the prediction period.

Warranty Cost Prediction Based on Parametric CDF

Let X and Y denote the variables defining the warranty terms. For example, for a refrigerator under a one-dimensional (1D) policy, X would be the age, and for an automobile under two-dimensional (2D) policy, X would be the age and Y would be the mileage (see **Multiattribute Warranty Policies; Warranty Cost Analysis**). $F(x) = F(x; \theta)$

$= \Pr(X \leq x)$ and $F(x, y) = F(x, y; \theta) = \Pr(X \leq x, Y \leq y)$ are used to represent the cumulative distribution function (**cdf**) of X and the joint cdf of X and Y , respectively. Let the average manufacturing cost per unit be denoted by c_m and the average warranty cost (the cost incurred by the seller for servicing a claim) estimated from the data in warranty database be denoted by c_s . Under a 1D warranty with only first failure coverage [2], an estimate of the expected average cost per unit to the manufacturer for servicing a warranty within period $X = x_w$, denoted by $\hat{C}_1(x_w)$, is c_s times the proportion of units expected to fail within x_w ; that is, $\hat{C}_1(x_w) = c_s \hat{F}(x_w)$. Similarly, under a 2D warranty with only first failure coverage, the expected cost within the warranty coverage region ($X = x_w, Y = y_w$), denoted by $\hat{C}_2(x_w, y_w)$, is given by $\hat{C}_2(x_w, y_w) = c_s \hat{F}(x_w, y_w)$.

Under a 1D nonrenewing free replacement warranty [3, 4] for nonrepairable items, the expected warranty cost within period $X = x_w$ is c_m times the expected number of replacements in the interval $[0, x_w]$; that is, $\hat{C}_3(x_w) = c_m \hat{M}(x_w)$, where $M(x)$ is the *renewal function* (see **Renewal Theory**) associated with cdf $F(x)$. Similarly, under a 2D policy, the expected warranty cost within the warranty coverage region ($X = x_w, Y = y_w$) is $\hat{C}_4(x_w, y_w) = c_m \hat{M}(x_w, y_w)$, where $M(x, y)$ is a bivariate renewal function associated with $F(x, y)$ [5]. If the item is repairable (see **Nonparametric Methods for Analysis of Repair Data**), under a 1D warranty, the expected warranty cost is given by

$$\hat{C}_5(x_w) = c_s \int_0^{x_w} \mu(t) dt \quad (1)$$

where $\mu(t)$ is the failure *intensity function* for a *nonhomogeneous Poisson process* (see **Intensity Functions for Nonhomogeneous Poisson Processes**). The expected warranty cost for a 1D nonrenewing pro-rata rebate warranty with linear rebate function $q(x) = (1 - x/x_w)c_b$, $0 \leq x \leq x_w$ [3, 4] is given by

$$\hat{C}_6(x_w) = c_b \left[\hat{F}(x_w) - \frac{\hat{\mu}_{x_w}}{x_w} \right] \quad (2)$$

where c_b is the sale price (cost to buyer) per unit, and $\hat{\mu}_{x_w} (= \int_0^{x_w} x f(x) dx)$ is the partial expectation of X with **pdf** $f(x) = dF(x)/dx$. More detail on these methods and also on other policies are provided in Blischke and Murthy [3, 4, 6] and Murthy and Blischke [7]. The above methods require selection of parametric lifetime distributions that adequately

2 Warranty Cost Prediction Based on Warranty Data

represent the observed and often right censored warranty claims data (*see Warranty Claims and Costs: Statistical Analysis of; Warranty Modeling*). If the warranty database provides information on subsystem level and by failure mode, then the lifetime distributions and average costs can be estimated by subsystem and failure mode wise by their rectification action (repair or replace). These estimates can be used to predict the warranty cost for the overall product.

Menzefricke [8] derived a model to predict warranty costs and to estimate the mean and variance of total warranty cost to the manufacturer during a given period of time. Meeker and Escobar [9, Chapter 12] discussed methods for computing *predictions* and *prediction bounds* for the number of failures in a future time interval. Suppose a single group of units N is placed into service, and $r > 0$ units have failed in the interval $(0, x_c)$. If $F(x; \theta)$ describes the lifetime distribution, the point prediction for the number of additional failures, K , in a future interval (x_c, x_w) is $\hat{K} = (N - r)\hat{\rho}$, where $\hat{\rho} = \frac{\hat{F}(x_w; \theta) - \hat{F}(x_c; \theta)}{1 - \hat{F}(x_c; \theta)}$ is the conditional probability of failing in the interval (x_c, x_w) , given that a unit survived until x_c . The naive $100(1 - \alpha)\%$ upper prediction bound for K , $\tilde{K}(1 - \alpha)$, is computed as the smallest integer k such that $\text{BINCDF}(k; N - r, \hat{\rho}) \geq 1 - \alpha$, (BINCDF means binomial cumulative distribution function). Therefore, the predicted additional warranty cost until x_w would be $c_s \hat{K}$ with naive $100(1 - \alpha)\%$ upper prediction bound $c_s \tilde{K}(1 - \alpha)$. This method can be extended for predicting future claims and costs for more general situations in which multiple groups of units enter into service over a longer period of time (*see* [9]).

Age-Based Warranty Cost Prediction (Nonparametric Approach)

The *age-based* (or *age-specific*) evaluation of warranty cost has engendered considerable interest [10–14]. The age-based evaluation of warranty costs means **estimation** and prediction of warranty costs assuming that the expected number of claims per product at age t depends on t . Here we consider a month as the unit for age without loss of generality. If necessary, the unit “month” can be easily replaced by “week”, “day” and so on. Let S denotes the number of months of sale, M denotes the length of the observation period ($S \leq M$), and W denotes the length of the

warranty period. We assume that observed claims are stratified into k groups ($k = 1, 2, \dots, K$) based on the cost and define $c(k)$ as the average cost of a claim in group k , N_s as the number of products sold in month s , and $\{r_{st}(k), s = 1, 2, \dots, S; t = 1, 2, \dots, \min(M - s + 1, W); k = 1, 2, \dots, K\}$ as the reported number of claims of cost $c(k)$ at age t for those products sold in month s .

Let $\lambda_t(k)$ be the expected number of claims of cost $c(k)$ for a unit at age t , $t = 1, 2, \dots, \min(W, M)$. If claims occur randomly, the expected value of $r_{st}(k)$ is $N_s \lambda_t(k)$. Then the moment estimate of $\lambda_t(k)$ [12] can be expressed as

$$\hat{\lambda}_t(k) = \frac{\sum_{s=1}^{\min(S, M-t+1)} r_{st}(k)}{\sum_{s=1}^{\min(S, M-t+1)} N_s} = \frac{r_{\cdot t}(k)}{R_t} \quad (3)$$

where $R_t = \sum_{s=1}^{\min(S, M-t+1)} N_s$ is the total number of products sold up to month $\min(S, M - t + 1)$, and $r_{\cdot t}(k) = \sum_{s=1}^{\min(S, M-t+1)} r_{st}(k)$ is the total number of claims of cost $c(k)$ for units of age t reported up to month $\min(S, M - t + 1)$, for $t = 1, 2, \dots, \min(W, M)$, $k = 1, 2, \dots, K$. The estimate of cumulative total expected warranty cost per product up to age t is $\hat{C}_t = \sum_{k=1}^K c(k) \hat{\lambda}_t(k)$, where $\hat{\lambda}_t(k) = \sum_{i=1}^t \hat{\lambda}_i(k)$, $t = 1, \dots, \min(W, M)$.

To predict warranty cost, we define

$$Q_t(k) = \frac{1}{N} \sum_{i=1}^t \sum_{s=1}^{\min(S, M-t+1)} r_{si}(k) \quad (4)$$

as the average number of claims of cost $c(k)$ up to age t , per product sold, where $N = \sum_{s=1}^S N_s$, and

$$Q_t = \sum_{k=1}^K c(k) Q_t(k) \quad (5)$$

as the average total cost per product up to age t , $t = 1, 2, \dots, \min(W, M)$. Then $\hat{\lambda}_t(k)$ and $\hat{C}_t$ can be used to estimate $Q_t(k)$ and Q_t , respectively. Estimation of the variance and normal-approximation confidence limits for Q_t are explained in Kalbfleisch *et al.* [10] and Kalbfleisch and Lawless [11] for a more general case with reporting lag. This method provides an age-specific prediction of warranty costs for a finite population of units. If we wish to predict warranty costs for t greater than the length of claim-observed period,

M , $\lambda_t(k)$ must be estimated for $t > M$ using sources other than the current data. When the objective is to predict costs over calendar time, the predicted total cost of claims in month m , $\hat{Q}_m^*$, can be expressed by

$$\hat{Q}_m^* = \sum_{k=1}^K \sum_{s=\max(1, m-W+1)}^{\min(m, S)} c(k) N_s \hat{\lambda}_{m-s+1}(k), \quad m = 1, 2, \dots, M \quad (6)$$

For a future month m , the values of N_s for $s = S+1, \dots, m$, may also have to be estimated if the sales of the units are continued after month S , such that $m > S$.

Warranty Cost Prediction Based on Bayesian Approach

There is a limited literature on Bayesian approaches (see **Bayesian Reliability Analysis**) that addresses forecasting warranty claims and costs. Chen *et al.* [15] and [16] discussed methods for forecasting warranty claims and costs by using Markov mesh models which can be represented as a Bayesian dynamic linear models (DLM's). They proposed the DLM with leading indicators that used to incorporate data from previous model year and presented examples for forecasting the number of repairs and the cost of reporting a unit. Stephens and Crowder [17] also adopted the Bayesian approach with Markov chain Monte Carlo (MCMC) (see **Hierarchical Markov Chain Monte Carlo (MCMC) for Bayesian System Reliability**) sampling for predicting future warranty exposure.

Example

1. Suppose that the lifetime T of a product is distributed Weibull that is,

$$f(t; \eta, \beta) = \frac{\beta}{\eta} \left(\frac{t}{\eta}\right)^{\beta-1} \exp \left[-\left(\frac{t}{\eta}\right)^{\beta} \right], \quad t > 0 \quad (7)$$

with scale parameter $\eta = 5$ years and shape parameter $\beta = 1.5$. To model warranty claim data and estimate model parameters, see **Warranty Claims and Costs: Statistical Analysis of; Warranty Modeling; Warranty Cost Analysis**. Consider that for this product the manufacturing cost per unit, $c_m = \$300$, the sale price per unit, $c_b = \$500$, and the average cost incurred by the seller for servicing a claim, $c_s = \$50$.

- (a) Assuming only first failure coverage warranty, the predicted warranty cost per unit to the manufacturer for servicing a warranty for a warranty period of 2 year is $C_1(x_w = 2) = c_s \hat{F}(2) = 50 \times 0.2235 = \11.18 .
 - (b) Considering a nonrenewing free replacement warranty, the predicted warranty cost up to a warranty period of 2 year is $C_3(x_w = 2) = c_m \hat{M}(2) = 300 \times 0.2407 = \72.21 . For evaluation of renewal function, $M(\cdot)$, see Blischke and Murthy [3, 6].
 - (c) If the item is repairable (using **minimal repair**), under the Weibull nonhomogeneous Poisson process, the predicted warranty cost up to a warranty period of 2 year is $C_5(x_w = 2) = c_s \int_0^2 \mu(t) dt = c_s \times (2/\eta)^\beta = 50 \times 0.2530 = \12.65 .
 - (d) Assuming a nonrenewing pro-rata rebate warranty with linear rebate function $q(x) = (1 - x/x_w)c_b$, $0 \leq x < x_w$, from equation (2) the predicted warranty cost per unit to manufacturer for a warranty period of 2 year is $C_6(x_w = 2) = c_b \left[\hat{F}(2) - \frac{1}{2} \int_0^2 t f(t) dt \right] = 500 \times 0.0936 = \46.80 .
2. Pal and Murthy [5] showed that the warranty claim data of a motorcycle collected within the 2D warranty coverage region (age $X = 180$ days, usage $Y = 8000$ km) follows Gumbel's bivariate exponential model as

$$F(x, y) = 1 - e^{-\left(\frac{x}{\theta_x}\right)} - e^{-\left(\frac{y}{\theta_y}\right)} + e^{-\left\{\left(\frac{x}{\theta_x}\right)^\delta + \left(\frac{y}{\theta_y}\right)^\delta\right\}^{1/\delta}}, \quad x, y \geq 0 \quad (8)$$

with the estimates of the parameters $\hat{\theta}_x = 357.0$, $\hat{\theta}_y = 10077.5$ and $\hat{\delta} = 3.15$. The average warranty cost per claim for this product is $c_s = 428.4$ rupees. Therefore, the predicted warranty cost for the extended warranty coverage region ($X = x_w = 365$ days, $Y = y_w = 12000$ km) is $\hat{C}_2(365, 12000) = c_s \hat{F}(365, 12000) = 428.4 \times 0.586 = 251.04$ rupees.

3. Suppose that a product has a 6-month warranty period and that $N_s = 1000$ units are sold on each month $s = 1, 2, \dots, S$. Assume that claims are observed during 6-month period ($S = 6$) and

classified into three groups ($K = 3$) based on warranty cost, and the estimates of the expected number of claims $\hat{\lambda}_t(k)$ defined by equation (3) for three groups are $\{\hat{\lambda}_t(1), \hat{\lambda}_t(2), \hat{\lambda}_t(3)\} = \{(0.004, 0.006, 0.008, 0.007, 0.013, 0.018), (0.004, 0.008, 0.005, 0.006, 0.009, 0.015), (0.003, 0.009, 0.008, 0.009, 0.005, 0.007)\}$. For simplicity of illustration, assume that on each month the same amount 1000 units are continued to sale up to 12 month. If the average cost of a claim in groups one, two and three are $c(1) = \$10.0$, $c(2) = \$20.0$ and $c(3) = \$30.0$, respectively, then using equation (6) the predicted cost at month 12 ($m = 12$) is $\hat{Q}_{12}^* = \sum_{k=1}^3 \sum_{s=7}^{12} c(k) N_s \hat{\lambda}_{12-s+1}(k) = \2730 .

Summary

This article reviewed some methods for predicting warranty cost during a given future value(s) of the variable(s) defining the warranty terms. It modeled the predicted cost as a function of the lifetime distribution of the product and the average cost per claim based on observed costs of previous claims or incorporating costs data from previous model years.

References

- [1] Murthy, D.N.P. & Djamaludin, I. (2002). New product warranty: a literature review, *International Journal of Production Economics* **79**, 231–260.
- [2] Iskandar, B.P. & Blischke, W.R. (2003). Reliability and warranty analysis of a motorcycle based on claims data, in *Case Studies in Reliability and Maintenance*, W.R. Blischke & D.N.P. Murthy, eds, John Wiley & Sons, New York, pp. 623–656.
- [3] Blischke, W.R. & Murthy, D.N.P. (1996). *Product Warranty Handbook*, Marcel Dekker, New York.
- [4] Blischke, W.R. & Murthy, D.N.P. (2000). *Reliability Modeling, Prediction, and Optimization*, John Wiley & Sons, New York.
- [5] Pal, S. & Murthy, G.S.R. (2003). An application of Gumbel's bivariate exponential distribution in estimation of warranty cost of motor cycles, *International Journal of Quality and Reliability Management* **20**(4), 488–502.
- [6] Blischke, W.R. & Murthy, D.N.P. (1994). *Warranty Cost Analysis*, Marcel Dekker, New York.
- [7] Murthy, D.N.P. & Blischke, W.R. (2001). Warranty and reliability, in *Handbook of Statistics: Advances in Reliability*, N. Balakrishnan & C.R. Rao, eds, Elsevier Science, Vol. 20, pp. 541–583.
- [8] Menzefricke, U. (1993). Prediction of warranty costs during a given period of time, *INFOR* **31**, 127–135.
- [9] Meeker, W.Q. & Escobar, L.A. (1998). *Statistical Methods for Reliability Data*, John Wiley & Sons, New York.
- [10] Kalbfleisch, J.D., Lawless, J.F. & Robinson, J.A. (1991). Methods for the analysis and prediction of warranty claims, *Technometrics* **33**, 273–285.
- [11] Kalbfleisch, J.D. & Lawless, J.F. (1996). Statistical analysis of warranty claims data, in *Product Warranty Handbook*, W.R. Blischke & D.N.P. Murthy, eds, Marcel Dekker, New York, Chapter 9, pp. 231–259.
- [12] Lawless, J.F. (1998). Statistical analysis of product warranty data, *International Statistical Review* **66**(1), 41–60.
- [13] Suzuki, K., Karim, M.R. & Wang, L. (2001). Statistical analysis of reliability warranty data, in *Handbook of Statistics: Advances in Reliability*, N. Balakrishnan & C.R. Rao, eds, Elsevier Science, Vol. 20, pp. 585–609.
- [14] Karim, M.R. & Suzuki, K. (2005). Analysis of warranty claim data: a literature review, *International Journal of Quality and Reliability Management* **22**(7), 667–686.
- [15] Chen, J., Lynn, N.J. & Singpurwalla, N.D. (1995). Markov mesh models for filtering and forecasting with leading indicators, in *Analysis of Censored Data, IMS Lecture Notes-Monograph Series, Vol. 27*, H.L. Koul & Deshpande, J.V., eds, Institute of Mathematical Statistics, pp. 39–54.
- [16] Chen, J., Lynn, N.J. & Singpurwalla, N.D. (1996). Forecasting warranty claims, in *Product Warranty Handbook*, W.R. Blischke & D.N.P. Murthy, eds, Marcel Dekker, New York, Chapter 31, pp. 803–817.
- [17] Stephens, D. & Crowder, M. (2004). Bayesian analysis of discrete time warranty data, *Applied Statistics* **53**, 195–217.

Related Articles

Intensity Functions for Nonhomogeneous Poisson Processes; Multiattribute Warranty Policies; Non-parametric Methods for Analysis of Repair Data; Renewal Theory; Warranty Claims and Costs; Statistical Analysis of; Warranty: Usage and Wear Process for; Warranty Cost Analysis; Warranty Modeling; Warranty Servicing.

M. REZAUL KARIM AND KAZUYUKI SUZUKI

Warranty Servicing

Introduction

When a warranty is given with the sale of a product, the manufacturer has to service all claims made by the customers. Warranty costs (*see* **Warranty Cost Analysis; Warranty Cost Prediction Based on Warranty Data**) are the additional expenses that the manufacturer incurs in servicing these **warranty claims**. Models to estimate warranty costs for many different types of warranties can be found in Blischke and Murthy [1, 2], and a review of the literature on **warranty cost** modeling from 1990 to 2000 is given in Murthy and Djamaludin [3].

Warranty costs can be reduced by performing effective warranty servicing. To do this, the manufacturer needs to address various issues at both the strategic and operational levels. In this article, we develop a framework that links all the issues involved in warranty servicing and we review the models that have been developed to make optimal servicing decisions. The outline of the article is as follows. The Section titled “Framework for Warranty Servicing” deals with the servicing framework and, in the section titled “Servicing Strategies”, we review servicing strategies that involve repair *versus* replace decisions for the manufacturer at product failures. Finally, we state our conclusions.

Framework for Warranty Servicing

The warranty servicing process is shown in Figure 1.

Warranty Claims

The number of warranty claims (*see* **Warranty Claims and Costs: Statistical Analysis of**) that a manufacturer has to service is a function of (a) the terms of the **warranty policy**, (b) the number of items sold, and (c) the number of product failures. The number of failures depends on the inherent product reliability (influenced by the design and production decisions taken by the manufacturer) and on the usage mode and environment (influenced by the consumer).

Warranty Policies. The two most common warranty policies (*see* **Multiattribute Warranty Policies;**

Warranty Modeling) for consumer durables and industrial or commercial products are the prorata warranty (PRW) and the free replacement warranty (FRW). At each product failure, the manufacturer refunds a fraction of the purchase price to the customer under a PRW and repairs or replaces the item free of charge under a FRW. Both of these warranties may also be characterized as either renewing or nonrenewing. In the renewing case, all replaced or repaired items are covered by a warranty identical to that of the original purchased product. In the nonrenewing case, the coverage for a replaced or repaired item is only for the remainder of the original warranty period. Another important characteristic of a warranty is its dimensionality. In a one-dimensional warranty, the product is covered for a certain time interval and, in the two-dimensional case, the warranty coverage is characterized by a two-dimensional planar region with one axis representing time and the other axis representing usage. A taxonomy for warranty policies is proposed by Blischke and Murthy [2].

Product Reliability. A product is a complex item that fails due to the failure of one or more of its many components. The product can be restored to an operational state by replacing or repairing all the failed components (*see* **Replacement Strategies; Imperfect Repair; General Minimal Repair Models; Repairable Systems Reliability**) and the time taken to do this is usually small relative to the time between failures. The manufacturer needs to be able to estimate the number of product and component failures for repair-facility planning.

At the product level, failure occurrences can be modeled by a stochastic point process with a given **intensity function**. In order to model component level failure, a renewal process (*see* **Markov Renewal Processes in Reliability Modeling**) is used in the nonrepairable case and a point process is required for a repairable component. Details of these modeling techniques can be found in Blischke and Murthy [1, 2].

Product Usage. Product usage intensity (such as number of miles traveled per year by a car, number of times a washing machine is used per week – *see* **Warranty: Usage and Wear Process for**) and operating environment vary across the customer population. This variation changes the rate at which products fail and the number of warranty claims that

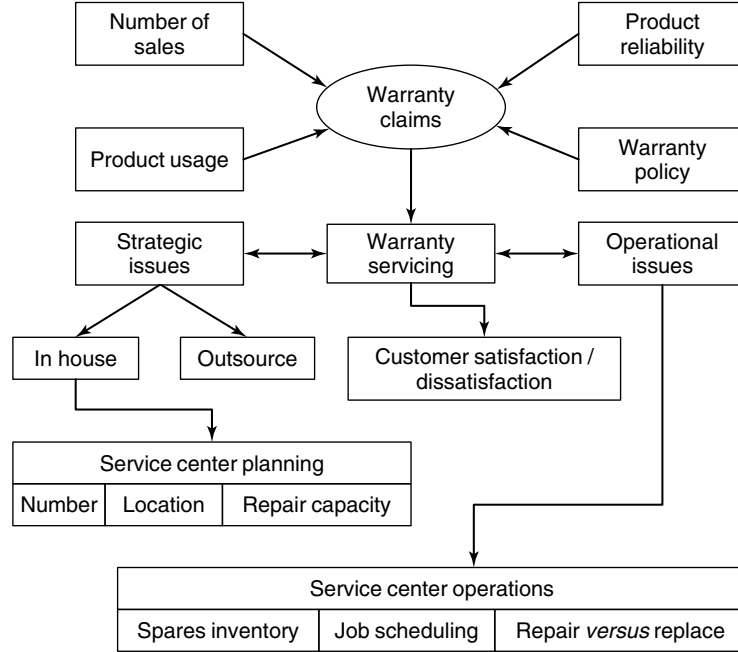


Figure 1 Warranty servicing process

need to be serviced. Cost models for products sold with a FRW under varying usage intensity are studied by Kim *et al.* [4].

Product Sales and Warranty Claims. Let L denote the length of the product life cycle, which is the time period from when the product is first introduced into the market until the instant it is replaced in the market by a new and better product. Let W denote the length of the warranty period. If the product is sold with non-renewing FRW policy, then warranty claims occur at random points over the interval starting at time zero when the product is first introduced and ending at time $(L + W)$ when the warranty ceases for the last items sold.

If the intensity function (given by the rate of occurrence of failures (ROCOF)) for product failures is known then the expected warranty claims rate (expected number of claims per unit time) over the interval $[0, L + W)$ can be determined from this ROCOF and the sales rate (sales per unit time), see Blischke and Murthy [1, 2]. Typical plots of the sales rate and the corresponding expected claims rate are given in Figure 2. Both curves show an initial increase, reach a maximum level, and then decrease

to zero with the expected claims rate always lagging the sales rate.

Warranty Servicing

The strategic issues in warranty servicing that the manufacturer needs to consider are (a) the number of service centers and their location, (b) the capacity and staffing for each service center (to ensure prompt response times for customer satisfaction), and (c) whether to own these centers or to outsource the servicing to an independent agent.

The tactical and operational issues are (a) the transportation of the material needed for warranty

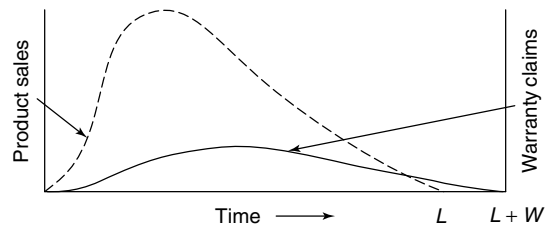


Figure 2 Sales rate and expected warranty claims

servicing, (b) management of spare parts inventory, (c) job scheduling, and (d) optimal repair/replace decisions for failed items.

Strategic Issues. The location of service centers (for carrying out repairs) and warehouses (for stocking spares needed) depend on the geographical distribution of customers who have bought the product, the type of product, and its reliability characteristics.

Sometimes, it is necessary to carry out on-site evaluation and repair of a failed product (such as a lift in a building). Otherwise, the failed item is taken to either the retailer (in the case of most consumer durables) or directly to a designated service center. Depending on the complexity of the product, it might also have to be moved between service centers operating at different levels to have its repair completed. The problem of deciding where an item should be repaired in a multiechelon repair facility has received some attention in the logistics literature and is known as *level of repair analysis* (LORA). Alfredsson [5] deals with decisions with regard to LORA and the spare parts and test equipment needed to support a system. Barros and Riley [6] deal with optimization of the LORA.

Models to determine the number of levels, the location of the service centers and their capacities must take into account the transportation time and cost for moving failed and repaired items between service centers, and the cost of operating the service centers (equipment and skilled persons needed at each center).

The manufacturer might also need a network of warehouses with a multiechelon structure involving one or more levels. Two key issues are the location of the warehouses and the capacity of each warehouse, defined through the quantities of different components that need to be stocked. This problem has received some attention in the logistics literature. However, the models need to be modified to take into account the location of the service centers and the reliability characteristics of the product.

The optimal warehouse location must take into account the transportation time and cost for moving parts between warehouses (in the case of a multiechelon structure) and from warehouses to service centers, the cost of operating the warehouses, and the capacity of each warehouse based on the demand for spares from the various service centers that are serviced by the warehouse.

The service center capacity requirement is determined by the service demand, which depends on the geographical distribution of sales and on product reliability. In solving the location problem, the issues of coverage (so that all the customers are covered), distance that a failed item must travel, and distance that a repairman has to travel in case of a field visit or that a customer has to travel to bring a failed item to a service center or collection point have to be addressed. Given the coverage, the model based on sales over a region can be used to determine the demand and the service center capacity.

The management of inventory levels for service support logistics deals with ordering policies. There is a trade-off between the cost associated with, and the benefits derived from, having an inventory, and several ordering strategies have been proposed and studied. These include the reorder point control, fixed interval control, and the min-max control (see Nahmias [7] and Sherbrooke [8]). The total demand for spares (of a component) at a warehouse over the product life cycle can be modeled using the techniques described in the section titled "Product Sales and Warranty Claims".

If the manufacturer decides to outsource the warranty servicing to independent agents then properly drafted contracts are essential. The interests of the two parties (manufacturer and agent) are different and this raises several new issues that are highlighted in the following example.

Contract A: The manufacturer pays a lump sum to the agent and in return the agent has to service all claims during the warranty period at no additional cost to the manufacturer.

Contract B: The service agent charges the manufacturer for each warranty service.

The new and relevant issues here include:

Informational asymmetry

The manufacturer has better knowledge of product reliability compared with the agent and similarly the agent has better information regarding field failures than the manufacturer. This asymmetry can lead to each party deciding on actions that are optimal from their own individual perspective but overall suboptimal.

Moral hazard

This situation arises when the agent shirks in the effort expended (under Contract A) or carries out

4 Warranty Servicing

overservicing (under Contract B) and the manufacturer is unable to observe the service agent's effort. In the former case, this can lead to customer dissatisfaction which can affect the manufacturer's reputation and sales.

Monitoring

The manufacturer can obtain new information by monitoring the agent's actions. Such information will allow the manufacturer to assess the warranty servicing carried out by the agent but this results in additional effort and cost to the manufacturer.

Adverse selection

This issue arises when the manufacturer has to choose one or more service agents to carry out the warranty servicing from a pool of service agents. The service agents can misrepresent their ability and competencies and the manufacturer is not able to assess them prior to the signing of the contract. This can lead to the selection of inappropriate agents to service warranty.

Incentives

The manufacturer can provide proper incentives to the service agents so that the actions of the agents are in the best interests of the manufacturer and avoid the need for monitoring. A proper contract provides the right incentives for the service agent to provide the optimal effort.

Agency cost

The structuring, administering, and enforcing of contracts result in a cost that is referred to as the *agency cost*.

Risks

The manufacturer and the service agents have partly differing goals and risk preferences and both of these have an impact on their individual actions.

Agency theory (see Eisenhardt [9] and van Ackere [10]) can provide the framework needed to deal with the above issues, but it has yet to be applied to warranty servicing contracts.

Operational Issues. The key issues in spare parts inventory management are, making decisions about the components that should be carried as spare parts, the inventory levels of the spares, and the timing and amount of spares that should be ordered. To do this, information is required about the expected

number of component failures over time and this, in turn, depends on the level of sales over time and on **component reliability**.

Most models dealing with spare part inventories assume very simple forms for the depletion of inventory. In the warranty-servicing context, the depletion of the inventory of a particular component occurs in an uncertain manner that depends on the sales rate over the region serviced by the servicing center and the reliability of the product. Optimal decisions with regard to inventory levels and ordering policies need to take this into account.

In warranty servicing logistics, the three kinds of material that needs to be transported are (a) failed units from a lower level to a higher level in the case of a multiechelon service structure, (b) repaired items from service centers to customers or pickup points where customers can collect them, and (c) spares to and from warehouses. In each case, the amounts to be moved are random variables related to sales and product reliability.

The material transportation can be carried out either by the manufacturer or by an independent agent. In the latter case, a contract between the manufacturer and the agent must take into account the transportation cost, the frequency of transportation, upper limits on quantities to be transported, time limits, penalties for delays in delivery, breaches of contract, and so on. The Agency theory issues discussed in the section titled "Strategic Issues" are also relevant here and different contract options may be evaluated using this framework.

Products may be differentiated according to whether they are brought to a service center when they fail or whether a repairman goes out to rectify the failure. In the former case, the scheduling of repair jobs is an important issue that not only has an impact on the overall cost of providing the service but also on customer satisfaction. Job scheduling has been extensively studied [11–13]. The latter case results in the "traveling repairman problem", and a number of solutions have been proposed [14–16] to schedule jobs to reduce traveling time. If the warranty on the product includes penalties for servicing delays, then the scheduling also needs to take this into account.

The remaining operational issue in warranty servicing on repair *versus* replace decisions is discussed in the next section.

Servicing Strategies

Whenever a repairable product fails under warranty, the manufacturer has the option of either repairing the failed item or replacing it with a new one. The first option costs less than the second but the repaired item has a greater probability of failing during the remainder of the warranty period. It is therefore important for the manufacturer to choose an appropriate servicing strategy in order to minimize the expected cost of servicing the warranty for each product sold (*see Warranty Cost Analysis*).

Servicing strategies for products sold with one-dimensional warranties have received considerable attention. In a cost-repair limit strategy, an estimate of the cost to repair a failed product is made and by comparing it with some specified limit, the failed item is either repaired or replaced by a new one. Optimal repair limit strategies are discussed in Blischke and Murthy [1] and Zuo *et al.* [17].

In most of the remaining models in the literature, average repair costs are used to make the repair *versus* replace decisions. Biedenweg [18] and Nguyen and Murthy [19, 20] assume that repaired items have independent and identically distributed lifetimes different from that of a new item and consider strategies where the warranty period is divided into distinct intervals for repair and replacement. Nguyen [21] introduces the first servicing model with minimal repair (*see* Barlow and Hunter [22] and **General Minimal Repair Models**), the warranty period split into a replacement interval followed by a repair interval. The length of the first interval is selected optimally to minimize the expected warranty cost.

Jack and Van der Duyn Schouten [23] show that this strategy is suboptimal and that the optimal servicing strategy is in fact characterized by three distinct intervals – $[0, x]$, $[x, y]$ and $(y, W]$ where W is the warranty period. The optimal strategy is to carry out minimal repairs in the first and last intervals and to use either minimal repair or replacement by new in the middle interval depending on the age of the item at failure. (Refer to Jiang *et al.* [24] for the optimality of this strategy.) This strategy is difficult to implement, so Jack and Murthy [25] proposed a near optimal strategy involving the same three intervals but with only the first failure in the middle interval resulting in a replacement and all other failures being minimally repaired. (Refer to Sheu and Yu [26] for the case with random minimal repair cost.)

When the cost of replacement is high compared to the cost of a minimal repair then strategies involving replacement are not appropriate. In this case, **imperfect repair** strategies (where the failure characteristics of the repaired item are better than those after minimal repair but are not the same as a new item) can be used. The advantage of using imperfect repair (*see Imperfect Repair*) is that the degree of improvement in the reliability after repair is a decision variable under the control of the manufacturer. The literature on imperfect repair has focused mainly on its use in the context of preventive and corrective maintenance of unreliable systems. See, for example, Wang [27], Li and Shaked [28], Doyen and Gaudoin [29], Zequeira and Berenguer [30], Sheu *et al.* [31], and Tang and Lam [32]. Chukova *et al.* [33] study the successive operating times of a product under different types of imperfect repair.

Yun *et al.* [34] consider two servicing strategies (servicing strategies 1 and 2 respectively) involving minimal and imperfect repairs. In both strategies a failed product is subjected to at most one imperfect repair over the warranty period.

In strategy 1, all item failures occurring in the interval $[0, x)$ are minimally repaired. The first failure in the interval $[x, y]$ is rectified by imperfect repair with the proportional reduction factor in the hazard rate $\delta(t)$ depending on the age t at failure. The cost of an imperfect repair is an increasing function of the item age and the proportional reduction in the hazard rate. All subsequent failures in the interval $[x, y]$ are minimally repaired, as are those that occur in the remaining interval $(y, W]$. This strategy is characterized by the set $\{x, y, \Delta(x, y)\}$ consisting of the two parameters x and y , and the function $\Delta(x, y) \equiv \{\delta(t), x \leq t \leq y\}$. These are the decision variables that need to be selected optimally to minimize the expected warranty servicing cost $J(x, y, \Delta(x, y))$. Strategy 1 is similar to that proposed by Jack and Murthy [25] except that an imperfect repair is performed instead of a replacement in the interval $[x, y]$.

In strategy 2, the proportional reduction in the hazard rate does not change with time over the interval $[x, y]$. It is a parameter (δ) to be selected optimally. In this case, the strategy is characterized by the three parameters x , y , and δ , and these are to be selected optimally to minimize the expected warranty servicing cost $J(x, y, \delta)$.

Servicing strategies for products sold with two-dimensional warranties have also received some

attention. Iskandar and Murthy [35] examined two strategies similar to those in Nguyen and Murthy [19, 20] but with minimal repair. Iskandar *et al.* [36] discussed a servicing strategy similar to that given in Jack and Murthy [25]. (Refer to Chukova and Johnston [37] for the case without a restriction assumed by Iskandar *et al.* [36].)

Conclusions

There are still many aspects of warranty servicing that need further investigation. These include (a) the design of incentive contracts based on Agency theory when the servicing is outsourced, (b) models to study the various operational and strategic issues involved in warranty logistics, and (c) new servicing strategy models for products sold with two-dimensional warranties.

References

- [1] Blischke, W.R. & Murthy, D.N.P. (1994). *Warranty Cost Analysis*, Marcel Dekker, New York.
- [2] Blischke, W.R. & Murthy, D.N.P. (1996). *Product Warranty Handbook*, Marcel Dekker, New York.
- [3] Murthy, D.N.P. & Djameludin, I. (2002). Product warranty – a review, *International Journal of Production Economics* **79**, 231–260.
- [4] Kim, C.S., Djameludin, I. & Murthy, D.N.P. (2001). Warranty cost analysis with heterogeneous usage intensity, *International Transactions in Operational Research* **8**, 337–347.
- [5] Alfredsson, P. (1997). Optimization of multi-echelon repairable item inventory systems with simultaneous location of repair facilities, *European Journal of Operational Research* **99**, 584–595.
- [6] Barros, L. & Riley, M. (2001). A combinatorial approach to level of repair analysis, *European Journal of Operational Research* **129**, 242–251.
- [7] Nahmias, S. (1997). *Production and Operational Analysis*, McGraw-Hill, New York.
- [8] Sherbrooke, C.C. (1992). *Optimal Inventory Modeling of Systems: Multi-Echelon Techniques*, John Wiley & Sons, New York.
- [9] Eisenhardt, K.M. (1989). Agency theory: an assessment and review, *Academy of Management Review* **14**, 57–74.
- [10] van Ackere, A. (1993). The principal/agent paradigm: its relevance to various fields, *European Journal of Operational Research* **70**, 88–103.
- [11] Hajri, S., Liouane, N., Hammadi, S. & Borne, P. (2000). A controlled genetic algorithm by fuzzy logic and belief functions for job-shop scheduling, *IEEE Transactions on Systems, Man and Cybernetics, Part B (Cybernetics)* **30**, 812–818.
- [12] Jianer, C. & Miranda, A. (2001). A polynomial time approximation scheme for general multiprocessor job scheduling, *SIAM Journal on Computing* **31**, 1–17.
- [13] Ponnambalam, S.G., Ramkumar, V. & Jawahar, N. (2001). A multiobjective genetic algorithm for job shop scheduling, *Production Planning and Control* **12**, 764–774.
- [14] Afrati, F., Cosmadakis, S., Papadimitriou, C.H., Papa-georgiou, G. & Papakostantinou, N. (1986). The complexity of the traveling repairman problem, *Informatique Theorique et Applications* **20**, 79–87.
- [15] Agnihothri, S.R. (1998). Mean value analysis of the traveling repairman problem, *IIE Transactions* **20**, 223–229.
- [16] Yang, C.E. (1989). A dynamic programming algorithm for the traveling repairman problem, *Asia-Pacific Journal of Operational Research* **6**, 192–206.
- [17] Zuo, M.J., Liu, B. & Murthy, D.N.P. (2000). Replacement–repair policy for multi-state deteriorating products under warranty, *European Journal of Operational Research* **123**, 519–530.
- [18] Biedenweg, F.M. (1981). *Warranty analysis: consumer value vs manufacturers cost*, Unpublished Ph.D. Thesis, Stanford University.
- [19] Nguyen, D.G. & Murthy, D.N.P. (1986). An optimal policy for servicing warranty, *Journal of the Operational Research Society* **37**, 1081–1088.
- [20] Nguyen, D.G. & Murthy, D.N.P. (1989). Optimal replace-repair strategy for servicing items sold with warranty, *European Journal of Operational Research* **39**, 206–212.
- [21] Nguyen, D.G. (1984). *Studies in warranty policies and product reliability*, Unpublished Ph.D. Thesis, The University of Queensland.
- [22] Barlow, R.E. & Hunter, L. (1960). Optimal preventive maintenance policies, *Operations Research* **8**, 90–100.
- [23] Jack, N. & Van der Duyn Schouten, F. (2000). Optimal repair-replace strategies for a warranted product, *International Journal of Production Economics* **67**, 95–100.
- [24] Jiang, X., Jardine, A.K.S. & Lugtigheid, D. (2006). On a conjecture of optimal repair-replacement strategies for warranted products, *Mathematical and Computer Modelling* **44**, 963–972.
- [25] Jack, N. & Murthy, D.N.P. (2001). A servicing strategy for items sold under warranty, *Journal of the Operational Research Society* **52**, 1284–1288.
- [26] Sheu, S.H. & Yu, S.L. (2005). Warranty strategy accounts for bathtub failure rate and random minimal repair cost, *Computers and Mathematics with Applications* **49**, 1233–1242.
- [27] Wang, H. (2002). A survey of maintenance policies of deteriorating systems, *European Journal of Operational Research* **139**, 469–489.
- [28] Li, H. & Shaked, M. (2003). Imperfect repair models with preventive maintenance, *Journal of Applied Probability* **40**, 1043–1059.

-
- [29] Doyen, L. & Gaudoin, O. (2004). Classes of imperfect repair models based on reduction of failure intensity or virtual age, *Reliability Engineering and System Safety* **84**, 45–56.
- [30] Zequeira, R.I. & Bérenguer, C. (2006). Periodic imperfect preventive maintenance with two categories of competing failure modes, *Reliability Engineering and System Safety* **91**, 460–468.
- [31] Sheu, S.H., Lin, Y.B. & Liao, G.L. (2006). Optimum policies for a system with general imperfect maintenance, *Reliability Engineering and System Safety* **91**, 362–369.
- [32] Tang, Y.Y. & Lam, Y. (2006). A δ -shock maintenance model for a deteriorating system, *European Journal of Operational Research* **168**, 541–556.
- [33] Chukova, S., Arnold, R. & Wang, D.Q. (2004). Warranty analysis: an approach to modeling imperfect repairs, *International Journal of Production Economics* **89**, 57–68.
- [34] Yun, W.-Y., Murthy, D.N.P. & Jack, N. (2006). Warranty servicing with imperfect repair, *International Journal of Production Economics* In Print.
- [35] Iskandar, B.P. & Murthy, D.N.P. (2003). Repair-replace strategies for two-dimensional warranty policies, *Mathematical and Computer Modelling* **38**, 1233–1241.
- [36] Iskandar, B.P., Murthy, D.N.P. & Jack, N. (2005). A new repair-replace strategy for items sold with a two-dimensional warranty, *Computers and Operations Research* **32**, 669–682.
- [37] Chukova, S. & Johnston, M.R. (2006). Two-dimensional warranty repair strategy based on minimal and complete repairs, *Mathematical and Computer Modelling* **44**, 1133–1143.

Further Reading

- Murthy, D.N.P. & Blischke, W.R. (2005). *Warranty Management and Product Manufacture*, Springer-Verlag, London.
- Murthy, D.N.P., Solem, O. & Roren, T. (2003). Product warranty logistics: issues and challenges, *European Journal of Operational Research* **156**, 110–126.

Related Articles

Warranty Claims and Costs: Statistical Analysis of; Warranty: Usage and Wear Process for; Multiattribute Warranty Policies; Warranty Modeling ; Replacement Strategies; Imperfect Repair; General Minimal Repair Models; Markov Renewal Processes in Reliability Modeling; Warranty Cost Analysis; Warranty Cost Prediction Based on Warranty Data; Repairable Systems Reliability.

D.N. PRABHAKAR MURTHY AND JACK NAT

Sampling Inspection of Products

Statistical Product Inspection in Industrial Quality Control

Industrial **quality control** can be systematized in different ways. A popular systematization is arranged by *subject matter*: *design control*, **process control**, *product control*, where the term “product” is used in a broad sense meaning any completed identifiable entity like materials, parts, products, services, software, or data, see the definition by ISO 9000, [1]. In manufacturing, product control is often subdivided into *incoming material control (input control)* and *finished product control (output control)*. Product control is situated at the interface of two parties: a *supplier* or *vendor* or *producer* as the party providing product and a *customer* or *consumer* as the party receiving product. Supplier and customer may be part of the same organization (internal supplier) or of different organizations (external supplier).

Another systematization of quality control follows the *control objectives*: *preventive* control is distinguished from *monitoring* or *inspection* of results. Prevention is characterized by the interest of avoiding quality problems by creating an environment favorable for quality. Monitoring collects information on the results of routine business for assessment purposes. The assessment leads to a decision between two alternatives *accept* or *reject*. In a reasonable industrial environment, acceptance is the regular case and means *continue with routine business* whereas rejection calls for some extraordinary *intervention*.

Design control is essentially preventive in nature. Process and product control have both preventive and inspection aspects. Inspection in product control occurs as *receiving inspection* of incoming material, parts, products, or as *final inspection* of outgoing finished product. The inspection serves as a basis for an assessment of material or product with a subsequent decision on the alternatives accept or reject.

Inspection and decision in product control can take two directions. A *type A* procedure is directed toward a single assembled *lot* of product. A *type B* procedure is directed toward the *process* of manufacturing, transport, storage which provides the product. The

distinction between type A and type B procedures is adopted by ISO 2859-2 [2], for instance.

A type A procedure is concerned with the decision alternative *acceptance of the lot* and *rejection of the lot*, with the following interpretations. *Acceptance of the lot* (routine case): (a) final inspection: forwarding the lot to another department for processing, shipping the lot to a customer. (b) receiving inspection: accept the lot from the shipper for reselling or processing. *Rejection of the lot* (intervention): 100% rectification of the lot, i.e., sorting out and repairing or replacing improper items; return the lot to the manufacturing department or external supplier; sell product at a reduced price; scrap product.

A *type B* procedure is a combined lot and process assessment procedure. The decisions *acceptance* and *rejection* refer both to the individual lot, in the way exemplified for type A procedures, and to the process behind the lot with the following interpretations. *Acceptance*: continue the process of manufacture or delivery in a routine manner without intervention. *Rejection*: intervene into the process, e.g., send the inspection result as an alert to the manufacturing department, launch an inspection of the manufacturing process to detect assignable causes of quality variation, complain to the product provider requiring quality control interventions.

Deming [3] defines statistical quality control (SQC) as the “application of statistical principles and techniques in all stages of design, production, maintenance, and service, directed toward the economic satisfaction of demand”. Statistical methods are used for design, process, product control, both with preventive and monitoring (inspective) interest. Inspective SQC is based on *sampling*. The two most popular sampling based procedures are the **control chart** in **statistical process control** (SPC) and **acceptance sampling** (*statistical product inspection, SPI*) in product control. In SPI, the decision on acceptance or rejection of incoming or outgoing product is based on a sample taken from a lot, or from the process behind the lots.

The subsequent sections analyze the basics of SPI from a historical, practical, and methodological point of view.

Historical Background of Statistical Product Inspection

In the Taylorist production environments growing in the first decades of the twentieth century manufacturing was split up into a number of simple consecutive steps. Quality control was identified with inspection of parts at the end of each production step, terminating in a final inspection of finished product. After World War I, expanding mass production made 100% inspection infeasible and involved sampling techniques. The systematic development of statistical lot inspection schemes started in the 1920s. Early contributions are due to Becker [4] and to a group integrating H.F. Dodge and W.A. Shewhart in the “inspection engineering” department at Bell Telephone Laboratories founded in 1925, see the documentation by Dodge [5, 6]. SPI schemes prospered as the main branch of SQC in theoretical development and in industrial practice until the 1950s. Most popular exponents are the *Dodge–Romig Tables* and the *Military Standard* series, in particular MIL-STD-105 and MIL-STD-414, which originated in inspection schemes of the U.S. armament industry during World War II. At that time, lot inspection had a direct quality regulating purpose, as expressed by Dodge [7]: “The purpose of the inspection is to determine the acceptability of individual lots submitted for inspection, i.e., sorting good lots from bad.”

Industry faced new challenges in the decades after World War II: market saturation, tightening quality demands, legal obligations, consumer and environmental protection, globalization, increasing cost of material and labor, technological progress, increasing complexity of products. Quality control had to rearrange priorities: from product control to design and process control, from inspective to preventive methods. The traditional emphasis on **product inspection** was severely criticized. Deming [8] formulates number 3 of his 14 principles for transformation of western management: “Cease dependence on inspection to achieve quality. Eliminate the need for inspection on a mass basis by building quality into the product in the first place.” The new orientation of industry changed SQC. The field of statistics in quality control was considerably enlarged, exceeding the classical SQC core disciplines SPC and SPI. Statistics can be useful at any of the stages of the quality loop. SPI schemes lost reputation and were even attacked as completely obsolete, see [9] or [8].

However, inspection still has a place in modern product quality control. ISO 9004 [10], mentions inspection as a tool of product control under the issues *receiving inspection planning*, *verification of incoming material and parts*, *completed product verification*. The latter issue explicitly mentions 100% inspection (screening) and sampling inspection of lots. SPI can play a reasonable role if it is designed not as a quality regulating but as an adjunct tool in a product control system based on prevention and on cooperative supplier–customer relations. The specific challenges for statistical lot inspection in modern industrial environments are discussed in more detail by **Acceptance Sampling in Modern Industrial Environments**.

Quality Modeling in Product Inspection

A clear idea of item quality, lot quality, and process quality is indispensable for reasonable SPI.

Consider discrete items (units) $i = 1, 2, \dots$. Each item i is ascribed a characteristic X_i which is influenced by random events during production, transport, and storage. The quality of item i is supposed to be determined exclusively by X_i , the *quality indicator of item i* . A lot is a finite set of items (units) $1, \dots, N$ of finished product considered at a fixed time. Hence the quality of a given lot is fixed and depends exclusively on the quality of the items in the lot. It is quantified by the *lot quality indicator* Q as a function $Q = Q(X_1, \dots, X_N)$ of the quality indicators of the items in the lot. The choice of the lot quality indicator reflects the economic appraisal of the lot: the economic value of the lot should depend on the quality $X_1, \dots, X_N$ of the items in the lot only through the lot quality indicator Q . The most important lot quality indicator is the *lot mean* $Q = 1/N \sum_{i=1}^N X_i = \bar{X}_{\text{lot}}$.

Each lot of items $1, \dots, N$ is assembled from a process of manufacturing, handling, transport, and so on. The process is the original source of variation in the item quality indicators X_i and in the lot quality indicator $Q(X_1, \dots, X_N)$. Hence it is important to conceive an idea of process quality. From a stochastic point of view, process quality is determined by a parameter of the probability distribution of the random item quality indicator X_i . In many cases, the suitable process quality indicator is the expectation $\theta_i = E[X_i]$ of the item quality indicators resulting

from the process, i.e., the average process level. For specific applications, other distribution parameters may be of interest, as for instance **quantiles** in **reliability analysis**.

The following quality models are the most important instances of the above framework in acceptance sampling.

Conforming–nonconforming quality model

Item quality indicators: binary variables $X_i \in \{0, 1\}$ where $X_i = 0$, if item i is conforming, $X_i = 1$, if item i is nonconforming. Lot quality indicator: lot proportion nonconforming $P = 1/N \sum_{i=1}^N X_i = \bar{X}_{\text{lot}}$. Process quality indicator: process proportion nonconforming $\pi = P(X_i = 1)$. Relation between lot and process quality indicator: $\pi = E[P]$, i.e., the expected lot proportion nonconforming equals the process proportion nonconforming.

Nonconformities (discrete measurement) quality model

Item quality indicators: X_i is the number of nonconformities on item i . Lot quality indicator: lot average number of nonconformities $K = 1/N \sum_{i=1}^N X_i = \bar{X}_{\text{lot}}$. Process quality indicator: process average number of nonconformities $\lambda = E[X_i]$. Relation between lot and process quality indicator: $\lambda = E[K]$, i.e., the expected lot average number of nonconformities equals the process average number of nonconformities.

Continuous measurement quality model

Item quality indicators: X_i is a continuous scale measurement, e.g., weight, length, volume, surface, diameter, pressure, tensile strength, and lifetime. Various types of lot and process quality models according to the specific type of application exist, see [11].

Specification range model

A conforming–nonconforming indicator X_i is built from a discrete or continuous measurement X'_i by prescribing a *specification range* $\mathcal{R}$ for X'_i where $X_i = 0$, if $X'_i \in \mathcal{R}$, $X_i = 1$, if $X'_i \notin \mathcal{R}$. Usual specification ranges are $\mathcal{R} = (-\infty; UL]$ with *upper specification limit* UL , $\mathcal{R} = [LL; +\infty)$ with *lower specification limit* LL , $\mathcal{R} = [LL; UL]$ with *two-sided specification limits* LL , UL .

Sampling procedures are classified according to the underlying quality models. *Sampling by attributes* is based on the conforming–nonconforming model or

on the nonconformities quality model. *Sampling by variables* is based on the specification range model. The continuous measurement model is considered in the framework of reliability sampling by **Reliability Sampling**.

Sampling Plans and Sampling Techniques

SPI is implemented according to a *sampling plan*. A *sampling plan* S is defined by the following *design components*:

- (SP1) Sample-size parameters $n_1, n_2, \dots$
- (SP2) A sampling rule which prescribes how to sample.
- (SP3) Sample statistics $t_1, t_2, \dots$ to evaluate the samples.
- (SP4) Decision parameters.
- (SP5) A decision rule which infers a decision (acceptance or rejection) from the sample statistics and the decision parameters.

The simplest instance is a *single sampling plan* $S = (n, A)$, see **Single Sampling by Attributes and by Variables** a single random sample of size n is taken, the sample quality indicators $Y_1 = X_{i_1}, \dots, Y_n = X_{i_n}$ are evaluated by a statistic $t(Y_1, \dots, Y_n)$, the decision is *acceptance* if t lies in the *acceptance region* A , otherwise the decision is *rejection*. Double, multiple and sequential sampling plans allow for several samples to be taken consecutively until a decision is achieved, see **Multiple Sampling Plans**.

The sample should provide unbiased and representative information on the quality of the inspection target (lot or process). This requirement is met by a *random sample*, as described in the sequel.

In a type A procedure, the samples are drawn without replacement from a specified lot of items $1, \dots, N$. Formally, random sampling without replacement from a lot $1, \dots, N$ means that all $\binom{N}{n}n!$ sequences of pairwise different items $i_1, \dots, i_n$ have the same probability to constitute the sample. In practice, a reliable way to guarantee this property proceeds as follows: assign physically numbers 1 through N to the items in the lot, then extract n pairwise different numbers $1 \leq i_1, \dots, i_n \leq N$ from a random number generator as implemented in statistical packages. If the items cannot be numbered explicitly, a variety of provisions can help to exclude

biases and to come close to an ideally random sample, in particular, not to sample from regional or temporal clusters, e.g., from the top or the bottom of the lot, from the first or last items produced.

The sample for a type B procedure may be from a lot as for a type A procedure, or directly from the process without previous assembling into lots. In final inspection, the items can be sampled directly at the end of the manufacturing line. In receiving inspection, the items may be sampled from a continuous stream of just-in-time delivered material or product. Formally, a process random sample has two properties: the item quality indicators $X_{i_1}, \dots, X_{i_n}$ are independent, and the process quality parameter θ remains invariant over the sample. The practical provisions to achieve these properties are the same as in control charting, *see Control Charts, Overview*. In particular, the possibility of shifts in the process quality parameter within the sample should be excluded, and excessive proximity of items should be avoided to exclude **correlation**.

The Operating Characteristic Function of a Sampling Plan

The *operating characteristic function (OC function)* of a sampling plan S is the probability of achieving the decision *acceptance* as a function of the sampling plan parameters (sample sizes, acceptance region), and the quality indicator of the inspection target (lot or process).

The inspection target of a type A procedure is an individual lot of size N with a given value $Q = q$ of the lot quality indicator, e.g., lot proportion nonconforming, or lot average number of nonconformities. In this case, the OC function is the conditional probability

$$L_S(q) = P(\text{accept the lot under } S | Q = q) \quad (1)$$

of acceptance under the condition that the underlying process has provided a lot of quality $Q = q$. The unconditional probability measure P characterizes the underlying process.

The inspection target of a type B procedure is a process characterized by the process quality indicator θ , e.g., process proportion nonconforming, or process average number of nonconformities. In this case, the OC function is the parametric probability

$$L_S(\theta) = P_\theta(\text{accept the lot/process under } S) \quad (2)$$

As a general rule, type B OC functions can be used as an approximation of type A OC functions for large lot size N . For a detailed analysis of OC functions in acceptance sampling, in particular, the relation between process and lot models, *see* [11]. Specific OC functions are discussed in **Single Sampling by Attributes and by Variables, Multiple Sampling Plans, Rectification Sampling Schemes, Continuous Sampling, Reliability Sampling, Variables Sampling under Measurement Error, and Attributes Sampling under Classification Error**.

Protection is a central issue in product assessment: the consumer seeks protection against accepting unsatisfactory quality, the producer wants good quality to be accepted. The protective potential of sampling plans is characterized by properties and quantities related to the OC function. Important quantities are listed and explained by Table 1.

Definition of Quality Levels

Rational product assessment needs a definition of acceptable (satisfactory) and unacceptable (rejectable, unsatisfactory) product quality. Let Γ be the range of the product quality indicator (lot quality indicator $\gamma = q$ or process quality indicator $\gamma = \theta$), e.g., $\Gamma = [0; 1]$ for the proportion nonconforming, $\Gamma = [0; +\infty)$ for the average number of nonconformities or the average measure value in the continuous measurement model. The parameter range is partitioned by $\Gamma = \Gamma_I \cup \Gamma_{II}$ into the disjoint sets Γ_I (acceptable values of the product quality indicator) and Γ_{II} (unacceptable values of the product quality indicator). The separating boundary of Γ_I and Γ_{II} is the *breakeven quality* γ_0 , e.g., the maximum acceptable proportion nonconforming or the maximum acceptable average number of nonconformities.

The choice of a breakeven quality may be based on technical, security or economic reasons. Economic approaches to lot assessment always imply the specification of a breakeven quality. A widely used economic scheme suggested by Moriguti [12] evaluates the decision acceptance and rejection by linear loss functions $C_A(\gamma)$, $C_R(\gamma)$ of the product quality indicator γ . Then the breakeven quality γ_0 is the intersection point of C_A and C_R , *see Economic Sampling Schemes*. Many authors consider the following particular instance of Moriguti's model, *see* [9] or [8]. Let $\gamma = p$ be the lot proportion nonconforming where rejection means rectification of the lot.

Table 1 Characteristic quantities of sampling plans with OC $L(\gamma)$, (lot quality indicator $\gamma = q$ or process quality indicator $\gamma = \theta$)

Notation	Definition
γ_ρ	ρ 100% point defined by $L(\gamma_\rho) \stackrel{!}{=} \rho$, i.e., a quality level γ_ρ is accepted with probability ρ .
CQ	Consumer's quality. A poor quality level from the consumer's point of view which should be accepted with small probability $L(\text{CQ})$ only.
CR	Consumer's risk: the probability $\text{CR} = L(\text{CQ})$ of accepting the quality level CR.
LTPD	Lot tolerance percent defective. A term for the CQ under inspection for lot proportion nonconforming.
LQL	Limiting quality level. Alternative term for CQ.
RQL	Rejectable quality level. Alternative term for CQ.
PQ	Producer's or supplier's quality. Good quality level from the provider's point of view which should be rejected with small probability $1 - L(\text{PQ})$ only.
PR	Producer's or supplier's risk: the probability $\text{PR} = 1 - L(\text{PQ})$ of rejecting the quality level PR.

Let k_1 be the cost of inspecting an item, k the cost of replacing a nonconforming by a conforming item, k_2 the cost from processing a nonconforming item. Then $C_A(p) = (k_2 + k)Np$, $C_R(p) = (k_1 + kp)N$. Hence the breakeven quality is $p_0 = k_1/k_2$.

It is a basic weakness that, except economic schemes like the one described by **Economic Sampling Schemes**, the majority of the popular SPI schemes do not require the explicit specification of breakeven qualities. Instead, the design of such schemes strongly relies on quantities as listed by Table 1. Some of these quantities have some similarity with a breakeven quality, but they are not properly defined as such. See **Acceptance Sampling in Modern Industrial Environments** for a further analysis of the lack of precise definitions of quality levels in customary SPI.

The Design of Sampling Plans

A *design algorithm* allows determining a sampling plan from some specified input parameters. The design algorithm defines a *sampling scheme*. A *sampling table* is a ready-to-use representation of a sampling scheme in the form of a list indexed by values of the input parameters. The input parameters are of the following types: (a) lot descriptors, e.g., lot size; (b) process quality parameters; (c) lot quality parameters; and (d) economic parameters.

An enormous variety of design rules and **sampling schemes** is discussed in literature. We give a brief survey organized by the three major classification issues. For a more extensive treatment see [13].

A design rule is based on an *evaluation index* $I(S, q)$ which evaluates a sampling plan S under lot quality q . Two types of indices have to be distinguished: (a) Statistical indices based on the OC $L_S(q)$ alone. (b) Economic indices integrating cost and profit parameters, e.g., sampling cost, loss or profit incurred from rejecting/accepting a lot, the latter mostly modeled on the basis of Moriguti's [12] linear scheme, see **Economic Sampling Schemes**.

A second classification relates to the assumptions on prior knowledge on product quality, i.e., the assumptions on the distribution function F of the lot quality indicator Q . Three approaches have to be distinguished: (a) no assumptions (b) partial knowledge represented by parameters of the distribution of Q , e.g., the mean, variance, or quantiles. This amounts to assuming F to be a member of a subset Γ of all distribution functions. (c) Complete knowledge about the distribution function of Q (Bayes approach). Assumptions on process quality are particularly relevant for the way of building economic objective functions from economic indices, subsequently exemplified for a loss index $I(S, q)$. Approach (a) leads to the worst case function $I(S) = \max_q I(S, q)$. Approach (b) leads the restricted worst case function $I(S) = \max_{F \in \Gamma} E_F[I(S, Q)]$. Approach (c) leads to the Bayes objective function $I(S) = E_F[I(S, Q)]$. See **Prior Information in Sampling Schemes** for a more detailed analysis of prior knowledge on product quality in SPI schemes.

A third classification relates to the purposes or optimality principles pursued. Three kinds can be distinguished: (a) *Statistical principles* require certain properties of the OC, e.g., prescribing consumer's

risk, consumer's quality, or producer's risk and producer's quality, see Table 1. (b) *Economic principles* require to maximize (minimize) an economic profit (loss) function $I(S)$. See **Economic Sampling Schemes** for an instance of an economic loss function. (c) *Convenience principles* emphasize pragmatic aspects of usability and simplicity of a sampling scheme.

Special SPI Schemes

A large variety of sampling schemes is discussed in literature, see [14] as an extensive reference source. The following schemes are used, usable or useful.

Dodge–Romig plans

Dating back to Dodge's work at Bell Telephone in the late 1920s and through the 1930s. Published in book form in 1941, with a revised edition in 1959, see [7]. Rectification sampling plans of single and double type for lot proportion nonconforming with a mixed economic-statistical design. Strengths: transparent design, prior knowledge on product quality reflected. Weaknesses: no economic motivation for critical design parameters, inadequate for tight quality requirements. See **Rectification Sampling Schemes** for further analysis.

Prescription of two points of the OC function

A simple and intuitive statistical design suggested by Peach-Littauer [15] in 1946. Consumer's quality level CQ, producer's quality level PQ, and upper bounds β for the consumer's risk (CR), α for the producer's risk (PR) are prescribed. Strength: general, only based on the OC. Weaknesses: prior knowledge on lot quality is not reflected, inadequate for tight quality requirements.

Military standard

Originating in sampling tables developed for the US army during World War II by a group of experts headed by H.F. Dodge. Recent versions: MIL-STD-105E [16], issued in 1989 for attributes sampling for lot proportion nonconforming and lot average number of nonconformities, MIL-STD-414 [17] for sampling by variables for lot proportion nonconforming issued in 1957. In revised form, adopted as national standards in many countries and as international standards ISO 2859 and ISO 3951. Strengths: (a) Simple, widely used and accepted by industry. (b) A proper *sampling system*, including single and

multiple sampling, with three degrees of inspection severity and rules for *switching* among these. Weaknesses: (a) No consistent design principle. (b) Vague definition of the quality target AQL (*acceptable quality level*). (c) Prior knowledge on lot quality not reflected. (d) Inadequate for tight quality requirements. See **Attributes Sampling Schemes in International Standards; Variables Sampling Schemes in International Standards** for further analysis.

α -Optimal Sampling Scheme

A transparent economic design suggested by Collani [18] based on Moriguti's [12] linear loss model. Strengths: Consistent economic design, incomplete prior knowledge on lot quality is considered. Weakness: Hitherto no experience from industrial implementation recorded. See **Economic Sampling Schemes** for further analysis.

Joyce Orsini's rules

An economic convenience approach propagated by Deming [8]. Strengths: simple, admitting for no-sampling alternatives, based on an economic and process model. Weaknesses: no thorough design, no experience from industrial implementation recorded.

References

- [1] International Organization for Standardization. (2000). *ISO 9004: Quality Management Systems – Fundamentals and Vocabulary*, International Organization for Standardization, Geneva.
- [2] International Organization for Standardization. (1985). *ISO 2859-2: Sampling Procedures for Inspection by Attributes - Part 2: Sampling Plans Indexed by Limiting Quality (LQ) for Isolated Lot Inspection*, International Organization for Standardization, Geneva.
- [3] Deming W.E. (1971). Some statistical logic in the management of quality, in *Proceedings of the All India Conference on Quality Control*, New Delhi, pp. 98–119.
- [4] Becker, R., Plaut, H. & Runge, I. (1927). *Anwendungen der Mathematischen Statistik auf Probleme der Massenfabrication*, Springer, Berlin.
- [5] Dodge, H.F. (1969). Notes on the evolution of acceptance sampling plans. Parts I,II,III, *Journal of Quality Technology* **1**, 77–88, 155–162, 225–232.
- [6] Dodge, H.F. (1970). Notes on the evolution of acceptance sampling Plans. Part IV, *Journal of Quality Technology* **2**, 1–8.
- [7] Dodge, H.F. & Romig, H.G. (1959). *Sampling Inspection Tables*, 2nd Edition, John Wiley & Sons, New York.
- [8] Deming, W.E. (1986). *Out of the Crisis*, MIT Press, Cambridge.

-
- [9] Feigenbaum, A.V. (1961). *Total Quality Control*, McGraw-Hill, New York.
- [10] International Organization for Standardization. (2000). *ISO 9004: Quality Management Systems – Guidelines for Performance Improvements*, International Organization for Standardization, Geneva.
- [11] Göb, R. (2001). Methodological foundation of statistical lot inspection, in *Frontiers in Statistical Quality Control*, H.-J. Lenz & P.-Th. Wilrich, eds, Physica Verlag, Heidelberg, New York, pp. 3–24.
- [12] Moriguti, S. (1955). Notes on sampling inspection plans. Reports of Statistical Applications Research, *Japanese Union of Scientists and Engineers* **3**, 99–121.
- [13] Arnold, B.F., Göb, R. (2003). Sample method and quality control, in *Probability and Statistics, Encyclopedia of Life Support Systems (EOLSS), Developed Under the Auspices of the UNESCO*, R. Viertl, ed, Eolss Publishers, Oxford.
- [14] Schilling, E.G. (1982). *Acceptance Sampling in Quality Control*, Marcel Dekker, New York, Basel.
- [15] Peach, P. & Littauer, S.B. (1946). A note on sampling inspection, *The Annals of Mathematical Statistics* **17**(1), 81–84.
- [16] United States Department of Defense. (1989). *Military Standard, Sampling Procedures and Tables for Inspection by Attributes (MIL-STD-105E)*, US Government Printing Office, Washington DC.
- [17] United States Department of Defense. (1957). *Military Standard, Sampling Procedures and Tables for Inspection by Variables for Percent Defective (MIL-STD-414)*, US Government Printing Office, Washington DC.
- [18] von Collani, E. (1986). The α -optimal sampling scheme, *Journal of Quality Technology* **18**, 63–66.

RAINER GÖB

Single Sampling by Attributes and by Variables

Introduction

Single sampling plans are the simplest instance of a statistical **product inspection** (SPI) procedure. A single sampling plan is defined by a pair (n, c) consisting of the **sample size** n and the *acceptance number* c , and a *sample statistic* t . The decision rule of a single sampling plan prescribes the following steps: (a) n items are sampled at random; (b) the value of the sample statistic t is calculated from the elements of the sample; and (c) *acceptance* is warranted if $t \leq c$, otherwise the decision is *rejection*.

The study of single sampling plans is illustrative for the understanding of **acceptance sampling**. Single sampling exemplifies all basic features of acceptance sampling in a particularly simple and transparent instance. From theoretical optimality criteria, e.g., expected sampling amount, more refined sampling algorithms with several iterative sample runs like double, multiple, and **sequential sampling** plans (see **Multiple Sampling Plans**) are sometimes superior to single sampling. However, these optimality studies ignore administrative barriers and costs which oppose iterative sampling on shop floor. Simplicity and administrative efficiency have made single sampling plans the most popular acceptance sampling technique by far in industrial practice.

The subsequent explanations consider single sampling plans in general. The special case of single sampling plans under rectifying inspection is considered by **Rectification Sampling Schemes** in detail.

Item Quality, Lot Quality, and Process Quality

Consider discrete items $i = 1, 2, \dots$. The product quality model established by **Sampling Inspection of Products** distinguishes three types of quality indicators: (a) The random *quality indicator* X_i associated with each item i . (b) The *lot quality indicator* which is a function of the quality indicators $X_1, \dots, X_N$ of the items $1, \dots, N$ in the

lot. (c) The quality indicator of the lot generating process. The present article restricts attention to the following instances of this product quality model.

Conforming–nonconforming quality model.

Item quality indicators: binary variables $X_i \in \{0, 1\}$ where $X_i = 0$, if item i is conforming, $X_i = 1$, if item i is nonconforming. Lot quality indicator: lot proportion nonconforming $P = 1/N \sum_{i=1}^N X_i = \bar{X}_{\text{lot}}$. Process quality indicator: process proportion nonconforming $\pi = P(X_i = 1)$.

Nonconformities (discrete measurement) quality model.

Item quality indicators: X_i is the number of nonconformities on item i . Lot quality indicator: lot average number of nonconformities $K = \frac{1}{N} \sum_{i=1}^N X_i = \bar{X}_{\text{lot}}$. Process quality indicator: process average number of nonconformities $\lambda = E[X_i]$.

Specification range model.

A conforming–nonconforming indicator X_i is built from a discrete or continuous measurement X'_i by prescribing a *specification range* $\mathcal{R}$ for X'_i where $X_i = 0$, if $X'_i \in \mathcal{R}$, $X_i = 1$, if $X'_i \notin \mathcal{R}$. Usual specification ranges are $\mathcal{R} = (-\infty; UL]$ with *upper specification limit* UL , $\mathcal{R} = [LL; +\infty)$ with *lower specification limit* LL , $\mathcal{R} = [LL; UL]$ with *two-sided specification limits* LL, UL .

Sampling procedures are classified according to the underlying quality models. *Sampling by attributes* is based on the conforming–nonconforming model or on the nonconformities quality model. *Sampling by variables* is based on the specification range model.

Single Sampling Plans and their Operating Characteristic Function

The data basis for decision by a single sampling plan (n, c) is a single *sample* of n items $\iota_1, \dots, \iota_n$ with a corresponding sequence of item quality indicators $X_{\iota_1}, \dots, X_{\iota_n}$. The sample is evaluated by the sample statistic $t = t(X_{\iota_1}, \dots, X_{\iota_n})$. The sample should provide unbiased and representative information on the quality of the inspection target (lot or process). This requirement is met by a *random sample*. Random sampling from lots and processes is discussed by **Sampling Inspection of Products**.

2 Single Sampling by Attributes and by Variables

Table 1 Sample statistics and OC functions of single sampling plans (n, c) by attributes

Sampling for proportion nonconforming		
Sample statistic	$t = \sum_{j=1}^n X_{t_j}$ number of nonconforming items in the sample $X_{t_1}, \dots, X_{t_n}$	
Type A	Inspection target:	lot proportion nonconforming p
	Hypergeometric OC:	$P(t \leq c) = L_{N,n,c}(p) = \sum_{m=0}^c \frac{\binom{Np}{m} \binom{N(1-p)}{n-m}}{\binom{N}{n}}$
Type B	Inspection target:	process proportion nonconforming π
	Binomial OC:	$P(t \leq c) = L_{n,c}(\pi) = \sum_{m=0}^c \binom{n}{m} \pi^m (1-\pi)^{n-m}$
Poisson approximation	Inspection target:	lot or process proportion nonconforming $q = p$ or $q = \pi$
	Poisson OC:	$P(t \leq c) = L_{n,c}^*(q) = \sum_{m=0}^c \frac{(nq)^m}{m!} \exp(-nq)$
Sampling for average number of nonconformities		
Sample statistic	$t = \sum_{j=1}^n X_{t_j}$ number of nonconformities on items in the sample $X_{t_1}, \dots, X_{t_n}$	
Type A	Inspection target:	lot average number of nonconformities k
	Negative hypergeometric OC:	$P(t \leq c) = \sum_{m=0}^c \frac{\binom{n+m-1}{n-1} \binom{N(1+k)-n-m-1}{N-n-1}}{\binom{N(1+k)-1}{N-1}}$
Type B	Inspection target:	process average number of nonconformities λ
	Poisson OC:	$P(t \leq c) = L_{n,c}^*(\lambda) = \sum_{m=0}^c \frac{(n\lambda)^m}{m!} \exp(-n\lambda)$

The *operating characteristic Function (OC function)* of a sampling plan (n, c) is the probability of achieving the decision *acceptance* as a function of the sampling plan parameters n, c , and the parameters of the inspection target (lot or process).

Type A product inspection is distinguished from type B product inspection, *see Sampling Inspection of Products*. The inspection target of a type A procedure is an individual lot of size N with a given value $Q = q$ of the lot quality indicator, e.g., lot proportion nonconforming, or lot average number of nonconformities. In this case, the OC function is the conditional probability

$$L_{N,n,c}(q) = P(\text{accept} | Q = q) = P(t \leq c | Q = q) \quad (1)$$

where the unconditional probability measure P characterizes the underlying process.

The inspection target of a type B procedure is the product generating process characterized by the process quality indicator θ , e.g., process proportion

nonconforming, or process average number of nonconformities. In this case, the OC function is the parametric probability

$$L_{n,c}(\theta) = P_\theta(\text{accept}) = P_\theta(t \leq c) \quad (2)$$

As a general rule, type B OC functions can be used as an approximation of type A OC functions for large lot size N . For a detailed analysis of OC functions in acceptance sampling, in particular, the relation between process and lot models, see [1].

The sample statistics and OC functions of single sampling plans by attributes are listed in Table 1. Figure 1 displays the OC functions for plans for lot proportion nonconforming and for lot average number of defects. The plans are taken from MIL-STD-105E, normal inspection, inspection level II, AOQ of 1% for proportion nonconforming, 10% for average number of defects. The examples show that even for moderate lot size the difference between the type A and type B OCs is very small. Sampling by variables is discussed in the next section.

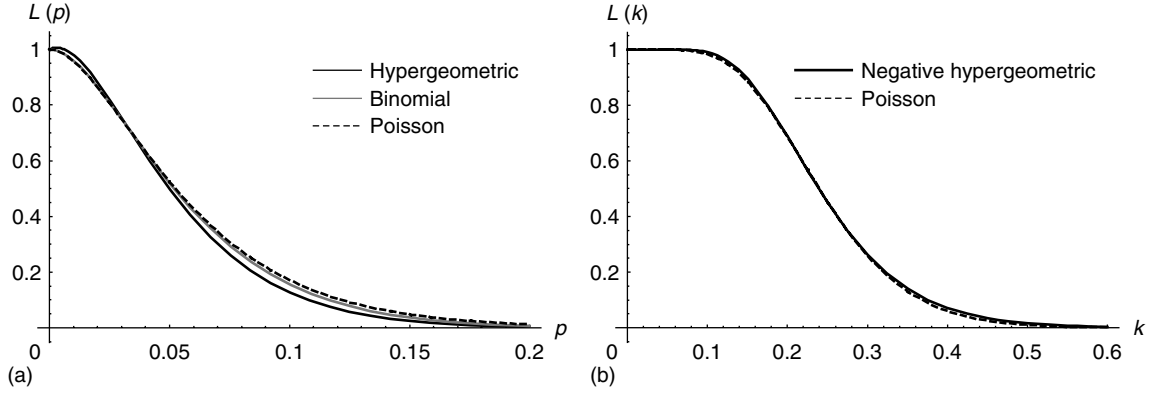


Figure 1 OC functions of the single sampling plan $(n, c) = (32, 1)$ for proportion nonconforming p and $(n, c) = (32, 7)$ for average number of defects k . For type A plans the lot size is $N = 160$

Table 2 Process proportion nonconforming π in the model of sampling by variables. Φ is the distribution function of the standard normal distribution $N(0, 1)$

Specifications	X' with cdf F	X' with normal distribution $N(\mu, \sigma^2)$
One-sided upper UL	$\pi_U = F(X' > UL)$	$\pi_U = 1 - \Phi\left(\frac{UL - \mu}{\sigma}\right) = \Phi\left(\frac{\mu - UL}{\sigma}\right)$
	$\pi_{L,U} = \pi_L + \pi_U$	$\pi_{L,U} = \Psi_\sigma\left(\mu - \frac{LL + UL}{2}\right)$ where
Two-sided $LL < UL$	$= F(X' < LL)$ $+ F(X' > UL)$	$\Psi_\sigma(x) = \Phi\left(\frac{1}{\sigma}\left(\frac{LL - UL}{2} - x\right)\right)$ $+ \Phi\left(\frac{1}{\sigma}\left(x + \frac{LL - UL}{2}\right)\right)$

Single Sampling by Variables

Sampling by variables is based on the specification range model exposed above. Intrinsically, sampling by variables is a type B procedure where the inspection target is the process proportion nonconforming π . In the specification range model, π can be expressed with the cumulative distribution function (cdf) of the underlying continuous measurement quality indicators X' in the way expressed by Table 2. Accordingly, the implementation of sampling by variables requires that at least the shape of the cdf of the X'_i from the process is stable and known. This requires a good and reliable knowledge about the properties of the process. Without such knowledge, variables plans should not be used, see the remarks on variables sampling by [2].

In a large variety of cases, a **normal distribution** $N(\mu, \sigma^2)$ can serve at least as a good approximation of the distribution of X'_i . Thus the design of sampling plans by variables for this situation is particularly interesting. Table 2 shows the functional relationship between the process proportion nonconforming π , and process mean μ and process variance σ^2 . Hence statistical inference on π is equivalent to statistical inference on μ and σ^2 based on an independent random sample of measurements $X'_{t_1}, \dots, X'_{t_n}$. The inference assumes that μ and σ^2 are constant for the lot or process segment under consideration, and hence in particular for the sample. Unsatisfactory process proportion nonconforming π can result either from an inappropriate value of μ or from an inappropriate value of σ^2 . The sample statistics used for a single sampling plan (n, c) are based on estimating μ by the sample mean $\bar{X}'_{\text{sample}}$ and σ^2 , if it is unknown, by the

4 Single Sampling by Attributes and by Variables

Table 3 Sample statistics and OC functions of single sampling plans (n, c) by variables, process variance σ^2 known. Φ is the distribution function of the standard normal distribution $N(0, 1)$, $\bar{X}'_{\text{sample}} = \frac{1}{n} \sum_{j=1}^n X'_{t_j}$ is the sample mean

One-sided upper specification limit UL	
Sample statistic	$t = \frac{\bar{X}'_{\text{sample}} - UL}{\sigma} \sqrt{n}$
OC function	$L_{n,c}(\pi_U) = \Phi\left(c + \Phi^{-1}(1 - \pi_U)\sqrt{n}\right) = \Phi\left(c - \Phi^{-1}(\pi_U)\sqrt{n}\right)$
Two-sided specification limits $LL < UL$	
Sample statistic	$t = \left \frac{\bar{X}'_{\text{sample}} - \frac{LL+UL}{2}}{\sigma} \sqrt{n} \right $
OC function	$L_{n,c}(\pi_{L,U}) = \Phi\left(c + \frac{\Psi_{\sigma}^{-1}(\pi_{L,U})}{\sigma} \sqrt{n}\right) - \Phi\left(-c + \frac{\Psi_{\sigma}^{-1}(\pi_{L,U})}{\sigma} \sqrt{n}\right)$ where $\Psi_{\sigma}^{-1}(\pi_{L,U})$ is the positive inverse of the function Ψ_{σ} defined by Table 2

Table 4 Sample statistics and OC functions of single sampling plans (n, c) by variables under unknown process variance σ^2 estimated by the sample variance $S^2 = \frac{1}{n-1} \sum_{j=1}^n (X_{t_j} - \bar{X}_{\text{sample}})^2$. F_{δ} is the distribution function of the noncentral t distribution with $n - 1$ degrees of freedom and noncentrality parameter δ

One-sided upper specification limit UL	
Sample statistic	$t = \frac{\bar{X}'_{\text{sample}} - UL}{S} \sqrt{n}$
OC function	$L_{n,c}(\pi_U) = F_{\Phi^{-1}(\pi_U)\sqrt{n}}(c) = 1 - F_{\Phi^{-1}(1-\pi_U)\sqrt{n}}(-c)$
Two-sided specification limits $LL < UL$	
Sample statistic	$t = \left \frac{\bar{X}'_{\text{sample}} - \frac{LL+UL}{2}}{S} \sqrt{n} \right $
OC function	$L_{n,c}(\pi_{L,U}) = F_{\delta}(c) - F_{\delta}(-c)$ where $\delta = \frac{\Psi_{\sigma}^{-1}(\pi_{L,U})}{\sigma} \sqrt{n}$, and where $\Psi_{\sigma}^{-1}(\pi_{L,U})$ is the positive inverse of the function Ψ_{σ} defined by Table 2

sample variance S^2 , see Tables 3 and 4. The sample statistics and OC functions for single sampling plans by variables are displayed by Table 3 for the situation of known process variance σ^2 , and by Table 4 for the situation of unknown process variance σ^2 . The tables consider the case of a one-sided upper limit UL and two-sided limits $LL < UL$. The case of a one-sided lower limit LL can be reduced to the case of an upper limit by considering measurements $Y'_i = -X'_i$ and the upper limit $UL = -LL$.

A type A interpretation of variables sampling, to inspect for lot proportion nonconforming, requires a considerable caution. Two aspects are intuitively evident: there is no functional relationship between the process mean μ and lot proportion nonconforming,

and the test statistics suggested by Tables 3 and 4 do not provide suitable information on lot proportion nonconforming. In fact, the following phenomena can occur under a type A use of sampling by variables. (a) Rejection of the lot although the actual sample contains no nonconforming items. (b) Rejection of lots containing 100% conforming items. (c) In case of two-sided specification limits: acceptance of lots containing 100% nonconforming items under 100% inspection. These phenomena have been recorded and analyzed by many authors, e.g. [3–5]. If the focus is on the proportion nonconforming in individual single lots, sampling by attributes should rather be used instead of sampling by variables. To eliminate the specific problems of the two-sided case, various

authors have suggested to use two separate one-sided variables sampling plans instead of a two-sided variables plan, see [3].

Designing a Single Sampling Plan as a Test of Significance

Before implementing **inspection policies**, *acceptable (satisfactory)* and *unacceptable (unsatisfactory)* quality levels should be defined. Consider a lot or process quality indicator of the type *The smaller the better*, e.g., lot or process proportion nonconforming, lot or process average number of defects. In this case, a so-called breakeven quality level q_0 distinguishes acceptable quality with $q < q_0$ from unacceptable quality with $q > q_0$. The motivation of a breakeven quality by an economic model is discussed in **Sampling Inspection of Products and Economic Sampling Schemes**.

With respect to a breakeven quality, decision on acceptance and rejection amounts to a test of the hypothesis $H: q \leq q_0$ against the alternative $K: q > q_0$. The ideal test with $L(q) = 1$ for $q < q_0$ and $L(q) = 0$ for $q > q_0$ can be achieved by 100% inspection only. Sampling inspection involves error, only the probabilities of erroneous decisions can be controlled. In industrial practice, submission of

acceptable quality can be considered as the regular case. Hence, the test of H against K can be considered as a test of significance where only the probability of erroneous rejection of H is controlled by requiring $1 - L(q) \leq \alpha$ for $q \leq q_0$. Although this design approach determines only one of the two sampling plan parameters it can be useful in specific situations.

Modern quality requirements are very strict. **Six Sigma** approaches discuss proportion nonconforming values down to a few parts per million. For such situations, only acceptance number $c = 0$ is tolerable. The sample size of acceptance number zero plans can be calculated from the requirement $1 - L(q) \leq \alpha$ for $q \leq q_0$.

Inspection for lot proportion nonconforming p .

Let p_0 be the breakeven quality. For a large lot, the hypergeometric OC $L_{N,n,c}(p)$ can be approximated by the binomial OC $L_{n,c}(p)$, see Table 1. The requirement $1 - L_{n,0}(p_0) \leq \alpha$ leads to the sample size $n = \ln(1 - \alpha) / \ln(1 - p_0)$.

Inspection for lot average number of defects k .

Let k_0 be the breakeven quality. For a large lot, the negative hypergeometric OC $L_{N,n,c}(k)$ can be approximated by the Poisson OC $L_{n,c}(k)$, see Table 1. The requirement $1 - L_{n,0}(k_0) \leq \alpha$ leads to the sample size $n = -\ln(1 - \alpha) / k_0$.

Table 5 Variables sampling plans (n, c) with prescribed values $L_{n,c}(\pi_1) \geq 1 - \alpha$, $L_{n,c}(\pi_2) \leq \beta$. Φ is the distribution function of the standard normal distribution $N(0, 1)$

One-sided upper specification limit UL , process variance σ^2 known	
Sample size n	Positive integer nearest to $\left(\frac{\Phi^{-1}(1 - \alpha) - \Phi^{-1}(\beta)}{\Phi^{-1}(1 - \pi_1) - \Phi^{-1}(1 - \pi_2)} \right)^2$
Acceptance number c	$\frac{\Phi^{-1}(1 - \pi_1)\Phi^{-1}(\beta) - \Phi^{-1}(1 - \pi_2)\Phi^{-1}(1 - \alpha)}{\Phi^{-1}(1 - \pi_1) - \Phi^{-1}(1 - \pi_2)}$
Two-sided specification limits $LL < UL$, process variance σ^2 known (approximate solution)	
Sample size n	Positive integer nearest to $\left(\frac{\Phi^{-1}(1 - \alpha) - \Phi^{-1}(\beta)}{\Phi^{-1}(1 - \pi_1) - \Phi^{-1}(1 - \pi_2)} \right)^2$
Acceptance number c	$\frac{UL - LL}{2\sigma} \sqrt{n} + \frac{\Phi^{-1}(1 - \pi_1)\Phi^{-1}(\beta) - \Phi^{-1}(1 - \pi_2)\Phi^{-1}(1 - \alpha)}{\Phi^{-1}(1 - \pi_1) - \Phi^{-1}(1 - \pi_2)}$
One-sided upper specification limit UL , process variance σ^2 unknown (approximate solution)	
Sample size n	Positive integer nearest to $\frac{4\Phi^{-1}(1 - \alpha)^2 \left(1 + \frac{1}{8}(\Phi^{-1}(1 - \alpha) + \Phi^{-1}(\beta))^2 \right)}{(\Phi^{-1}(1 - \pi_1) - \Phi^{-1}(1 - \pi_2))^2}$
Acceptance number c	$\left(\Phi^{-1}(1 - \alpha) + \Phi^{-1}(\beta) \right) \frac{\sqrt{n}}{2}$

Design by Prescribing Values of the OC Function

Statistical designs of sampling plans consist in prescribing certain properties of the OC function. Important OC-related characteristic quantities of sampling plans like producer's quality (PQ), producer's risk (PR), consumer's quality (CQ), and consumer's risk (CR) are discussed in **Sampling Inspection of Products**. A simple and intuitive statistical design by prescribing two points of the OC function was suggested by Peach and Littauer [6] in 1946. A good quality level q_1 should be accepted with a high-probability $L(q_1) \geq 1 - \alpha$. A poor quality level q_2 should be accepted with a low-probability $L(q_2) \leq \beta$. Often, $q_1 = \text{PQ}$ is identified as the producer's quality, and the requirement $L(q_1) \geq 1 - \alpha$ imposes the upper limit $1 - L(\text{PQ}) \leq \alpha$ on the PR. $q_2 = \text{CQ}$ is identified as the consumer's quality, and the requirement $L(q_2) \leq \beta$ imposes the upper limit $L(\text{CQ}) \leq \beta$ on the CR. The two points design is made unique by requiring to choose the sampling plan (n, c) with the smallest sample size n whose OC satisfies both the constraints.

For single sampling by attributes, the acceptance number c is an integer. The two points design can be determined by a simple iterative search algorithm. An efficient algorithm is suggested by Peach and Littauer [6]. In case of single sampling by variables, the solution can be given explicitly, see Table 5. A derivation is given by Uhlmann [7]. The table considers one-sided and two-sided plans for process variance σ^2 known, and one-sided plans for process variance σ^2 unknown. The situation is considerably

involved for two-sided plans under unknown process variance σ^2 where often the application of two one-sided plans is recommended. For the latter case, [7] derives an approximate solution for the two points design.

References

- [1] Göb, R. (2001). Methodological foundation of statistical lot inspection, in *Frontiers in Statistical Quality Control*, H.-J. Lenz & P.-Th. Wilrich, eds, Physica-Verlag, Heidelberg, New York, Vol. 6, pp. 3–24.
- [2] Juran, J.M. & Godfrey, A.B. (1988). *Juran's Quality Handbook*, 5th Edition, McGraw-Hill, New York.
- [3] Schilling, E.G. (1982). *Acceptance Sampling in Quality Control*, Marcel Dekker, New York, Basel.
- [4] von Collani, E. (1990). ANSI/ASQC Z1.4 versus ANSI/ASQC Z1.9, *Economic Quality Control* **5**, 60–64.
- [5] Göb, R. (1996). An elementary model of statistical lot inspection and its application to sampling by variables, *Metrika* **44**, 135–163.
- [6] Peach, P. & Littauer, S.B. (1946). A note on sampling inspection, *The Annals of Mathematical Statistics* **17**(1), 81–84.
- [7] Uhlmann, W. (1982). *Statistische Qualitätskontrolle*, B. G. Teubner, Stuttgart.

Further Reading

Dodge, H.F. & Romig, H.G. (1959). *Sampling Inspection Tables*, 2nd Edition, John Wiley & Sons, New York.

ELART VON COLLANI, RAINER GÖB AND
KODAKANALLUR K. SURESH

Multiple Sampling Plans

Introduction

Under a **single sampling plan** (*see Single Sampling by Attributes and by Variables*), all items for inspection are sampled in one go, and the decision of acceptance or rejection is based on the evidence from this unique sample. From an intuitive point of view, there seems to be an opportunity to reduce the total inspection amount, if lot quality is either very good or very bad: Sample and inspect a few items; if the sample quality is very good accept, if it is very bad, reject; if the sample quality is doubtful, sample a few further items. Continue with this procedure, until the quality of the cumulative sample clearly indicates acceptance or clearly indicates rejection. If lot quality is on one of the extreme sides, the total sampling amount required for a sufficiently reliable decision should be smaller than under a single sampling plan.

The above idea leads to so-called iterative or multistage sampling plans with the following two special cases: (a) A *multiple m -stage sampling plan*, where an upper limit of m successive samples is imposed. The simplest case is a *double sampling plan* (two-stage sampling plan) with a maximum of two sampling steps. (b) A *sequential sampling plan*, where sampling may continue until 100% inspection of the lot.

The idea of iterative sampling is old and had already appeared in Dodge's technical reports at Bell Telephone Laboratories in the 1920s and 1930s, *see* [1]. Double and multiple sampling plans are provided by the traditional industrial standards. MIL-STD-105E provides double and seven-stage sampling plans (*see Attributes Sampling Schemes in International Standards*). ISO 2859-1 contains double and five-stage plans (*see Attributes Sampling Schemes in International Standards*). The Dodge–Romig tables

contain double sampling plans (*see Rectification Sampling Schemes*). Sequential sampling plans are provided by ISO 8422 and ISO 8423 (*see Sampling in Industrial Standards*).

Structure and Operation of Multiple Sampling Plans

In principle, multiple sampling plans can be devised for any of the discrete item quality models of statistical process inspection (SPI) (statistical lot inspection) (*see Sampling Inspection of Products* for a survey). As most established schemes of multiple sampling do, we restrict attention to the customary attributes quality models described by Table 1. For a formal exposition of these models *see Sampling Inspection of Products*. Sequential sampling procedures under the specification range model are provided by the standard ISO 8423 (*see Sampling in Industrial Standards*). We consider multiple sampling plans in a type A framework, where the decision target is a finite lot of items $1, \dots, N$. The methods may also be used for type B purposes where sampling and decision is directed toward the lot generating process. *See Sampling Inspection of Products* for the distinction of type A and type B procedures.

A multiple m -stage sampling plan is defined by a triple $(\mathbf{n}, \mathbf{a}, \mathbf{r})$, where $\mathbf{n} = (n_1, \dots, n_m)$ is a sequence of **sample sizes** $n_1, \dots, n_m \geq 1$, $\mathbf{a} = (a_1, \dots, a_m)$ is a sequence of integer *acceptance numbers* $0 \leq a_1 \leq \dots \leq a_m$, $\mathbf{r} = (r_1, \dots, r_m)$ is a sequence of integer *rejection numbers* $1 \leq r_1 \leq \dots \leq r_m$. The m parameters relate to the m stages of the sampling plan as exposed by Table 2. In most **sampling schemes**, e.g., in MIL-STD-105E and ISO 2859, the differences $r_l - a_l, l = 1, \dots, m - 1$, of the rejection and acceptance number have few variation in l .

The sampling and decision algorithm of a multiple m -stage sampling plan is explained by Table 3

Table 1 Attributes quality models

	Conforming–nonconforming	Nonconformities
Item quality indicator	Item is conforming (nondefective) or nonconforming (defective)	Number of nonconformities (defects) on the item
Lot quality indicator	Lot proportion nonconforming p	Lot average number of nonconformities k
Sample statistic t	Number of nonconforming items in the sample	Number of nonconformities on items in the sample

2 Multiple Sampling Plans

Table 2 The parameters of a multiple m -stage sampling plan for the inspection of a lot of size N

Stage	Sample size	Cumulative sample size	Acceptance number	Rejection number
1	n_1	$n_1^{(c)} = n_1$	a_1	$r_1, r_1 \geq a_1 + 2$
2	n_2	$n_2^{(c)} = n_1 + n_2$	a_2	$r_2, r_2 \geq a_2 + 2$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
m	n_m	$n_m^{(c)} = n_1 + \cdots + n_m$	a_m	$r_m, r_m = a_m + 1$
$0 \leq a_1 \leq a_2 \leq \cdots \leq a_m, \quad 1 \leq r_1 \leq r_2 \leq \cdots \leq r_m, \quad n_1 + \cdots + n_m \leq N$				

Table 3 The sampling and decision algorithm of a multiple m -stage sampling plan with parameters given by Table 2

Stage	Operation	Decision algorithm
1	Sample n_1 items Calculate sample statistic t_1	$t_1 \leq a_1 \implies$ terminate, accept $t_1 \geq r_1 \implies$ terminate, reject $a_1 + 1 \leq t_1 \leq r_1 - 1 \implies$ continue with stage 2
2	Sample n_2 items Calculate sample statistic t_2	$t_1 + t_2 \leq a_2 \implies$ terminate, accept $t_1 + t_2 \geq r_2 \implies$ terminate, reject $a_2 + 1 \leq t_1 + t_2 \leq r_2 - 1 \implies$ continue with stage 3
$\vdots$	$\vdots$	$\vdots$
m	Sample n_m items Calculate sample statistic t_m	$t_1 + \cdots + t_m \leq a_m \implies$ terminate, accept $t_1 + \cdots + t_m \geq r_m \implies$ terminate, reject

where the sample statistics t_l conform to the quality models as explained by Table 1. Because of $a_m = r_m - 1$, the algorithm terminates with certainty at stage m , either with acceptance or with rejection.

In most cases, it is not the proper act of sampling in the sense of selecting items from the lot, which involves costs, but rather the act of inspecting the items to calculate the test statistic (number of nonconforming items or total number of nonconformities in the sample). Inspection costs can be reduced by *curtailed inspection*. Two curtailment techniques exist: (a) Total curtailment: At each stage l , the items in the sample are successively inspected, and as soon as the cumulative number $t_1 + \cdots + t_{l-1} + t_{l,j}$ of nonconforming items or the number of nonconformities found until the j th inspected item equals r_l , inspection is stopped and the lot is rejected. (b) Second-stage curtailment: Starting with curtailment in stage 2 only, so as to obtain an unbiased estimator of the lot quality indicator from the first stage.

Performance Measures of Multiple Sampling Plans

The OC (operating characteristic) function of a type A m -stage sampling plan (n, a, r) is the conditional probability $L_{n,a,r}(q) = P(\text{accept the lot} | \text{lot quality indicator} = q)$ of accepting the lot under the condition that the lot quality indicator (lot proportion nonconforming or lot average number of defects) adopts a value q (see **Sampling Inspection of Products** for a general introduction). Table 4 explains and displays the OC function $L_{n,a,r}(p)$ as a function of the lot proportion nonconforming p under the conforming–nonconforming quality model, and the OC function $L_{n,a,r}(k)$ as a function of the lot average number of nonconformities k under the nonconformities quality model. The OC functions and further performance measures, see below, are expressed by means of stage l probabilities $A_{n,a,r,l}(q)$, $R_{n,a,r,l}(q)$, $T_{n,a,r,l}(q)$ as explained by Table 4. From a type A view, the probabilities

Table 4 Stage l probabilities and OC functions of m -stage sampling plans (n, a, r)

Generic formulas	
$A_{n,a,r,l}(q)$	$A_{n,a,r,l}(q) = P(\text{plan } (n, a, r) \text{ accepts in stage } l \mid \text{lot quality indicator} = q) =$ $\sum_{\substack{0 \leq j_1 \leq n_1, \dots, 0 \leq j_l \leq n_l \\ a_i + 1 \leq j_1 + \dots + j_i \leq r_i - 1, i=1, \dots, l-1 \\ j_1 + \dots + j_l \leq a_l}} \rho_{n,a,r,j_1, \dots, j_l}(q)$
$R_{n,a,r,l}(q)$	$R_{n,a,r,l}(q) = P(\text{plan } (n, a, r) \text{ rejects in stage } l \mid \text{lot quality indicator} = q) =$ $\sum_{\substack{0 \leq j_1 \leq n_1, \dots, 0 \leq j_l \leq n_l \\ a_i + 1 \leq j_1 + \dots + j_i \leq r_i - 1, i=1, \dots, l-1 \\ j_1 + \dots + j_l \geq r_l}} \rho_{n,a,r,j_1, \dots, j_l}(q)$
$T_{n,a,r,l}(q)$	$T_{n,a,r,l}(q) = P(\text{plan } (n, a, r) \text{ terminates in stage } l \mid \text{lot quality indicator} = q) =$ $A_{n,a,r,l}(q) + R_{n,a,r,l}(q)$
OC function	$L_{n,a,r}(q) = P(\text{sampling plan } (n, a, r) \text{ accepts} \mid \text{lot quality indicator} = q) =$ $A_{n,a,r,1}(q) + \dots + A_{n,a,r,m}(q)$
ASN function	$ASN_{n,a,r}(q) = E[\text{random cumulative sample size} \mid \text{lot quality indicator} = q] =$ $\sum_{l=1}^m n_l \sum_{j=l}^m T_{n,a,r,j}(q)$
m -stage sampling plan (n, a, r) for lot proportion nonconforming p	
Hypergeometric	$\rho_{n,a,r,j_1, \dots, j_l}(p) = \prod_{r=1}^l \frac{\binom{Np - (j_1 + \dots + j_{r-1})}{j_r} \binom{N - (n_1 + \dots + n_{r-1}) - (Np - (j_1 + \dots + j_{r-1}))}{n_r - j_r}}{\binom{N - (n_1 + \dots + n_{r-1})}{n_r - j_r}}$
Binomial	$\rho_{n,a,r,j_1, \dots, j_l}(p) = \binom{n_1}{j_1} \dots \binom{n_l}{j_l} p^{j_1 + \dots + j_l} (1-p)^{n_1 + \dots + n_l - j_1 - \dots - j_l}$
Poisson	$\rho_{n,a,r,j_1, \dots, j_l}(p) = \frac{n_1^{j_1}}{j_1!} \dots \frac{n_l^{j_l}}{j_l!} p^{j_1 + \dots + j_l} \exp(-(n_1 + \dots + n_l)p)$
m -stage sampling plan (n, a, r) for lot average number of nonconformities k	
Poisson	$\rho_{n,a,r,j_1, \dots, j_l}(k) = \frac{n_1^{j_1}}{j_1!} \dots \frac{n_l^{j_l}}{j_l!} k^{j_1 + \dots + j_l} \exp(-(n_1 + \dots + n_l)k)$

given by Table 4 for the two attributes quality models are approximations: the hypergeometric distribution is approximated by the binomial distribution, and the negative hypergeometric distribution is approximated by the Poisson distribution, compare the approximations discussed by **Single Sampling by Attributes and by Variables**. The approximations are tolerable for large lot size N . From a type B view, where the lot generating process is the inspection target, the values given by Table 4 are exact. For a detailed analysis, see [2] or [3].

Under multiple sampling, the actually reached cumulative sample size SN is a random variable in which $SN = n_l^{(c)} = n_1 + \dots + n_l$ iff the procedure terminates at stage l , $1 \leq l \leq m$. A suitable performance measure is the expectation $E[SN]$, the so-called average sample number (ASN), which as a function of the sampling plan parameters and the lot quality indicator is denoted as $ASN_{n,a,r}(q)$.

$ASN_{n,a,r}(q)$ can be calculated by the termination probabilities $T_{n,a,r,1}(q), \dots, T_{n,a,r,m}(q)$ in the manner described by Table 4.

Under curtailed inspection, the random cumulative sample size is $SN = n_l^{(c)} = n_1 + \dots + n_l$, if the procedure terminates at stage l with acceptance, and $SN = n_1 + \dots + n_{l-1} + j$, if the procedure terminates at stage l with rejection, where j is the smallest number such that $t_1 + \dots + t_{l-1} + j = r_l$. Formulas for the $ASN = E[SN]$ under curtailment are provided by Montgomery [4] and in more technical detail by Uhlmann [2] and Hald [5].

Sequential Sampling

Sequential sampling plans and the corresponding decision algorithm are defined in a way analogous to the definition of multiple sampling plans, except

4 Multiple Sampling Plans

that there is no upper limit m imposed on the number of sampling stages, i.e. the triple (n, a, r) consists of infinite sequences. Tables 2 and 3 remain valid for explaining the parameters and the decision algorithm, except that there is no terminating restriction.

Usually, sequential sampling plans are simplified by the following two assumptions: (a) The stage sample sizes are identical, i.e. $n = n_1 = n_2 = \dots$ (b) The differences of rejection and acceptance numbers remain invariant, i.e., $d = r_l - a_l$ for $l = 1, 2, \dots$. Under the simplifications (a) and (b), a sequential sampling plan can be defined by a tuple (n, s, h_1, h_2) where n is the invariant sample size, s is a *slope*, h_1, h_2 are *increments*, and where $a_l = sln - h_1$ are the acceptance numbers, $r_l = sln + h_2$ the rejection numbers. The operation of a sequential sampling plan (n, s, h_1, h_2) is illustrated by Table 5. Plans with sample size $n \geq 2$ are called *group sequential plans*. *Item-by-item sequential plans* (*unit sequential plans*) with $n = 1$, i.e. only one item is drawn at each stage, are very common. If there are no economic or administrative restrictions for the sample (group) size n in a group sequential plan, Cowden [6] suggests the use of the smallest group size under which acceptance is possible in the first stage (first sample), i.e. n is the smallest integer exceeding h_1/s .

As for multiple sampling, the essential performance measures are the OC and the ASN function. The OC function of a type A sequential sampling plan (n, s, h_1, h_2) is the conditional probability $L_{n,s,h_1,h_2}(q) = P(\text{accept the lot} \mid \text{lot quality indicator} = q)$ of accepting the lot under the condition that the lot quality indicator (lot proportion nonconforming or lot average number of defects) adopts a value q .

The ASN function is $ASN_{n,s,h_1,h_2}(q) = E[\text{cumulative sample size} \mid \text{lot quality indicator} = q]$. In principle, the formulas given in Table 4 can be adapted to the sequential case by substituting finite by infinite sums. However, the resulting expressions are difficult both for theoretical and for numerical analysis. Better closed expressions and approximations are provided by [5].

Design by Prescribing Two Points of the OC Function

The design of sampling plans based on prescribing two points of the OC function is described for simple sampling plans by **Single Sampling by Attributes and by Variables**. The same approach can be used for designing iterative sampling plans: a good quality level q_1 should be accepted with a high probability $L(q_1) \geq 1 - \alpha$, and a poor quality level q_2 should be accepted with a small probability $L(q_2) \leq \beta$, where $0 < \beta < 1 - \alpha$.

For multiple sampling plans, the two-points design does not determine a unique plan among all plans (n, a, r) . Additional constraints are necessary to assure uniqueness, e.g. requiring that the sample sizes satisfy a geometric progression $n_{i+1} = cn_i$. See [7] for a survey.

For sequential item-by-item sampling plans $(1, s, h_1, h_2)$, the simple approximate solution developed by [8] is displayed by Table 6. The results obtained by Wald [8] include approximations to the OC and ASN functions of the two-points design plan $(1, s, h_1, h_2)$.

Iterative Sampling with Rectification

Under rectifying sampling for lot proportion nonconforming, a rejected lot is subject to 100% inspection.

Table 5 The sampling and decision algorithm of a sequential sampling plan (n, s, h_1, h_2)

Stage	Operation	Decision algorithm	
1	Sample n items Calculate sample statistic t_1	$t_1 \leq s \cdot n - h_1$	$\implies$ terminate, accept
		$t_1 \geq s \cdot n + h_2$	$\implies$ terminate, reject
		$s \cdot n - h_1 < t_1 < s \cdot n + h_2$	$\implies$ continue with stage 2
$\vdots$	$\vdots$	$\vdots$	
l	Sample n items Calculate sample statistic t_l	$t_1 + \dots + t_l \leq sln - h_1$	$\implies$ terminate, accept
		$t_1 + \dots + t_l \geq sln + h_2$	$\implies$ terminate, reject
		$-h_1 < t_1 + \dots + t_l - sln < h_2$	$\implies$ continue with stage $l + 1$
$\vdots$	$\vdots$	$\vdots$	

Table 6 Design of an item-by-item sequential sampling $(1, s, h_1, h_2)$ plan by prescribing two points of the OC function

Design rule	$0 < q_1 < q_2 < 1, \quad 0 < \beta < 1 - \alpha, \quad L_{1,s,h_1,h_2}(q_1) \geq 1 - \alpha \geq \beta \geq L_{1,s,h_1,h_2}(q_2)$	
Item-by-item plan $(1, s, h_1, h_2)$ for lot proportion nonconforming p		
Increments	$h_1 = \frac{\log(1-\alpha)-\log(\beta)}{\log\left(\frac{p_2}{p_1}\right)+\log\left(\frac{1-p_1}{1-p_2}\right)}, \quad h_2 = \frac{\log(1-\beta)-\log(\alpha)}{\log\left(\frac{p_2}{p_1}\right)+\log\left(\frac{1-p_1}{1-p_2}\right)}$	
Slope	$s = \frac{\log(1-p_1)-\log(1-p_2)}{\log\left(\frac{p_2}{p_1}\right)+\log\left(\frac{1-p_1}{1-p_2}\right)}$	
Item-by-item plan $(1, s, h_1, h_2)$ for lot average number of nonconformities k		
Increments	$h_1 = \frac{\log(1-\alpha)-\log(\beta)}{\log(k_2)-\log(k_1)}, \quad h_2 = \frac{\log(1-\beta)-\log(\alpha)}{\log(k_2)-\log(k_1)}$	
Slope	$s = \frac{k_2-k_1}{\log(k_2)-\log(k_1)}$	
Approximate OC and ASN functions		
Generic OC	$L_{1,s,h_1,h_2}(q) = \frac{\exp(h_2t)-1}{\exp(h_2t)-\exp(-h_1t)}, \quad t = t(q)$	
$q = p$	$t(p)$ defined as the inverse $p(t) = \frac{\exp(st)-1}{\exp(t)-1}$	
$q = k$	$t(k)$ defined as the inverse $k(t) = \frac{st}{\exp(t)-1}$	
Generic ASN	$\frac{h_2-(h_1+h_2)L_{1,s,h_1,h_2}(q)}{q-s},$ to be used with $q = p$ or $q = k$	

All nonconforming items discovered in the sample and, in case of rejection, in the lot are replaced by conforming ones. The performance measures of iterative rectification plans are defined analogously to the performance measures of single rectification plans (*see Single Sampling by Attributes and by Variables*): as a function of the proportion nonconforming p in the submitted lot, the *average outgoing quality* $AOQ(p)$ is the expected value of the lot proportion nonconforming after application of a rectification plan, and the *average total inspection* $ATI(p)$ is the expected value of inspected items. The *average outgoing quality limit* is the maximum $AOQL = \max_{0 \leq p \leq 1} AOQ(p)$. Table 7 displays exact formulas

for multiple sampling plans, and approximate formulas for sequential item-by-item plans designed by prescribing two points of the OC function as shown in Table 6, see [7] for a derivation of the approximations.

Comparison of Single and Iterative Sampling

The intuitive motivation of iterative sampling is a reduction of expected sample size under extreme (very good or very bad) lot quality. For a simple quantitative comparison, one can compare the ASN functions of single and iterative plans which are

Table 7 Performance measures of iterative rectification plans for lot proportion nonconforming p

Multiple m -stage plan $(n, a, r)^{(a)}$	
$AOQ(p)$	$AOQ(p) = \sum_{l=1}^m p \left(1 - \frac{n_1 + \dots + n_l}{N}\right) A_{n,a,r,l}(p)$
$ATI(p)$	$ATI(p) = \sum_{l=1}^m (n_1 + \dots + n_l) A_{n,a,r,l}(p) + N (1 - L_{n,a,r,l}(p))$
Item-by-item sequential plan $(1, s, h_1, h_2)^{(b)}$	
$AOQ(p)$	$AOQ(p) \approx p L_{1,s,h_1,h_2}(p)$
$ATI(p)$	$ATI(p) \approx \frac{L_{1,s,h_1,h_2}(p)(\log(\beta)-\log(1-\alpha))}{p \log\left(\frac{p_2}{p_1}\right) + (1-p) \log\left(\frac{1-p_2}{1-p_1}\right)} + N (1 - L_{1,s,h_1,h_2}(p))$

^(a) See Table 4 for the probabilities $A_{n,a,r,l}(p)$ and the OC $L_{n,a,r,l}(p)$

^(b) Notation and quantities from Table 6

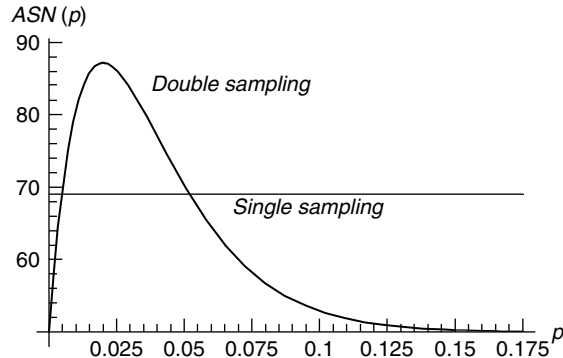


Figure 1 Average sample number $ASN(p)$ for single sampling plan $(69, 1)$ and double sampling plan with $n = (50, 100)$, $a = (0, 3)$, $r = (2, 4)$

matched to have approximately the same OC function. An instance of this comparison under inspection for lot proportion nonconforming p is displayed by Figure 1. For the single sampling plan $(69, 1)$ and the double sampling plan with $n = (50, 100)$, $a = (0, 3)$, $r = (2, 4)$, the OC functions are very close. With respect to the ASN, the double plan is superior to the single plan for p smaller than 0.005 and p larger than 0.06. For the large intermediate range from 0.5 to 6% single sampling is better.

Opposed to the advantage of iterative sampling for extreme lot quality are several disadvantages of iterative sampling: administrative difficulties;

incompetence of inspectors to execute the algorithm correctly; hidden setup sampling costs from repeated sampling, in particular under sequential plans. Iterative methods should be preferred over single sampling only if the product quality history inspires data-based confidence that lot quality is really on one of the extreme sides.

References

- [1] Dodge, H.F. (1969). Notes on the evolution of acceptance sampling plans. Parts I, II, III, *Journal of Quality Technology* **1**, 77–88, 155–162, 225–232.
- [2] Uhlmann, W. (1982). *Statistische Qualitätskontrolle*, B.G. Teubner, Stuttgart.
- [3] Schilling, E.G. (1983). *Acceptance Sampling in Quality Control*, Marcel Dekker, New York.
- [4] Montgomery, D.C. (2005). *Introduction to Statistical Quality Control*, 5th Edition, John Wiley & Sons, New York.
- [5] Hald, A. (1981). *Statistical Theory of Sampling Inspection by Attributes*, Academic Press, London.
- [6] Cowden, D.J. (1957). *Statistical Methods in Quality Control*, Prentice-Hall, Englewood.
- [7] Duncan, A.J. (1986). *Quality Control and Industrial Statistics*, 5th Edition, Irwin Publishing, Homewood.
- [8] Wald, A. (1947). *Sequential Analysis*, John Wiley & Sons, New York.

KODAKANALLUR KRISHNASWAMY SURESH
AND RAINER GÖB

Sampling in Industrial Standards

Evolution of Sampling Standards

During the first few decades of the twentieth century, many industrialized countries including the United States, United Kingdom, Germany, and France had established national standardization bodies that were the forerunners of those that exist today. Early developmental focus was primarily on product standards, especially driven by the demand for goods created during and between the two world wars. With the development of new, generically applicable **quality control** methods such as **control charts** and sampling inspection by the engineers and mathematicians of the Western Electric Company and Bell Telephone Laboratories in the 1920 and 1930s, it would only be a matter of time until the need for standardization of such statistical methods would be necessary in industrial production.

The basis of the earliest standards was a direct result of the expansion of the telephone industry in the United States and the associated demand for high-quality products that made up the complex telephone systems. Through the efforts of Shewhart, Dodge, Romig, and others, sampling methodologies were developed as alternatives to 100% inspection. These efforts resulted in published rules for the use of control charts and tables for sampling inspection in technical journals and books. By the early 1930s, these new methodologies were becoming better known both within the United States and across the ocean, leading to the creation of one of the first official national standards on the subject with Pearson's development of British Standard (BS) BS-600 (see [1], Part I).

With the outbreak of World War II in 1939, the demand for quality products to support the war effort greatly increased. Foreseeing that the United States would likely be drawn into the conflict, the American Standards Association took the initiative to create two standards in May 1941 with respect to quality control and the use of control charts (see [1], Part II). With the United States' formal entry into the war in December 1941, an unprecedented demand for quality goods and products by the military forces became reality. Two options were considered to address this

problem: (a) emphasize training of contractors and require their use of process quality controls, or (b) develop and apply a system of **acceptance sampling** inspection procedures for use by government inspectors. Owing to the urgent need for a practical solution, the primary focus shifted to the development of sampling inspection procedures, although work on process quality control was continued and resulted in a third standard on the subject in April 1942.

By July 1942, the US Army Service Forces had produced the first forerunner to a sampling inspection standard *via* its document "Standard Inspection Procedures – Quality Control". The goals of these attribute-based sampling plans were to enable the military to quickly obtain enormous quantities of necessary materials that were both reliable and safe, while attempting to treat all vendors fairly and alike. To facilitate implementation, a relatively simple approach was devised that was based on acceptable quality levels and a method of switching between normal and tightened inspection in response to submitted product quality levels. Research on sampling inspection remained a priority during the war years, leading to the development of sampling plans based on producer's risk quality and consumer's risk quality for both sampling by attributes and sampling by variables, and the creation of the theory for **sequential sampling**. Sampling plans based on this research were adopted into a US Navy standard issued in 1945.

Following the war, the focus shifted to adopting a more unified approach to **attribute sampling** inspection within the military branches (see [1], Part III). This led to the creation of the Joint Army–Navy standard JAN-STD-105 in 1949. In 1950, this standard was revised and replaced by the now familiar military standard MIL-STD-105A. This standard experienced technical amendments over the years (version B in 1958, C in 1961, D in 1963, and finally E in 1989) but maintained the same basic features. *See Attributes Sampling Schemes in International Standards* for further details.

In parallel to the evolution of the military standard attribute sampling plans, work also commenced on the development of compatible inspection by variables sampling plans. Different branches of the US military had created variables sampling plans as early as 1952 (Navy Ordnance Standard NAVORD OSTD 80) and 1954 (Ordnance Standard (ORD-M603-10)). These standards were replaced by the now well-known MIL-STD-414 in 1957

(see [2] and [3]), based largely on the research in [4]. MIL-STD-414 was designed to complement MIL-STD-105 for product characteristic measurements that could be assumed to follow a **normal distribution**. See **Variables Sampling Schemes in International Standards** for further details.

To address another common quality control situation, the US military also created standards for **continuous sampling** procedures for inspection by attributes, designed for processes producing a continuous stream of products rather than a series of lots. The first standardized procedures were produced by different military branches in the early 1950s, culminating in the issuance of handbooks H-106 and H-107 in the late 1950s, covering multilevel and single-level continuous sampling procedures respectively (see [1], Part IV, and [5]). The procedures from the two handbooks were amalgamated in 1962 to form MIL-STD-1235. Like its counterpart MIL-STD-105, the standard underwent technical amendments over the years (version A in 1974, B in 1981, and finally C in 1988) but maintained the same basic features. See **Continuous Sampling** for further details.

Movement from Military to National to International Standards

With the growth of international commerce in the years following World War II, the sampling standards that were originally focused on military procurement of goods found themselves being increasingly used by industry in general for all sorts of products. With the broadened application of the standards, it was evident that a transition from the traditional military custodianship of the sampling standards to civilian bodies formed specifically for the purpose of developing and maintaining standards in general would eventually take place. Although in some countries sampling standards had been sponsored by national standards bodies as early as the 1930s and 1950s (notably the United Kingdom and Japan respectively), the transition began to pick up pace throughout the 1970s and 1980s. For example, in 1971 the American National Standards Institute (ANSI) introduced its version of MIL-STD-105D designated as ANSI Z1.4. The following year, ANSI had incorporated MIL-STD-414 under the new designation of ANSI Z1.9. Another important milestone occurred in 1974 with the International Organization for Standardization (ISO) issuing ISO 2859

as an international replacement for MIL-STD-105D. In the years following these initial developments, military organizations such as the US Department of Defense and the UK Ministry of Defence have adopted policies of discontinuing development and maintenance of standards where appropriate civilian standards are available. Most of the well-known military standards have been formally canceled as a result with notices to use the replacement national standards.

Over the past few decades, the development and maintenance of sampling standards has increasingly shifted away from national bodies toward international sponsorship. As most national standards bodies are members of ISO and generic sampling standards are of interest to many countries, it has become more economical for interested parties to collaborate in the development of these standards as delegates of their national standards bodies on international technical committees. Member countries then often adopt the resulting ISO standard as their national standard.

For the development of general statistical standards, ISO has created Technical Committee TC69, with Subcommittee SC5 specifically responsible in the area of acceptance sampling. ISO TC69/SC5 is the custodian of more than a dozen generic standards related to sampling inspection; these are briefly described in the following section. Refer to [6] for full titles and additional information.

International Sampling Standards (ISO)

ISO 2859 Sampling Procedures for Inspection by Attributes

ISO 2859 was the first generic sampling standard created by ISO and was developed to serve as the international replacement for MIL-STD-105D. Since its introduction in 1974 as a single document, the standard has now evolved into a multipart series. The primary application of the series is in the production of a continuing series of lots, providing single, double, multiple, and sequential **sampling schemes** for this purpose. However, Part 2 of the series provides sampling plans for isolated lots that were designed to intertwine with the design of the continuing series philosophy of Parts 1, 3, and 5. Part 4 of this series was designed to introduce an auditing component to the system to support the movement to formal **quality management** systems employing the sampling

schemes of this series. Part 10 of this series was created to provide a general introduction to acceptance sampling by attributes and provides a brief summary of the attribute sampling schemes and plans used in Parts 1 through 5. It also provides guidance on the selection of the appropriate inspection system for use in a particular situation. Refer to **Attributes Sampling Schemes in International Standards** for technical details.

ISO 3951 Sampling Procedures for Inspection by Variables

ISO 3951 was developed to serve as the international replacement for MIL-STD-414. Since its introduction in 1981 as a single document, the standard has now evolved into a multipart series. To the extent that is technically possible, the ISO 3951 series is consistent with the goals and philosophies of its counterparts in the ISO 2859 series. See **Variables Sampling Schemes in International Standards** for further details.

ISO 8422 and 8423 Sequential sampling plans (attributes and variables)

These two standards contain sequential sampling plans and procedures for inspection of discrete items by attributes and by variables for percent nonconforming for known standard deviation, respectively. These sampling standards are similar to the ISO 2859-5 and ISO 3951-5 sampling schemes but the plans are indexed in terms of the producer's risk point and the consumer's risk point rather than by acceptance quality limit (AQL), and therefore provide specific control of these risks whereas the AQL plans rely on switching rules to alter the risks in response to changing process quality. The standards are designed primarily for the analysis of samples taken from stable processes but the attributes standard may also be used for acceptance sampling of an isolated lot when its size is large, and the expected fraction nonconforming is small (significantly smaller than 10%). The sampling plans are applicable where the extent of nonconformity is expressed in terms of proportion nonconforming. The attributes plans may

also be used where nonconformity is expressed on a per item basis (per 100 items).

ISO 13448 Acceptance Sampling Procedures Based on the Allocation of Priorities Principle (APP)

Part 1 of this series gives guidance and explains the methodology of the allocation of priorities principle (APP). Its principal premise is to permit the use of any type of prior objective and subjective information when determining the appropriate sampling plan to be used in a given situation. As a result, when the information suggests that the production quality is very high and the supplier is trustworthy, no sampling inspection may be required. Conversely, when the opposite situation is suggested, 100% inspection may be warranted. Examples of such information are inspection results for previous lots, certification of the supplier's quality management system as being in conformity with ISO 9001, quality control data, and customers' subjective estimates of the supplier's capability to provide the desired quality, all of which may be summarized in a trust level. This allows a reduction or increase in the required **sample size** as the customer's trust in the producer increases or decreases, as the case may be. In addition, when successive inspections of the same lot are carried out by different parties (i.e., customer, supplier, and/or a third party), the APP procedures allow parties to choose sampling plans commensurate with their differing resources and inspection capabilities. The choice of sampling plans needs only to be coordinated among the involved parties with respect to the customer's or supplier's risks in order to reduce the possibilities of contradictory results and wasted resources when quality is acceptable.

Part 2 provides **single sampling** by attributes plans for inspection of lots in accordance with the guidelines of Part 1. All subjective and objective information pertaining to the supplier's capability to provide the desired quality may be taken into account by the customer or a third party when deciding on his sampling plan, thus allowing smaller sample sizes when the information is favorable. This information is categorized in six trust levels ranging from full trust to a lack of any reliable information with respect to the supplier's quality capability with a seventh special trust level for important characteristics requiring 100% inspection. Guided by the trust information and

4 Sampling in Industrial Standards

associated consumer risks, permissible plans may be selected from the standard according to an index of 14 NQL (normative [or guaranteed] quality limits) values ranging geometrically from 0.15 to 65 and over. Rules for successive sample inspections on the same lot by different parties (i.e., supplier, customer, and/or a third party) are also provided.

ISO 14560 Acceptance Sampling Procedures by Attributes – Specified Quality Levels in Nonconforming Items Per Million

This standard specifies, for quality levels expressed as nonconforming items per million items, procedures for estimating the quality level of a single entity (e.g., a lot) and, when the production process is in statistical control, for estimating the process quality level based on evidence from several samples. Procedures are also specified for using this information when selecting a suitable sampling plan so as to verify that the quality level of a given lot does not exceed a stated limiting quality level (LQL). For the case where no prior sample data is available, guidance is given for presuming a process quality level in selecting a plan. The standard also provides incentives for suppliers to improve their quality. The lot acceptance portion of this specification requires larger sample sizes when quality declines and smaller sample sizes when quality improves. The sampling plans are indexed by 24 LQL values ranging from 500 to 100 000 nonconforming items per million items. In the case of a continuing series of lots, sample size is related to a prior estimate of the process quality level and in the case of an isolated lot, sample size is related to a presumed process quality level or that of the first few lots in a series.

ISO 18414 Acceptance Sampling Procedures by Attributes – Accept-Zero Sampling System Based on Credit Principle for Controlling Outgoing Quality

This standard specifies a system of single sampling schemes for lot-by-lot inspection by attributes. All the sampling plans of the system are of accept-zero form, i.e., no lot is accepted if the sample from it contains one or more nonconforming items. The schemes depend on a suitably defined average outgoing quality limit (AOQL) chosen by the user. No restrictions are placed on the choice of the value of the AOQL

or on the sizes of successive lots in the series. This standard's approach intends to induce a supplier, through the economic and psychological pressure of lot nonacceptance and consequent loss of accumulated credit, to attempt to maintain a nonconformity-free process, while assuring, by means of the lowest practicable sample sizes, that the long-term percentage of nonconforming items delivered to the customer or marketplace does not exceed the AOQL. This objective is achieved by a progressive reduction in the sample size in response to good quality history (i.e., increasing credit). The schemes enable the required sample size to be determined on the basis of a simple calculation involving the lot size, the AOQL value, and the accumulated credit, so no tables are required. Owing to the nature of the schemes, this standard is not applicable where inspection is destructive.

ISO 21247 Combined Accept-Zero Sampling Systems and Process Control Procedures for Product Acceptance

This standard provides a set of accept-zero sampling systems and procedures for planning and conducting inspections to assess quality and conformance to specified requirements. It also provides requirements for alternative acceptance methods proposed by the supplier. Such alternative methods would be based upon establishing and implementing an internal prevention-based quality management system as a means of ensuring that all products conform to requirements specified by the contract and associated specifications and standards. In the absence of acceptable alternative methods, the standard presents three types of matched sampling systems for the sampling inspection of product submitted to the customer for acceptance. These systems provide for sampling inspection by attributes, sampling inspection by variables, and for continuous sampling by attributes. The three types of matched sampling systems are indexed by seven specified verification levels and five sample size code letters. Approximately the same level of protection is provided by each of the three types of systems, when using the same code letter and verification level. The supplier has the option to utilize the type of sampling system, at the same verification level, that best suits the production process. The sampling systems and procedures of the standard are not intended for use with destructive tests or where product screening is not feasible or desirable. In such

cases, it permits the sampling strategy to be specified in the contract or product specifications.

It is noteworthy that this standard was developed as a civilian replacement for US Department of Defense (DoD) Test Method Standard MIL-STD-1916 "DoD Preferred Methods for Product Acceptance" and was directly sponsored by DoD resources. The issuance of MIL-STD-1916 in 1996 marked a dramatic change in the US military's general philosophy regarding the role of sampling inspection in quality control, especially with respect to the AQL concept and its implied acceptable levels of non-conforming product. The standard's stated purpose was to encourage product suppliers to submit efficient and effective **process control** (prevention) procedures in place of prescribed sampling (detection) requirements, with emphasis on continuous improvement. MIL-STD-1916 is expected to be canceled as a result of the issuance of ISO 21247.

*Draft International Standard ISO/DIS 24153
Random Sampling and Randomization Procedures*

To ensure the integrity of the results obtained through the application of the various sampling standards, the validity of the random sampling procedures used is critical. Previously, sampling standards had simply specified the need for a sample to be drawn at random, very frequently without identifying how this should be done. Experience of many experts and practitioners has indicated that samples were often drawn in accordance with methods that were inherently biased. This standard was created to fill

the void and provide an authoritative reference for not only sampling inspection situations but various other experimental situations as well. It defines procedures for random sampling and **randomization**, including approaches based on mechanical devices, tables of random numbers, and portable computer algorithms. It also provides information needed to audit the results of a random sampling exercise when required by quality management personnel or regulatory bodies. It is expected that this document will become a full international standard in late 2007.

References

- [1] Dodge, H.F. (1969). Notes on the evolution of acceptance sampling plans. Parts I, II, III, IV, *Journal of Quality Technology* **1**, 77–88 155–162, 225–232; **2**, 1970: 1–8.
- [2] Kao, J.H.K. (1971). MIL-STD-414 sampling procedures and tables for inspection by variables for percent defective, *Journal of Quality Technology* **3**, 28–37.
- [3] Duncan, A.J. (1975). Sampling by variables to control the fraction defective: part 1, *Journal of Quality Technology* **7**, 34–42.
- [4] Lieberman, G.J. & Resnikoff, G.J. (1955). Sampling plans for inspection by variables, *Journal of the American Statistical Association* **50**, 457–516.
- [5] Banzhaf, R.A. & Brugger, R.M. (1970). MIL-STD-1235(ORD), single- and multi-level continuous sampling procedures and tables for inspection by attributes, *Journal of Quality Technology* **2**, 41–53.
- [6] International Organization for Standardization. (2007). Internet home page: <http://www.iso.org/>.

FRANK PALCAT

Attributes Sampling Schemes in International Standards

Introduction

Standardized sampling procedures for the inspection of attribute quality characteristics were developed during World War II by the United States Department of Defense. The first military standard MIL-STD-105A was published in 1950. Updated versions were issued as MIL-STD-105B, MIL-STD-105C and MIL-STD-105D. MIL-STD-105E [1] is the latest version issued in 1989.

Further national and international standards emerged as derivatives of MIL-STD. In 1971, the American Society for Quality (ASQ) introduced the American National Standards Institute ANSI/ASQ Z1.4 as a derivative of MIL-STD105D, the most recent version of ANSI/ASQ Z1.4 [2] issued in 2003 is a modified counterpart of MIL-STD-105E. The US military cancelled MIL-STD-105E by a note in February 1995 and recommended the use of ANSI/ASQ Z1.4. ISO 2859 was the first sampling standard created by the ISO (International Standards Organization) as a derivative of and international alternative to MIL-STD-105. Issued as a single document in 1974, the standard evolved into a series of presently six parts, labeled 1 through 5, and 10, see [3–8]. The present version of ISO 2859 has been adopted as a national standard in many countries, e.g., as BS 6001 in the United Kingdom or as DIN ISO 2859 in Germany.

MIL-STD-105E provided the pattern for all further developments in standardized attributes sampling. The article describes the purpose, design and operation of MIL-STD-105E in detail, and analyzes the relation among MIL-STD-105E and the very similar ISO 2859-1. The further parts of ISO 2859 are briefly reviewed. Further attributes sampling standards are discussed by **Sampling in Industrial Standards**.

Basics of MIL-STD-105 and Derivatives

Three issues are essential for statistical **product inspection** (SPI): the quality model, the control objective, and the quality indices used.

Quality Models

The national and international standards considered below are based on the two attributes quality models of SPI. (a) *Conforming–nonconforming*: items are classified as conforming (nondefective) or nonconforming (defective), the lot quality indicator is the lot proportion nonconforming p , the process quality indicator is the process proportion nonconforming π , and the sample statistic t is the number of nonconforming items in the sample. (b) *Nonconformities*: item quality is determined by the number of nonconformities (defects) of the item, the lot quality indicator is the lot average number of nonconformities k , the process quality indicator is the process average number of nonconformities λ , and the sample statistic t is the number of nonconformities on items in the sample. See **Sampling Inspection of Products** for more information on quality models of SPI.

The Control Objective

Type A product inspection is distinguished from type B product inspection (see **Sampling Inspection of Products**). Type A inspection is designed to control the quality of individual lots: the quality target is the lot quality indicator, and the decision on acceptance or rejection bears on the individual lot. Type B procedures are designed to control the quality of the lot generating process of manufacture, transport and so on. The quality target is the process quality indicator, and the decision on acceptance or rejection refers both to the individual lot and to the lot generating process.

Decidedly, MIL-STD-105 and its derivatives are intended for type B inspection. MIL-STD-105E [1, p. 1] states “These plans are intended primarily to be used for a continuing series of lots or batches.” The control tool directed toward the process of lots are the *switching rules*, as pointed out by ISO 2859-1 [3, p. 13]: “These schemes are intended primarily to be used for a continuing series of lots, that is, a series of lots to allow the switching rules to be applied.” The switching rules are a regime to change between different levels of inspection severity (normal, tightened, reduced) according to *past decisions* on lots, see the more detailed exposition below. However, the switching regime is not designed as a mechanism to convey *previous information* on process quality. The switching regime has two purposes: (a) penalizing poor process quality by high rejection rates,

eventually leading to discontinuation of inspection (rejection without inspection); thus, the consumer is protected against poor quality, and pressure is exerted on the provider to submit satisfactory quality; (b) reducing the inspection cost in cases of proven good process quality.

Using MIL-STD-105E or ISO 2859-1 for the inspection of isolated lots without applying the switching rules contradicts the intrinsic intention of the schemes. There are two options for the inspection of isolated lots under a type A perspective: use the LQ (limiting quality) tables in MIL-STD-105 or ISO 2859-1, or use the specific LQ scheme provided by ISO 2859-2 [4].

Quality Indices: AQL and LQ

The design of sampling plans or **sampling schemes** involves *quality indices*. Commonly, the quality indices are critical levels of the quality target, i.e., the lot or the process quality indicator. MIL-STD-105 and its derivatives use two quality indices: the first index is termed as the acceptance quality level (AQL), in MIL-STD and as *acceptance quality limit* in ISO, and the LQ.

MIL-STD-105E [1, p. 2] defines “When a continuous series of lots is considered, the AQL is the quality level which, for the purposes of sampling inspection, is the limit of a satisfactory process average.” ISO 2859-1 [3, p. 25] defines the AQL as the “quality level that is the worst tolerable process average when a continuing series of lots is submitted for **acceptance sampling**”. The AQL clearly refers to the process quality indicator. However, the definitions are not precise. The operational clauses “for the purposes of sampling inspection” and “submitted for acceptance sampling” are rather confusing. From a business management point of view, the “satisfactory” or “tolerable” process quality level should be defined independently of the inspection aspect as a *breakeven quality*, i.e., a threshold value which separates satisfactory from unsatisfactory process quality levels. In final analysis, the definition of the breakeven quality has to be based on economic criteria. The standards neither mention the concept of a breakeven quality nor give any guidelines to a rational selection thereof. As a consequence, the AQL is mostly not specified on grounds of economic expertise, but stipulated by contracts, or prescribed by convenience of industrial segments, or dictated by authorities, or chosen

arbitrarily. The situation is even more complicated by ISO 2859-1 where the AQL is implicitly identified with the producer’s quality (*see Sampling Inspection of Products*), and high probabilities of acceptance for lots of AQL quality are tabulated. Why should the worst tolerable quality level be accepted with high probability? *See Sampling Inspection of Products and Acceptance Sampling in Modern Industrial Environments* for a further discussion of quality levels.

In MIL-STD-105, the LQ is a secondary index to be used for determining sampling plans with a type A focus [1, p. 9]. The LQ is the primary quality index of ISO 2859-2 [4]. The LQ, also denoted as the consumer’s quality level (CQL) by ISO 2859-1 or as the LTPD (lot tolerance percent defective) by the Dodge–Romig scheme, expresses a limit for tolerable values of the lot quality indicator (lot proportion nonconforming or lot average number of nonconformities) from the consumer’s view. The probability $L_S(LQ)$ of accepting a lot with proportion nonconforming LQ under a sampling plan S is the *consumer’s risk* (CR) (*see Sampling Inspection of Products*), where L_S is the operating characteristic (OC) function of the sampling plan S . The CR should be small. ISO 2859-2 approximately uses the customary value 10%, i.e., $L_S(LQ) \approx 0.10$. However, such statistical requirements give no guidelines for choosing values of the LQ. The situation is as problematic as for the AQL.

The Structure of MIL-STD-105E

Many national and international attributes sampling standards reproduce the pattern of MIL-STD-105 with minor modifications only. An understanding of MIL-STD-105E is the key to the understanding of many further sampling systems.

MIL-STD-105E as a Sampling System

Schilling [9, p. 277] establishes the following hierarchy: a *sampling scheme* is an indexed collection of sampling plans, including an algorithm for choosing an appropriate plan; a *sampling system* is an indexed collection of sampling schemes, including an algorithm for switching between the schemes. In this sense, MIL-STD-105E is a proper sampling system. In the schemes, the sampling plans are

Table 1 Inspection modes of MIL-STD-105E

Sampling modes	Single, double, seven-stage	OC matched plans No particular advice for selection
Inspection levels	General levels: I, II, III Special levels: S-1, S-2, S-3, S-4	Affect fraction of sample size and lot size. II: regular. I: sample smaller than I, larger risk. III: sample larger than I, smaller risk. S-1, S-2, S-3, S-4: very small samples, large risk.
Inspection states	Normal, tightened, reduced, discontinued	Affect fraction of sample size and acceptance number. Governed by switching rules, see Table 4.

indexed by the lot size N and the AQL. In the system, the subschemes are indexed by the sampling modes, inspection levels and inspection states, see Table 1.

The Contents of MIL-STD-105E

The print version of MIL-STD-105E provides the following materials: (a) information on purpose and scope of the sampling system; (b) basic definitions, in particular of the quality models and quality indices; (c) advice on the implementation and operation of the system, in particular, formation of lots, inspection modes as given by Table 1, sampling technique, switching rules as given by Table 4, algorithm of determining a sampling plan from tables; (d) master tables for determining sampling plans. The first table contains the sample size code letters, see Table 2. The master tables are classified by the sampling mode and by the inspection state. As an example, the master table for **single sampling** under normal inspection is provided by Table 3; (e) Tables of the average outgoing quality limit (*see Rectification Sampling Schemes*), for single sampling under tightened inspection; (f) Tables of the LQ under single sampling if $L_S(LQ) = 0.05$ and if $L_S(LQ) = 0.10$; (g) A table of limit numbers to be used in the rule for switching from normal to reduced inspection, see the rules in Table 4; (h) Average sample number curves (*see Multiple Sampling Plans*) for double and multiple sampling; and (i) The OC curves for the sampling plans under all sample size code letters. The binomial OC is used for inspection for lot proportion nonconforming, and the Poisson OC is used for inspection for lot average number of defects (*see Single Sampling by Attributes and by Variables; Multiple Sampling Plans*).

Table 2 Table of sample size code letters

Lot or batch size	Special inspection levels				General inspection levels		
	S-1	S-2	S-3	S-4	I	II	III
2 to 8	A	A	A	A	A	A	B
9 to 15	A	A	A	A	A	B	C
16 to 25	A	A	B	B	B	C	D
26 to 50	A	B	B	C	C	D	E
51 to 90	B	B	C	C	C	E	F
91 to 150	B	B	C	D	D	F	G
151 to 280	B	C	D	E	E	G	H
281 to 500	B	C	D	E	F	H	J
501 to 1200	C	C	E	F	G	J	K
1201 to 3200	C	D	E	G	H	K	L
3201 to 10000	C	D	F	G	J	L	M
10001 to 35000	C	D	F	H	K	M	N
35001 to 150000	D	E	G	J	L	N	P
150001 to 500000	D	E	G	J	M	P	Q
500001 and over	D	E	H	K	N	Q	R

The Implementation and Operation of MIL-STD-105E

The implementation of MIL-STD-105E requires three subsequent decisions which are associated with different functions of a company.

1. Select one of the 26 AQL values ranging from 0.01% up to 1000%, see Table 3. Values larger than 10% are intended for inspection for the average number of defects (nonconformities) per item. The selection of the AQL value is essentially a business management task. The choice of a threshold for tolerable process quality requires a review and synthesis of all relevant aspects of technical, engineering, contractual, and economic aspects.

Table 3 Master table of single sampling plans (n, c) under normal inspection. Acceptance number $c = Ac$, rejection number $c + 1 = Re$. Down arrows prescribe first sampling plan below arrow, up arrows prescribe first sampling plan above arrow

Sample size code letter	Sample size	Acceptable quality level (normal inspection)															
		0.010	0.015	0.025	0.040	0.065	0.10	0.15	0.25	0.40	...	250	400	650	1000		
		Ac	Re	Ac	Re	Ac	Re	Ac	Re	Ac	Re	Ac	Re	Ac	Re	Ac	Re
A	2											10	11	14	15	21	22
B	3											14	15	21	22	30	31
C	5											21	22	30	31	44	45
D	8											30	31	44	45	↑	
E	13											44	45	↑			
F	20											↑					
G	32											0	1				
H	50											0	1				
J	80											↑					
K	125											↑					
L	200											1	2				
M	315											1	2				
N	500											2	3				
P	800											3	4				
Q	1250	0	1	↑	↓	1	2	3	4	5	6	7	8	10	11		
R	2000	↑		1	2	3	4	5	6	7	8	10	11	14	15		

Table 4 Switching algorithm of MIL-STD-105E

Start	Normal inspection
Normal ↓ tightened Tightened ↓ normal Tightened ↓ discontinued	2 out of 5 preceding lots have been rejected. 5 preceding lots have been accepted. In a tightened inspection series, a total of 5 lots is rejected. All of the conditions (a), (b), (c), (d) are satisfied: (a) 10 preceding lots on normal inspection, none rejected. (b) The total $t_{\Sigma} = t_1 + \dots + t_{10}$ of the sample statistics in 10 preceding samples satisfies $t_{\Sigma} \leq l$ where the limit number l is displayed by Table VIII [1, p. 32] as a function of the AQL and the total number n_{Σ} of sampled units in 10 preceding samples. (c) Production is at a steady state. (d) The authority responsible for sampling accepts reduced inspection. At least one of the conditions (a), (b), (c), (d) is satisfied: (a) A lot is rejected. (b) An inspection result terminates with $Ac < t < Re$, i.e., neither acceptance nor rejection is warranted. (c) Production is irregular or delayed. (d) Other conditions warrant that normal inspection shall be instituted.
Normal ↓ reduced	
Reduced ↓ normal	

2. Select the inspection level from I, II, III, S-1, S-2, S-3, S-4, and the sampling mode (single, double or multiple sampling). This step pertains to the **quality control** department and requires economic, organizational, and statistical expertise.
3. When the choices (1) and (2) are made, the selection of the inspection state and of the specific sampling plan is determined by the operational regime of MIL-STD-105E. This administrative phase is executed by the inspection force according to the following algorithm. (a) Determine the inspection state from the switching algorithm given by Table 4. (b) Determine the lot size N . (c) Use N and the inspection level to determine the sample size code letter X from Table 3. (d) Go to the relevant master table of sampling plans which is identified by the sampling mode and inspection state, see Table 3, for instance. In the table, identify the line with code letter X. In this line, find the sampling plan in the column labeled by the respective value of the AQL.

The Design of MIL-STD-105E

The structure of MIL-STD-105 dates back to sampling schemes developed for the US army from 1942 until 1945, *see Sampling in Industrial Standards* for a detailed historical analysis. Proceeding under wartime conditions the development was strongly oriented by pragmatic criteria. The sampling tables should be rapidly available, simply structured, convenient in implementation and operation. As a consequence, the design of MIL-STD-105 is not built on rigid and consistent theory, but on convenience principles, intuitive plausibility, and rules of thumb. We review the main results of the detailed and profound analysis of the design MIL-STD-105 by Hald [10, pp. 81–99].

Relation between Lot Size and Sample Size. Hald [10, p. 83] states that “there is no theoretical foundation between lot size and sample size in Mil. Std.” The relation was chosen on grounds of intuitive reasoning by the group which constructed the AQL sampling tables for the US army in 1942. A numerical comparison of sample size n_{II} and lot

size N for single sampling in inspection level II at normal inspection shows that $n_{II}/\sqrt{N}$ is approximately constant.

The sample sizes n_I and n_{III} for levels I and III are approximately $n_I \approx 0.4n_{II}$, $n_{III} \approx 1.6n_{II}$. The relation n_{S-i}/n_{II} for the sample sizes under the special inspection levels S-1, S-2, S-3, S-4 is decreasing in N . For largest lot size $N \geq 500\,000$ the maximum sample sizes are $n_{S-1} = 8$, $n_{S-2} = 13$, $n_{S-3} = 50$, $n_{S-4} = 125$, only. The inspection levels S-1, S-2, S-3, S-4 are intended for cases “where relatively small sample sizes are necessary and large sampling risk can or must be tolerated” [1, p. 7], e.g., in case of destructive inspection.

Range of Tabulated AQL Values. The 26 tabulated AQL values in the master tables are obtained by rounding from a geometric series $0.010 \cdot 10^{i/5}$, $i = 0, \dots, 25$. There is no deeper meaning behind this rule except that it seemed plausible to the group of experts which designed the scheme for the US army in the 1940s. MIL-STD-105E [1, p. 1] admits that the range is “not suitable for applications where quality level in the defective parts per million range can be realized”.

Relation between Sample Size and Acceptance Number. We restrict attention to single sampling, see the analysis of double and multiple sampling below. First, consider normal inspection. Analogously to the tabulation of the AQL values, the 15 tabulated values of the sample size n are obtained by rounding from a geometric series $2 \cdot 10^{i/5}$, $i = 0, \dots, 14$. Hence, in the master tables the products $n \times \text{AQL}$ are approximately constant on all left-bottom to right-top diagonals, see Table 3. Correspondingly, the acceptance numbers are fixed on these diagonals such that the OC values $L_S(\text{AQL})$ are also approximately fixed. However, this property is arbitrary and has no systematic motivation. The OC values $L_S(\text{AQL})$ should rather be fixed in the columns, i.e., for fixed quality threshold AQL the probability of acceptance at the threshold should also remain invariant. However, Hald [10, p. 85] demonstrates that in the columns for fixed AQL throughout $L_S(\text{AQL})$ increases from 0.882 to 0.986.

Tightened inspection generally prescribes smaller acceptance number at the same sample size. The acceptance number is determined so that the average outgoing quality limit (AOQL, see **Rectification**

Sampling Schemes) equals the AQL, see [10, p. 87]. This principle is meaningful only under rectifying inspection (see **Rectification Sampling Schemes**).

Single sampling plans are usually defined as a pair (n, c) of sample size n and acceptance number c (see **Single Sampling by Attributes and by Variables**). The master tables of MIL-STD-105E give both the acceptance number $c = \text{Ac}$ and the rejection number Re . Normal and tightened inspection assume $\text{Re} = c + 1$ as usual, i.e., the procedure terminates with acceptance and rejection. Reduced inspection assumes $\text{Ac} + 1 < \text{Re}$ so that the component (b) of the switching rule “reduced $\rightarrow$ normal”, Table 4, can take effect.

The Switching Rules. Hald [10, pp. 88–96] demonstrates the following properties of rules for switching from normal to tightened.

1. Consider a series of lots of constant proportion nonconforming p and constant size N . Then (a) the probability $p_T(p, N)$ of being on tightened inspection is a rapidly increasing function of p , (b) the probability $p_T(p, N)$ of being on tightened inspection is a decreasing function of N .
2. The probability of acceptance under normal inspection strongly exceeds the probability of acceptance under tightened inspection. Whereas the properties (a) and (2) are reasonable to exert pressure on the supplier, property (b) appears strange.

Double and Multiple Sampling Plans. The theory of multiple sampling is described by **Multiple Sampling Plans**. For each single sampling plan $S_{\text{single}} = (n, c)$ with acceptance number $c = \text{Ac} > 0$, MIL-STD-105E provides a double sampling plan S_{double} and a seven-stage sampling plan $S_{\text{seven-stage}}$ such that the sampling plans are approximately OC matched, i.e., $L_{S_{\text{single}}}(p) \approx L_{S_{\text{double}}}(p) \approx L_{S_{\text{seven-stage}}}(p)$.

Differences between MIL-STD-105E and ISO 2859-1

Because of the strong similarity of the standards, ISO 2859-1 can easily be understood by analyzing the few differences from MIL-STD-105E. Essentially the same differences hold between MIL-STD-105E and American National Standards Institute/American Society for Quality Control ANSI/ASQC Z1.4.

1. In the quality model, MIL-STD-105E uses the terms “defective” and “defect” whereas ISO 2859-1 uses the terms “nonconforming” and “nonconformity”.
 2. The term AQL is expanded as acceptable quality *level* in MIL-STD-105E, and as acceptable quality *limit* in ISO 2859-1.
 3. The gaps $Ac + 1 < Re$ between acceptance and rejection number under reduced inspection are given up such that $Ac + 1 = Re$ for all single sampling plans. Component (b) of the switching rule “reduced $\rightarrow$ normal”, Table 4, is eliminated.
 4. The components (a) and (b) of the switching rule “normal $\rightarrow$ reduced”, Table 4, are replaced by the component “The switching score is at least 30”. The sequence $s_1, s_2, \dots$, of switching scores is defined by Table 5.
 5. Double sampling plans in ISO 2859-1 are slightly changed so as reduce the average sample number.
 6. ISO 2859-1 contains tables of the probability $1 - L(AQL)$ of rejecting lots of AQL quality under single sampling. $1 - L(AQL)$ is termed as “*producer’s risk*”.
 7. ISO 2859-1 uses the term “*consumer’s risk quality*” (CRQ) instead of “limiting quality”. The consumer’s risk is fixed at 10%, i.e., $L(CRQ) = 0.10$.
- (b) widely accepted by industry, long-term practical experience;
 - (c) type B focus on the stream of lots;
 - (d) accrued quality experience is used for the choice of sampling plans in a planned and structured way through the switching rules.
- MIL-STD-105E and its derivatives like ISO 2859-1 have the following disadvantages.
- (a) no transparent and consistent design principles, design based on arbitrary reasoning and convenience;
 - (b) vague definition of the AQL, no clear guidelines for choosing AQL values;
 - (c) no foundation in an economic model, no economic evaluation of sampling;
 - (d) no mechanism for conveying prior information on the lot generating process into the design of sampling plans;
 - (e) discontinuation of inspection considered only in the sense of rejection without inspection, not in the sense of acceptance without rejection;
 - (f) design inconsistencies, in particular, considerable variation in the probability of accepting lots of AQL quality, probability of being on tightened inspection strongly decreasing in lot size;
 - (g) the range of AQL values, with a smallest value of 0.01% proportion nonconforming, not inadequate for modern parts per million quality requirements.

Appraisal of MIL-STD-105 and Derivatives

MIL-STD-105E and its derivatives like ISO 2859-1 have the following advantages.

- (a) simplicity, ease in implementation and operation;

Further Parts of ISO 2859

ISO 2859-2 [4], introduced in 1985, provides sampling plans indexed by LQ for isolated lot inspection.

Table 5 Recursive definition of the sequence $s_1, s_2, \dots$, of switching scores at subsequent inspections 1, 2, $\dots$ in ISO 2859-1)

Start	$s_1 = 0$
$s_t \rightarrow s_{t+1}$, single sampling	Let c be the acceptance number at inspection $t + 1$. If $c \geq 2$ and if the lot could be accepted under next smaller AQL: $s_{t+1} = s_t + 3$. If $c \geq 2$ and if the lot could not be accepted under next smaller AQL: $s_{t+1} = 0$. If $c \leq 1$ and if the lot is accepted: $s_{t+1} = s_t + 2$. If $c \leq 1$ and if the lot is not accepted: $s_{t+1} = 0$.
$s_t \rightarrow s_{t+1}$, double sampling	If the lot is accepted after the first sample: $s_{t+1} = s_t + 3$. If the lot is not accepted after the first sample: $s_{t+1} = 0$.
$s_t \rightarrow s_{t+1}$, five-stage sampling	If the lot is accepted after the third sample: $s_{t+1} = s_t + 3$. If the lot is not accepted after the third sample: $s_{t+1} = 0$.

The standard serves as an alternative in situations where the prerequisites of ISO 2859-1, a continuing series of lots and the opportunity to follow the switching regime, are not met. The basic idea is to provide consumer protection by imposing an upper limit on the CR, i.e., the probability $L_S(LQ)$ of accepting a lot of proportion nonconforming LQ. Except for the absence of the switching regime, the structure of ISO 2859-1 is analogous ISO 2859-1. The tabulated lot size ranges match those in ISO 2859-1. The suggested LQ values follow a geometric progression of ten values ranging from 0.5% to 32%. The design of the sampling plans is such that the CR is nominally around the value of 10%, i.e., $L_S(LQ) \approx 0.10$, but usually less. Detailed tables and graphs are provided giving OCs of the sampling plans and their relationship to those in ISO 2859-1, where applicable.

ISO 2859-3 [5], introduced in 1991, contains generic **skip-lot sampling** procedures. Skip-lot sampling procedures are designed to be used as an alternative to reduced inspection under the schemes in ISO 2859-1 where the production quality is high and there is confidence in the supplier's quality assurance system. The procedure involves determining randomly with a specified probability whether a lot presented for inspection will be accepted without inspection. See **Skip Lot and Chain Sampling** for more information on skip-lot procedures.

ISO 2859-4 [6] establishes sampling plans and procedures that can be used to assess whether the quality level of an entity (lot, process, and so on) conforms to a declared value. The sampling plans are designed to provide a risk of less than 5% of contradicting a correct declared quality level, with a 10% risk of failing to contradict an incorrect declared quality level that is related to the limiting quality ratio. Sampling plans are provided corresponding to three levels of discriminatory ability. These procedures are designed primarily for audit purposes and can be used when the quantity of interest is the number or fraction of nonconforming items, or the number of nonconformities, or the number of nonconformities per item.

ISO 2859-5 [7] contains **sequential sampling** schemes that supplement those in ISO 2859-1.

ISO 2859-10 [8] provides a general introduction to acceptance sampling by attributes and a brief

summary of the **attribute sampling** schemes and plans used in ISO 2859-1 through ISO 2859-5. It also provides guidance on the selection of the appropriate inspection system for use in a particular situation.

References

- [1] United States Department of Defense. (1989). *Military Standard, Sampling Procedures and Tables for Inspection by Attributes (MIL-STD-105E)*, U.S. Government Printing Office, Washington, DC.
- [2] American Society for Quality. (2003). *Sampling Procedures and Tables for Inspection by Attributes (ANSI/ASQ Z1.4)*, ASQ Quality Press, Milwaukee.
- [3] International Organization for Standardization. (1999). *ISO 2859-1: Sampling Procedures for Inspection by Attributes – Part 1: Sampling Schemes Indexed by Acceptance Quality Limit (AQL) for Lot-by-lot Inspection*, International Organization for Standardization, Geneva.
- [4] International Organization for Standardization. ISO 2859-2. (1985). *Sampling Procedures for Inspection by Attributes – Part 2: Sampling Plans Indexed by Limiting Quality (LQ) for Isolated Lot Inspection*, International Organization for Standardization, Geneva.
- [5] International Organization for Standardization. ISO 2859-3. (2005). *Sampling Procedures for Inspection by Attributes – Part 3: Skip-Lot Sampling Procedures*, International Organization for Standardization, Geneva.
- [6] International Organization for Standardization. ISO 2859-4. (2004). *Sampling Procedures for Inspection by Attributes – Part 4: Procedures for Assessment of Declared Quality Levels*, International Organization for Standardization, Geneva.
- [7] International Organization for Standardization. ISO 2859-5. (2005). *Sampling Procedures for Inspection by Attributes – Part 5: System of Sequential Sampling Plans Indexed by Acceptance Quality Limit (AQL) for Lot-by-lot Inspection*, International Organization for Standardization, Geneva.
- [8] International Organization for Standardization. ISO 2859-10. (2006). *Sampling Procedures for Inspection by Attributes – Part 10: Introduction to the ISO 2859 Series of Standards for Sampling for Inspection by Attributes*, International Organization for Standardization, Geneva.
- [9] Schilling, E.G. (1982). *Acceptance Sampling in Quality Control*, Marcel Dekker, New York.
- [10] Hald, A. (1981). *Statistical Theory of Sampling Inspection by Attributes*, Academic Press, London.

SAMINATHAN BALAMURALI, RAINER GÖB AND
CHI-HYUCK JUN

Variables Sampling Schemes in International Standards

Introduction

Issued in 1957, Military Standard 414 (MIL-STD-414) [1] was the first industrial standard on variables sampling, originally compatible with the attributes sampling standard MIL-STD-105A. Whereas the attributes sampling standard MIL-STD-105 evolved through several versions A, B, C, D, and E, MIL-STD-414 remained unchanged, except for the correction of two typos in 1968, and the compatibility between MIL-STD-105 and MIL-STD-414 was lost. The compatibility was restored by the civilian counterpart ANSI/ASQ Z1.9 whose latest version [2] issued in 2003 is compatible with MIL-STD-105E. The U.S. military cancelled MIL-STD-414 by a note in February 1999 and recommended the use of American National Standards Institute (ANSI)/ASQ Z1.9. For more historical background on industrial sampling standards (*see Sampling in Industrial Standards*).

On the international level, the International Standards Organization (ISO) introduced in 1989 the first variables sampling standard ISO 3951, compatible with the attributes sampling standard ISO 2859, which is a civilian counterpart of MIL-STD-105E. Like ISO 2859, ISO 3951 evolved into a multipart series. The most recent series issued in 2005 and 2006 contains three parts. ISO 3951-1 [3] is essentially compatible with ISO 2859-1, and in great part similar to ANSI/ASQ Z1.9. ISO 3951-2 [4] is an extension to multiple independent quality characteristics. ISO 3951-5 [5] is a sequential scheme for the inspection of a single quality characteristic. The part numbers are chosen to match ISO 2859 (*see Attributes Sampling Schemes in International Standards*), so some numbers are still unused. The present version

of ISO 3951 has been adopted as a national standard in many countries, e.g., as BS 6002 in the United Kingdom or as DIN ISO 3951 in Germany. Since no substantial up-to-date alternatives to the ISO standard exist in industrial variables sampling, the article concentrates on ISO 3951, with particular emphasis on ISO 3951-1 and on ISO 3951-2.

Basics of ISO 3951

Three issues are essential for statistical **product inspection** (SPI): the quality model, the control objective, and the quality indices used. ISO 3951 differs from the attributes sampling standards discussed by **Attributes Sampling Schemes in International Standards** in the quality model only. The control objectives and the quality indices are exactly the same.

The Quality Model of Variables Sampling

Variables sampling schemes are based on the *specification range* quality model (*see Sampling Inspection of Products; Single Sampling by Attributes and by Variables*). A conforming–nonconforming item quality indicator is built from a continuous measurement characteristic like weight, length, volume, by prescribing *specification limits* L and U for the measurement X . The model is explained by Table 1.

ISO 3951 is exclusively designed for the case of a continuous item characteristic X with **normal distribution** $N(\mu, \sigma^2)$, see [3, p. 2]. In the latter case, the process proportion nonconforming p is a function of the process mean μ , and the process standard deviation σ (*see Single Sampling by Attributes and by Variables*). Both the cases of known and unknown standard deviation σ are considered by ISO 3951: operating ISO 3951-1 under known process variance is called “ σ ” method, operating ISO 3951-1 under unknown process variance is called “ s ” method; see the classification of inspection modes in Table 2. ISO 3951-1 [3, p. 12] recommends that the inspection

Table 1 Quality models of variables sampling, based on a continuous measurement X

Specification limits	Item conforming iff	Process proportion nonconforming p
One-sided upper U	$X \leq U$	$p = p_U = P(X > U)$
One-sided lower L	$X \geq L$	$p = p_L = P(X < L)$
Two-sided $L < U$	$L \leq X \leq U$	$p = p_{L,U} = p_L + p_U = P(X < L) + P(X > U)$

2 Variables Sampling Schemes in International Standards

Table 2 Inspection modes of ISO 3951-1

Inspection method	“s” method, “σ” method	Prior knowledge on the process variance σ^2 : σ^2 unknown $\rightarrow$ “s” method σ^2 known $\rightarrow$ “σ” method
Inspection levels	General levels: I, II, III Special levels: S-1, S-2, S-3, S-4	Affect fraction of sample size and lot size II: regular. I: sample smaller than II, larger risk III: sample larger than II, smaller risk. S-1, S-2, S-3, S-4: very small samples, large risk.
Inspection states	normal, tightened, reduced, discontinued	Affect fraction of sample size and acceptance number. Governed by switching rules (see Table 5).

process starts operating with the “s” method where σ is estimated from each sample by the sample standard deviation s (see Table 3). If variability is under control, and a reliable long-term estimate of σ exists, one may switch to the “σ” method, which relies on the estimate as the true value of σ (see Table 3). The **estimation** of σ is explained by the International Organization for Standardization (p. 77. ISO 3951-1 [3]); p. 12, recommends to continue calculating sample standard deviations “for record purposes and to keep the **control charts** up-to-date”. The “σ” method allows for smaller sample sizes than the “s” method.

The classification of the parameters p , μ , σ^2 differs between MIL-STD-414 and ISO 3951. MIL-STD-414 defines the parameters as lot parameters. This view is inconsistent with the type B control objective of the standard, and leads to principal problems in the interpretation of variables sampling (see the discussion in **Single Sampling by**

Attributes and by Variables). ISO 3951 consistently defines p , μ , σ^2 as parameters of the lot generating process.

The Control Objective

As for the attributes sampling standards discussed by **Attributes Sampling Schemes in International Standards**, the control objective of ISO 3951 is of type B, i.e., directed toward “a continuing series of lots of discrete products, all supplied by one producer using one production process” [3, p. 1]. The lot generating production process should be “stable (under statistical control)” [3, p. 1]. The basic tool of type B control are the *switching rules* by which ISO 3951 intends to provide “automatic protection to the consumer should a deterioration in quality be detected” and “an incentive to reduce inspection costs should consistently good quality be achieved” [3, p. 1].

Table 3 ISO 3951-1 decision rules of a sampling plan (n, k) with sample measurements $X_1, \dots, X_n$

“s” method (process standard deviation σ unknown)	
Sample statistics	$\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$, $s = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2}$, $Q_L = \frac{\bar{X}-L}{s}$, $Q_U = \frac{U-\bar{X}}{s}$
Single lower limit L	$Q_L \geq k \Rightarrow$ accept, $Q_L < k \Rightarrow$ reject
Single upper limit U	$Q_U \geq k \Rightarrow$ accept, $Q_U < k \Rightarrow$ reject
Double limits $L < U$	ISO 3951-1 [3, pp. 18–25]
“σ” method (process standard deviation σ known)	
Sample statistics	$\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$, $Q_L = \frac{\bar{X}-L}{\sigma}$, $Q_U = \frac{U-\bar{X}}{\sigma}$
Single lower limit L	$Q_L \geq k \Rightarrow$ accept, $Q_L < k \Rightarrow$ reject
Single upper limit U	$Q_U \geq k \Rightarrow$ accept, $Q_U < k \Rightarrow$ reject
Double limits $L < U$	1. Determine f_σ from Table 6, ISO 3951-1 [3, p. 43] (see Table 6) 2. Calculate $\text{MPSD} = (U - L)f_\sigma$ 3. $\sigma > \text{MPSD}$: process variability unacceptable, discontinue inspection 4. $\sigma \leq \text{MPSD}$: accept, if $Q_L \geq k$ and $Q_U \geq k$; otherwise reject

ISO 3951-1 [3, pp. 10–11] permits type A inspection of isolated lots as a secondary control objective. The usefulness of this approach is dubious. There are principal reservations against variables sampling in a type A sense (*see* the discussion in **Single Sampling by Attributes and by Variables**).

Quality Indices: AQL and LQ

The quality indices of ISO 3951 and their definitions are the same as for ISO 2859: the AQL (*acceptance quality limit*), and the LQ (*limiting quality*). The AQL and the LQ are analyzed by **Attributes Sampling Schemes in International Standards**. In the master tables, the 16 AQL values range from 0.01 up to 10.0% (see Table 4). This is exactly that part of AQL values in MIL-STD-105 E and ISO 2859-1 intended for inspection for the proportion nonconforming.

Choice between Variables and Attributes Sampling

The specification range model explained by Table 1 can also serve as basis for the application of an attributes sampling standard as described by **Attributes Sampling Schemes in International Standards**. ISO 3951-1 [3, pp. 11–12], gives some guidelines for the choice between attributes and variables sampling.

Issues listed in favor of variables sampling: (a) Variables sampling is more informative since the underlying measurements X are recorded, and not only the binary conforming-nonconforming indicator. (b) Variables sampling enables smaller sample sizes for OC (operating characteristic) matched sampling plans, i.e., sampling plans which have approximately the same OC function. To illustrate this effect, ISO 3951-1 [3, p. 33] contains a table of OC matched sample sizes. (c) Variables sampling can easily be coordinated with control charts for continuous measurements.

Issues listed in favor of attributes sampling: (a) The implementation of attributes sampling is much simpler. (b) The quality model of variables sampling in the standards is restricted to at least approximately normally distributed continuous measurements X . The methods are inappropriate for skewed distributions of X . (c) Variables sampling has apparent peculiarities which may irritate the user, e.g., rejection based on a sample without nonconforming items.

The issue (c) is most serious. **Single Sampling by Attributes and by Variables** discusses inconsistencies of variables sampling in more detail. Particular inconsistencies are provoked by a type A view of variables sampling for the inspection of single isolated lots.

The structures of ISO 2859-1 and ISO 3951-1 are synchronized. Thus it is possible to start the inspection process under attributes sampling by ISO 2859-1, and to switch to variables sampling under ISO 3951-1 at a later stage when the necessary requirements, in particular approximately normally distributed measurements X , can be confirmed.

The Structure of ISO 3951-1

The structure of ISO 3951-1 is very similar to the structure of ISO 2859-1 which is analyzed by **Attributes Sampling Schemes in International Standards**.

ISO 3951-1 as a Sampling System

The hierarchy of sampling plans, sampling schemes, and sampling systems is discussed by **Attributes Sampling Schemes in International Standards**. Like ISO 2859-1, ISO 3951-1 is a sampling system with various inspection modes and guidelines or rules for choosing or switching between the modes. The inspection modes of ISO 3951-1 are explained by Table 2. Different from ISO 2859-1, ISO 3951-1 contains **single sampling** plans only.

The Contents of ISO 3951-1

The material provided by the print version of ISO 3951-1 can be classified into four parts. (a) sections 1 through 8: statements of purpose, definitions, general Explanations; (b) section 9 until 22, p. 31: advice on implementation and operation; (c) p. 32 until p. 75: tables and graphs; (d) p. 76 until end: normative and informative annexes.

To a considerable degree, the tables in parts (c) and (d) of ISO 3951-1 correspond to tables in ISO 2859-1. The table of sample-size code letters [3, p. 32.] is exactly the same as in ISO 2859-1. The master tables on pages 34 through 42 [3] are arranged in a way analogous to the master tables of ISO 2859-1. The tables are classified by the inspection state (normal,

Table 4 Master table of single sampling plans (n, k) for “s” method under normal inspection. Down arrows prescribe first sampling plan below arrow, up arrows prescribe first sampling plan above arrow

Sample size code letter	Sample size	Acceptance quality limit (% nonconforming)															
		0.010	0.015	0.025	0.040	0.065	0.10	0.15	0.25	0.40	0.65	1.0	1.5	2.5	4.0	6.5	10.0
		k	k	k	k	k	k	k	k	k	k	k	k	k	k	k	k
B	3													↓	0.954	0.818	0.526
C	4												↓	1.163	1.046	0.853	0.580
D	6											↓	1.395	1.275	1.108	0.902	0.587
E	9										↓	1.615	1.494	1.338	1.159	0.907	0.597
F	13									↓	1.830	1.712	1.565	1.405	1.189	0.938	0.614
G	18								↓	2.025	1.910	1.770	1.622	1.429	1.212	0.944	0.718
H	25								↓	2.102	1.969	1.829	1.652	1.457	1.225	1.035	0.809
J	35						↓	2.399	2.289	2.160	2.028	1.862	1.684	1.476	1.311	1.118	0.912
K	50						↓	2.569	2.461	2.336	2.209	2.052	1.885	1.693	1.543	1.372	1.193
L	70						↓	2.631	2.510	2.389	2.239	2.082	1.904	1.766	1.611	1.451	1.238
M	95						↓	2.736	2.670	2.553	2.410	2.261	2.093	1.822	1.676	1.484	↑
N	125		↓	3.037	2.937	2.824	2.711	2.574	2.432	2.274	2.154	2.021	1.886	1.710	↑		
P	160	↓	3.179	3.082	2.973	2.865	2.733	2.574	2.597	2.447	2.209	2.083	1.921	↑			
Q	200	3.310	3.215	3.109	3.004	2.877	2.747	2.603	2.495	2.377	2.258	2.106	↑				
R	250	3.350	3.247	3.146	3.023	2.898	2.760	2.657	2.545	2.432	2.289	↑					

tightened, reduced) and the inspection method (“s” method, “ σ ” method). As an example, consider the master table for normal inspection under the “s” method displayed by Table 4.

The Implementation and Operation of ISO 3951-1

The implementation of ISO 3951-1 requires five subsequent decisions.

1. Opting for variables sampling instead of for attributes sampling. See the comparison of variables and attributes sampling, above.
2. Choosing between the “s” method and the “ σ ” method. See the discussion of this choice within the framework of the quality model, above.
3. Selecting the AQL value. The choice of a threshold for tolerable process quality is essentially a business management task which requires a review and synthesis of all relevant aspects of technical, engineering, contractual, and economic aspects. The choice is hampered by the vagueness in the definition of the AQL (see the discussion in **Attributes Sampling Schemes in International Standards**).
4. Choosing among the inspection levels S-1, S-2, S-3, S-4 (special), I, II, III (general). As for MIL-STD-105E and ISO 2859-1, the inspection levels affect the sample sizes which is smallest for level S-1, and grows until level III. ISO 3951-1 [3, p. 13] suggests the use of level II “unless special circumstances indicate that another level is more appropriate”.

When the choices 1 through 4 have been made, the further steps – selection of the inspection severity state, of the specific sampling plan (n, k) , and the

operation of (n, k) – are uniquely determined by the operational regime of ISO 3951-1 according to the following algorithm. (a) Determine the inspection state from the switching algorithm (see Table 5). (b) Determine the lot size N . (c) Use N and the inspection level to determine the sample-size code letter X from the respective code letter table (see **Attributes Sampling Schemes in International Standards**). (d) Go to the relevant master table of sampling plans which is identified by the sampling mode and inspection state (see Table 4, for instance). In the table, identify the line with code letter X . In this line, find the sampling plan (n, k) in the column labeled by the respective value of the AQL. The 16 AQL values range from 0.01 up to 10.0%. (e) Operate the sampling plan (n, k) according to the algorithm explained by Table 3. The “s” method with double limits is rather involved, the reader is referred directly to ISO 3951-1 [3, pp. 18–25].

In case of double limits $L < U$, the algorithm contains a test on the process standard deviation, (see Table 3). If the process standard deviation σ is too large with respect to the length $U - L$ of the specification range, sampling cannot find out whether the process mean is inadequately positioned or not. In this case, inspection has to be discontinued until the process variability is reduced to a tolerable degree (see Table 6).

Supporting Procedures

The procedures of ISO 3951-1 are based on assumptions on the measurements X : normality, process variance stability, and, for the “ σ ” method, full knowledge of the process variance. These assumptions should be checked continuously when lots are submitted for inspection. ISO 3951-1 [3, p. 29] recommends

Table 5 Switching algorithm of ISO 3951-1. The rules “normal $\rightarrow$ tightened”, “tightened $\rightarrow$ normal”, “tightened $\rightarrow$ discontinued” are the same as in ISO 2859-1 (see **Attributes Sampling Schemes in International Standards**)

Normal	All of the conditions 1, 2, 3, 4 are satisfied:
↓	1. Ten preceding lots on normal inspection, none rejected
Reduced	2. Ten preceding lots would have been accepted under next smaller AQL
	3. Production is in statistical control
	4. The responsible authority accepts reduced inspection
Reduced	At least one of the conditions 1, 2, 3, 4 is satisfied:
↓	1. A lot is rejected
Normal	2. Production is irregular or delayed
	3. The responsible authority prefers normal inspection

6 Variables Sampling Schemes in International Standards

Table 6 Values f_σ for inspection under double specification limits in the “ σ ” method

Acceptance quality limit (% nonconforming)															
0.010	0.015	0.025	0.040	0.065	0.10	0.15	0.25	0.40	0.65	1.0	1.5	2.5	4.0	6.5	10.0
0.125	0.129	0.132	0.137	0.141	0.147	0.152	0.157	0.165	0.174	0.184	0.194	0.206	0.223	0.243	0.271

Table 7 Control states of ISO 3951-2

Single control	Only a single limit L or U considered, one AQL imposed on p_L or p_U , respectively
Combined control	Double limits $L < U$, same impact of $X < L$ and $X > U$, AQL imposed on $p_{L,U}$
Separate control	Double limits $L < U$, different impact of $X < L$ and $X > U$, AQL _{L} imposed on p_L , AQL _{U} imposed on p_U
Complex control	1. Double-lower: double limits $L < U$, different impact of $X \notin [L; U]$ and $X < L$, AQL _{L,U} imposed on $p_{L,U}$, AQL _{L} imposed on p_L 2. Double-upper: double limits $L < U$, different impact of $X \notin [L; U]$ and $X > U$, AQL _{L,U} imposed on $p_{L,U}$, AQL _{U} imposed on p_U

that control charts of the statistics $\bar{X}$ and s be kept. Additionally, one may keep histograms of the sample measurements to check the assumption of normality.

The Design of ISO 3951-1

The design of ISO 3951-1 replicates the design of ISO 2859-1. The analysis of the design of ISO 2859-1 in **Attributes Sampling Schemes in International Standards** applies implicitly to ISO 3951-1.

A Survey of ISO 3951-2

ISO 3951-2 as an Extension of ISO 3951-1

ISO 3951-2 is an extension of ISO 3951-1. The material and procedures of ISO 3951-1 are a proper subset of the material and procedures of ISO 3951-2. There are three major extensions. (a) Quality model: An arbitrary number $m \geq 1$ of independent continuous quality characteristic is admitted, i.e., on each item m independent measurements $X^{(1)}, \dots, X^{(m)}$ may be taken. Separate specification limits and specification ranges are associated with each of the m quality characteristics. (b) Control model: In the case of double

limits $L < U$, the option of different impacts of nonconformity by an excess below the lower limit L and of nonconformity by an excess over the upper limit U . This leads to four control states of *single*, *combined*, *separate*, and *complex control* as explained by Table 7. Each of the four control states can hold for each of m quality characteristics. (c) Sampling plan and decision algorithm: ISO 3951-2 contains the sampling plans (n, k) of ISO 3951-1 and the corresponding decision rule, see Table 3, for a single quality characteristic under single or combined control as the so-called “form k acceptability criterion”. For multiple quality characteristics and/or separate or complex control, ISO 3951-2 contains sampling plans (n, p^*) as the “form p^* acceptability criterion”.

The Implementation and Operation of ISO 3951-1.

To a large extent, the implementation and operation of ISO 3951-2 and ISO 3951-1 are the same (see the explanations above). The following extensions of the steps 1 through 4 are necessary. Step 2.: Choose between “ s ” and “ σ ” method for each of m quality characteristics separately. Step 3. For each quality characteristic with double specification limits, choose the appropriate control model from

Table 7, and select the corresponding AQL values. For each quality characteristic, the inspection level, the inspection severity state, and the table of sample-size code letters is handled as in ISO 3951-1.

The most important differences of ISO 3951-1 and ISO 3951-2 arise after the basic choices 1 through 4, while operating the master tables and the sampling plan, as it is necessary to handle multiple quality characteristics and the four control states from Table 7. The algorithm proceeds in the following steps 1 through 7.

1. Assign each quality characteristic $l \in \{1, \dots, m\}$ to one of the four control states $C_1, \dots, C_4$.
2. Specify the set $p_{0,1}, \dots, p_{0,r}$ of AQLs which are imposed to the quality characteristics.
3. Associated with the $p_{0,1}, \dots, p_{0,r}$, form r AQL classes $A_1, \dots, A_r$ of quality characteristic indices in the following way: (a) each characteristic $l \in C_1$ (single control) is put into the AQL class A_i associated with the imposed AQL $p_{0,i}$; (b) each characteristic $l \in C_2$ (combined control) is put into the AQL class A_i associated with the imposed AQL $p_{0,i}$; (c) each characteristic $l \in C_3$ (separate control) is put into the two AQL classes A_i, A_j associated with the AQL $p_{0,i}$ imposed on p_L and the AQL $p_{0,j}$ imposed on p_U ; (d) each characteristic $l \in C_4$ (combined control) is put into the two AQL classes A_i, A_j associated with the AQL $p_{0,i}$ imposed on $p_{L,U}$ and the AQL $p_{0,j}$ imposed on p_L or p_U , respectively.
4. From the p^* acceptability criterion master table identified by the inspection severity state, find the *acceptability value* p_i^* for each AQL $p_{0,1}, \dots, p_{0,r}$. The master tables are essentially organized like Table 4 where k is replaced by p^* .
5. For each class A_i and each $l \in A_i$ calculate an estimator $\hat{p}_{l,i}$ of the relevant probability nonconforming. The estimator differs for “s” method and “σ” method.
6. For each class A_i calculate $\hat{p}_i = 1 - \prod_{l \in A_i} (1 - \hat{p}_{l,i})$.
7. If all $\hat{p}_i \leq p_i^*$, accept, otherwise, reject.

ISO 3951-2 lacks a precise guideline on the relation of sampling and quality characteristics. It is unclear whether (a) one sample is taken, and all m quality characteristics are measured on each item in the sample, or whether (b) for each quality characteristic a separate sample is taken. Different

sample sizes for different quality characteristics can occur under the combined “s” and “σ” method where the process standard deviation is unknown for some and known for some other characteristics. Known process standard deviation requires smaller sample size. In this case, under method (a) the measurements would just be cancelled for the “σ” characteristics on some items. The basic orientation should be the requirement of independence of the m quality characteristics. Method (a) is appropriate if the m measurements $X^{(1)}, \dots, X^{(m)}$ taken on one item are independent. If this is not guaranteed, i.e., if the quality characteristics on one item may be correlated, method (b) has to be used.

ISO 3951-5

ISO 3951-5 [5] contains **sequential sampling** schemes that supplement the single sampling schemes provided by ISO 3951-1. The schemes are restricted to the case of known process standard deviation σ .

References

- [1] United States Department of Defense. (1957). *Military Standard, Sampling Procedures and Tables for Inspection by Variables for Percent Defective (MIL-STD-414)*, U.S. Government Printing Office, Washington, DC.
- [2] American Society for Quality. (2003). *Sampling Procedures and Tables for Inspection by Attributes (ANSI/ASQ Z1.9)*, ASQ Quality Press, Milwaukee.
- [3] International Organization for Standardization. ISO 3951-1. (2005). Sampling procedures for inspection by variables – Part 1: specification for single sampling plans indexed by acceptable quality limit (AQL) for lot-by-lot inspection for a single quality characteristic and a single AQL, International Organization for Standardization, Geneva.
- [4] International Organization for Standardization. ISO 3951-2. (2006). Sampling procedures for inspection by variables – Part 2: general specification for single sampling plans indexed by acceptable quality limit (AQL) for lot-by-lot inspection of independent quality characteristics, International Organization for Standardization, Geneva.
- [5] International Organization for Standardization. ISO 3951-5. (2006). Sampling procedures for inspection by variables – Part 5: sequential sampling plans indexed by acceptance quality limit (AQL) for inspection by variables (known standard deviation), International Organization for Standardization, Geneva.

SAMINATHAN BALAMURALI, RAINER GÖB AND
CHI-HYUCK JUN

Rectification Sampling Schemes

Introduction

Acceptance sampling plans usually require corrective action when the lots are rejected. The corrective action can be 100% inspection or screening of rejected lots where all observed nonconforming items are either removed for subsequent rework or returned to the supplier or replaced from a stock of known good items. Such **sampling schemes** are called *rectifying inspection schemes*, because the control activity affects the final product quality after control. Rectifying inspection schemes intend to restrict the long-run proportion nonconforming after control. This intention may be of interest both to the supplier and to the consumer. Thus rectifying inspection schemes are applied in incoming inspection, in-process inspection, or final inspection to give assurance regarding the average quality of material used in the next stage of manufacture or regarding the average quality of final product shipped to consumer.

We consider rectifying inspection for proportion nonconforming under the basic attributes quality model, i.e., items are classified as conforming or nonconforming, lot quality is determined by the lot proportion nonconforming p , and process quality is determined by the process proportion nonconforming $\pi = \bar{p}$. Rectifying inspection under a variables (specification range) quality model can be considered analogously, see [1]. See **Sampling Inspection of Products** for a general exposition of quality models in statistical **product inspection**.

The Dodge–Romig sampling tables [2] are the best-known rectification sampling scheme. The sequential rectification scheme by Anscombe [3] requires less than 100% inspection of rejected lots. In theory, Anscombe’s scheme is more economical, but it is more complicated in implementation. We concentrate on the simple Dodge–Romig scheme. For an exposition of Anscombe’s scheme, the reader is referred to [1].

Several types of rectification procedures exist. The two most important are (a) *replacing procedure*: all nonconforming items discovered in the course of sampling, or, in the case of rejection, by screening the

lot are replaced by conforming items; (b) *removing procedure*: all nonconforming items discovered in the course of sampling, or, in the case of rejection, by screening the lot are removed. The replacing procedure is more common, and is assumed by the Dodge–Romig scheme.

Rectifying inspection can be considered in a type A framework where the decision target is a finite lot of items $1, \dots, N$, or in a type B framework where sampling and decision is directed toward the lot generating process (see **Sampling Inspection of Products** for the distinction of type A and type B procedures). Originally, the Dodge–Romig scheme is decidedly intended for type A inspection as expressed by Dodge and Romig [4]: “The purpose of the inspection is to determine the acceptability of individual lots submitted for inspection, i.e., sorting good lots from bad.” Later, the second edition [2] of the sampling tables distinguishes type A and type B, and considers rectifying inspection in a type B sense for the control of the process output, too. In the sequel we adopt the original type A view.

Performance Measures and Characteristic Quantities

The basic performance measure of any sampling plan S for lot proportion nonconforming is its *operating characteristic* (OC) function, i.e., the probability

$$L_S(p) = P(\text{accept} \mid \text{lot proportion nonconforming} = p) \quad (1)$$

as a function of the actual value p of the lot proportion nonconforming. For double sampling plans, additional important measures are the probabilities

$$A_{S,l}(p) = P(\text{accept in stage } l \mid \text{lot proportion nonconforming} = p) \quad (2)$$

of accepting in stage l , $l = 1, 2$. Under a strict type A view, the OC functions and acceptance probabilities have to be calculated from appropriate hypergeometric probabilities (see **Single Sampling by Attributes and by Variables** for single and **Multiple Sampling Plans** for multiple sampling). The approximation schemes developed by Dodge and Romig [2] are based on approximating the hypergeometric probabilities by suitable Poisson ones in the way indicated

2 Rectification Sampling Schemes

by **Single Sampling by Attributes and by Variables** and **Multiple Sampling Plans**. Dodge and Romig [2] use the exact hypergeometric formulas for the calculations of tables only.

Rectification plans try to control the proportion nonconforming *after* application of the sampling plan, but to this end they incur a potentially high *total* inspection amount. These two contrary aspects are expressed by three specific performance measures, the *average outgoing quality* (AOQ), the *average outgoing quality limit* (AOQL), and the *average total inspection* (ATI). Table 1 explains these performance measures and provides formulas for the AOQ and the ATI under single and double sampling plans. Simple closed formulas for the AOQL exist in special cases only. For AOQ and ATI formulas under multiple and sequential sampling, consider **Multiple Sampling Plans**.

The Dodge–Romig scheme uses two further characteristics: the *lot tolerance percent defective* (LTPD) and the *consumer's risk* (CR). These characteristics are not specific for rectification schemes (see **Sampling Inspection of Products** for an explanation on OC related characteristics).

The Dodge–Romig Designs of Rectification Plans

The Dodge–Romig sampling tables [2] provide two designs for sampling plans: LTPD design and AOQL design. Both designs are based on minimizing the ATI as a function of the sampling plan under a prescribed value of the LTPD or the AOQL, respectively. The

ATI is a clear economic target, whereas the LTPD and the AOQL have no direct economic interpretation. Accordingly, the Dodge–Romig design may be classified as *semieconomic*. The designs are applied to single sampling plans (n, c) and to double sampling plans $(n_1, n_2, a_1, a_2, r_1, r_2)$ with the additional restriction $r_1 = r_2 = a_2 + 1$ (see **Single Sampling by Attributes and by Variables** and **Multiple Sampling Plans** for general expositions on single and multiple sampling plans).

As defined in Table 1, the ATI depends on the unknown lot proportion nonconforming p . Some prior knowledge on the lot proportion nonconforming is required. Dodge and Romig assume the so-called *process average proportion nonconforming* $\bar{p}$ which is defined as the expectation $E[P] = \bar{p}$ of the random lot proportion nonconforming P , see [2, p. 25]. In the lot and process quality model of **Sampling Inspection of Products**, $\bar{p} = \pi$ is the process proportion nonconforming. In a rigorous prior knowledge approach (see **Prior Information in Sampling Schemes** for a survey), one should consider the expected value $E[ATI(P)]$ of the random $ATI(P)$. Instead, Dodge and Romig [2] consider just $ATI(\bar{p})$. Various attempts have been made to justify this approach, in particular, considering $E[ATI(P)] \approx ATI(\bar{p})$ as a sufficiently good approximation under reasonable **prior distributions** of P , or assuming a one-point prior with outliers, see [5, p. 100].

The rationale of sampling inspection under knowledge of the expected proportion nonconforming $\bar{p} = \pi$ in the lot generating process is seriously questioned by the economic analyses of Feigenbaum [6]

Table 1 Characteristic quantities of a rectification sampling plan for lot proportion nonconforming

Generic definitions	
AOQ	Average outgoing quality. The expected value of the lot proportion nonconforming after application of a rectification plan as a function $AOQ(p)$ of the lot proportion nonconforming p before inspection.
AOQL	Average outgoing quality limit. Maximum $p_L = AOQL$ of the AOQ, i.e., $p_L = AOQL = \max_{0 \leq p \leq 1} AOQ(p)$.
ATI	Average total inspection. The expected value of inspected items, including 100% inspection in case of rejection.
Single sampling plan (n, c) with OC $L(p)$	
AOQ(p)	$AOQ(p) = p \left(1 - \frac{n}{N}\right) L(p) = p \frac{N - ATI(p)}{N}$
ATI(p)	$ATI(p) = n + (1 - L(p))(N - n)$
Double sampling plan $(n_1, n_2, a_1, a_2, r_1, r_2)$ with OC $L(p)$ and with probabilities $A_l(p)$ of terminating with acceptance in stage l , $l = 1, 2$	
AOQ(p)	$AOQ(p) = p \left\{ \left(1 - \frac{n_1}{N}\right) A_1(p) + \left(1 - \frac{n_1 + n_2}{N}\right) A_2(p) \right\}$
ATI(p)	$ATI(p) = n_1 A_1(p) + (n_1 + n_2) A_2(p) + N (1 - L(p))$

or Deming [7] (see **Sampling Inspection of Products**). In practice, however, process parameters are never completely stable. Even for a process under statistical control there may be occasional outliers in the process parameters. Protection against the effect of such outliers on product quality is a reasonable motivation for sampling inspection, see the arguments by Hald [5]. It remains to be analyzed, however, whether this motivation justifies the amount of sampling inspection imposed by the Dodge–Romig scheme. Unfortunately, the scheme gives no economic guidelines to answer this question.

The Dodge–Romig LTPD Design

The LTPD is considered as a relatively poor value p_t of p . Lots with $p = p_t$ should rarely be accepted. The Dodge–Romig LTPD design imposes the restriction $L_S(p_t) = 0.10$. From the consumer's perspective, the LTPD can be considered as the *consumer's quality* and the restriction $L_S(p_t) = 0.10$ as imposing a CR of 10% (see **Sampling Inspection of Products**). The Dodge–Romig LTPD plan is defined as the sampling plan S^* , which minimizes $ATI_S(\bar{p})$ among all sampling plans S satisfying $L_S(p_t) = 0.10$.

For the solution of this problem, Dodge and Romig [2] made two contributions: (a) An approximate solution based on approximating the exact hypergeometric acceptance probabilities by a Poisson distribution. See [5] for a detailed description of the derivation. Charts and tables support the implementation of the approximate solution. The algorithm is described by Schilling [8] in detail. (b) Tables of the exact solution based on the exact hypergeometric acceptance probabilities arranged as follows: one separate table for

each of the eight LTPD values 0.5, 1.0, 2.0, 3.0, 4.0, 5.0, 7.0, and 10.0%, where each table is arranged in six intervals for the process average $\bar{p}$ and where the maximum of $\bar{p}$ is half the respective LTPD. For larger process averages, sampling inspection is rather inappropriate compared to rejection without inspection. Lot size ranges from 1 to $N = 100\,000$. To understand the structure of the Dodge–Romig LTPD tables consider Tables 2 and 3, which provide an excerpt of the Dodge–Romig AOQL tables. The LTPD tables are arranged analogously.

Hald [5] analyzes the approximate solution and the tables in detail. The quality of the approximate solution decreases with increasing values of the LTPD p_t and of the fraction $\bar{p}/p_t$. In the tables, larger deviations from the exact solution occur only if the process average is close to the boundaries of the six intervals.

The Dodge–Romig AOQL Design

The essential objective of rectifying inspection is to impose a specified AOQL as an upper bound for the AOQ, i.e. the expected proportion nonconforming after application of the sampling plan. This objective is expressed by the AOQL design. A value $p_L = \text{AOQL}$ is prescribed. The Dodge–Romig AOQL plan is defined as the sampling plan S^* which minimizes $ATI_S(\bar{p})$ among all sampling plans S satisfying $\text{AOQL}_S = p_L$.

Quality cannot be inspected into a product, i.e. sampling plans cannot improve upon the quality provided by the lot generating process. Accordingly, prescribing an AOQL $p_L < \bar{p}$ smaller or equal to the process proportion nonconforming is absurd. If

Table 2 Dodge–Romig table of single sampling plans under AOQL design for an AOQL of 0.1%

Lot size	Process average 0–0.002%			Process average 0.003–0.020%			Process average 0.021–0.040%			Process average 0.041–0.060%			Process average 0.061–0.080%			Process average 0.081–0.100%		
	n	c	$pt\%$	n	c	$pt\%$	n	c	$pt\%$	n	c	$pt\%$	n	c	$pt\%$	n	c	$pt\%$
1–75	All	0	–	All	0	–	All	0	–	All	0	–	All	0	–	All	0	–
76–95	75	0	1.50	75	0	1.50	75	0	1.50	75	0	1.50	75	0	1.50	75	0	1.50
⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮
501–600	230	0	0.86	230	0	0.86	230	0	0.86	230	0	0.86	230	0	0.86	230	0	0.86
601–800	250	0	0.81	250	0	0.81	250	0	0.81	250	0	0.81	250	0	0.81	250	0	0.81
801–1000	270	0	0.76	270	0	0.76	270	0	0.76	270	0	0.76	270	0	0.76	270	0	0.76
⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮
50 001–100 000	365	0	0.58	835	1	0.46	1350	2	0.40	2480	4	0.33	3070	5	0.32	4270	7	0.30

4 Rectification Sampling Schemes

Table 3 Dodge–Romig table of double sampling plans $(n_1, n_2, a_1, a_2, r_1, r_2)$ under AOQL design for an AOQL of 0.1%. Acceptance numbers $c_1 = a_1$, $c_2 = a_2$, rejection numbers $r_1 = r_1 = a_2 + 1 = c_2 + 1$

Lot size	Process average 0–0.002%						Process average 0.003–0.020%						Process average 0.021–0.040%					
	Trial 1		Trial 2			$p_t\%$	Trial 1		Trial 2			$p_t\%$	Trial 1		Trial 2			$p_t\%$
	n_1	c_1	n_2	$n_1 + n_2$	c_2		n_1	c_1	n_2	$n_1 + n_2$	c_2		n_1	c_1	n_2	$n_1 + n_2$	c_2	
1–75	All	0	–	–	–	–	All	0	–	–	–	–	All	0	–	–	–	–
76–95	75	0	–	–	–	1.5	75	0	–	–	–	1.5	75	0	–	–	–	1.5
⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮
301–350	190	0	–	–	–	0.96	190	0	–	–	–	0.96	190	0	–	–	–	0.96
351–400	225	0	95	320	1	0.86	225	0	95	320	1	0.86	225	0	95	320	1	0.86
⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮
601–800	310	0	155	465	1	0.71	310	0	155	465	1	0.71	310	0	155	465	1	0.71
801–1000	350	0	185	535	1	0.66	350	0	185	535	1	0.66	350	0	185	535	1	0.66
⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮
50 001–100 000	670	0	770	1440	2	0.42	740	0	1230	1970	3	0.38	805	0	1725	2530	4	0.35

engineering or economic reasoning suggests a value p_L as a coercive bound for outgoing quality, and if $p_L < \bar{p}$, then lots from the process should be rejected without inspection. Throughout, the AOQL design assumes $p_L \geq \bar{p}$.

For the numerical implementation of the AOQL design [2] restrict attention to approximations of the exact hypergeometric acceptance probabilities by Poisson probabilities. Two contributions are made: (a) A simply calculable approximate solution, see [5] for a detailed description of the derivation. Charts and tables support the implementation of the approximate solution. The algorithm is described by [8] in detail. (b) Tables of the approximate solution arranged as follows: one separate table for each of the 13 AOQL values 0.1, 0.25, 0.5, 0.75, 1.0, 1.5, 2.0, 2.5, 3.0, 4.0, 5.0, 7.0, and 10.0%, where each table is arranged in six intervals for the process average $\bar{p}$ and where the maximum $\bar{p}$ is the respective AOQL. Lot size ranges from 1 to $N = 100\,000$. For illustration, Tables 2 and 3 display an excerpt of the Dodge–Romig AOQL tables for single and double sampling under a prescribed AOQL of 0.1%. Hald [5] analyzes the quality of the approximate solution and of the tables in detail.

Issues in the Implementation of the Dodge–Romig Scheme

The Dodge–Romig scheme is implemented in the following steps:

1. Decide between LTPD and AOQL plans.
2. Specify values of the LTPD or AOQL, respectively.
3. Specify the process average proportion nonconforming.
4. Decide between single and double sampling plans.
5. Determine the sampling plan from tables or from approximation formulas or from software.

The Dodge–Romig system is implemented in several production management packages, e.g. in a package distributed by Rockwell Automation. The final step 5 is simple and can easily be executed by shopfloor practitioners after suitable training. The first four steps have to be discussed.

Decision between LTPD and AOQL Tables

Originally, Dodge and Romig [9], subsumed LTPD plans under lot quality protection, AOQL under average quality protection. This distinction has often been interpreted in the sense that LTPD plans focus on the quality of single lots, whereas AOQL plans focus on the quality of a stream of lots. However, as stated later by Dodge and Romig in their amendments to the second edition of the sampling tables [2], both types may be considered in a type A view (focus on lots) or in a type B view (focus on the process). The different objectives of LTPD and AOQL plans are rather related to the distinction between *submitted*

and *outgoing* product: LTPD plans protect against accepting bad submitted quality, AOQL plans restrict the proportion nonconforming of outgoing product after inspection.

Specifying the LTPD or the AOQL

Both the AOQL and the LTPD are clearly defined from a formal point of view. However, the problem is to make these concepts meaningful and manageable in practice, in particular to relate them to the economy of **quality control**. The Dodge–Romig scheme gives no distinct economic guidelines on how to choose values of the AOQL or LTPD. Dodge and Romig [2, p. 5] recommend to the manufacturer the selection of an AOQL equal to the “acceptable quality level” specified by the consumer. Dodge and Romig [2, p. 6] recommend to “consider and compare the cost of inspection with the economic loss that would ensue if quality as bad as the LTPD were accepted often (or if the *average level* of quality were greater than the AOQL)”. In the end, these advices amount to relate the LTPD and the AOQL to the *break-even quality* in an economic profit–loss model for inspection (*see Sampling Inspection of Products*). However, the Dodge–Romig scheme contains no guidelines to specify such a model.

Specifying the Process Average

For a reasonable **estimation** of $\bar{p}$, knowledge on the process and data from the process are necessary.

- (a) The lot generating process is basically in a state of control, with only occasional out-of-control situations.
- (b) Quality records from the lot generating process exist. The estimation of $\bar{p}$ is not straightforward, since out-of-control outliers have to be excluded from the estimation. Dodge and Romig [2, p. 32] suggest a proportion nonconforming chart procedure:
 1. Plot the proportion nonconforming figures from the last 10 to 25 lots, excluding “nonrepresentative” lots with a strong out-of-control suspect.
 2. Verify a basic state of statistical control by comparing the data with 3σ limits.
 3. Replace points beyond 3σ limits by corresponding 2σ limits.

4. Use the average of the data as an initial estimate of $\bar{p}$.
5. Start the sampling inspection process. Update the initial estimate of $\bar{p}$ successively by using the estimates provided by the samples or screening procedures.

Small changes in the process average will have no effect on the sampling plan, see the efficiency studies by [5]. In particular, the tables provided by [2] (see Tables 2 and 3 above) are organized in intervals of the process average, so that only an excess over the interval boundaries will affect the design. Hald [5, p. 105] recommends that “in case of doubt the process average should be overestimated thus choosing a larger **sample size** than necessary because this leads to a smaller increase in the ATI than an underestimation of the same size.”

Deciding between Single and Double Sampling

The ordinary performance measure of multiple sampling is the average sample number (ASN) (*see Multiple Sampling Plans*). In the Dodge–Romig scheme, under given process average and given LTPD or AOQL, respectively, the ASN of single sampling is not in general smaller than the ASN of double plans. Under some constellations, even the first stage double sample size is larger than the single sampling size (compare Tables 2 and 3). However, for rectifying inspection, the interesting measure is not the expected total *sampling* amount, but the expected total *inspection* amount, i.e. the ATI. Dodge and Romig [2, pp. 30–31] demonstrate that the single sampling ATI consistently exceeds the double sampling ATI. The pattern is exactly the same as for the ASN-based comparison (*see Multiple Sampling Plans*): the advantage of double sampling is greatest for very small process proportion nonconforming. However, Dodge and Romig [2] reasonably attract attention to added costs from the administration of double sampling, naming “interruption of work” and “extra handling of product”.

Appraisal of the Dodge–Romig Scheme

We evaluate the usefulness of the Dodge–Romig scheme in modern industrial quality control. *See Acceptance Sampling in Modern Industrial Environments* for a general description of challenges

for statistical product inspection in modern quality control.

The Dodge–Romig scheme has two virtues, particularly in comparison with the Military Standard: (a) a definite, explicitly stated and transparent design; (b) *via* the process average, a mechanism that conveys previous information on product quality. In particular, a mechanism that reduces sample size in the case of improving product quality.

The following aspects of the Dodge–Romig scheme are unfavorable in a modern quality control environment: (a) the core design parameters LTPD and AOQL are not motivated by an economic model, in particular they are not clearly related to economic break-even qualities; (b) the range of the LTPD and AOQL considered in the tables—smallest LTPD at 0.5%, smallest AOQL at 0.1% – are much too high for modern rigorous product quality requirements; (c) the rationale of sampling under the assumption of a known process quality parameter is not made totally clear. Sampling inspection is not evaluated in comparison to the no sampling alternative.

References

- [1] Duncan, A.J. (1986). *Quality Control and Industrial Statistics*, 5th Edition, Irwin Publishing, Homewood.
- [2] Dodge, H.F. & Romig, H.G. (1959). *Sampling Inspection Tables*, 2nd Edition, John Wiley & Sons, New York.
- [3] Anscombe, F.J. (1949). Tables of sequential inspection schemes to control fraction defective, *Journal of Royal Statistical Society. Series A* **112**, 180–206.
- [4] Dodge, H.F. & Romig, H.G. (1929). A method of sampling inspection, *Bell System Technical Journal* **7**(4), 613–531.
- [5] Hald, A. (1981). *Statistical Theory of Sampling Inspection by Attributes*, Academic Press, London.
- [6] Feigenbaum, A.V. (1961). *Total Quality Control*, McGraw-Hill Book Company, New York.
- [7] Deming, W.E. (1986). *Out of the Crisis*, MIT Press, Cambridge.
- [8] Schilling, E.G. (1983). *Acceptance Sampling in Quality Control*, Marcel Dekker, New York.
- [9] Dodge, H.F. & Romig, H.G. (1941). Single sampling and double sampling inspection tables, *Bell System Technical Journal* **20**(1), 613–631.

RAINER GÖB

Continuous Sampling

Introduction

Acceptance sampling procedures are classified into type A and type B, *see Sampling Inspection of Products*. Type A procedures operate on a lot-by-lot basis: Incoming or outgoing lots of product are inspected to decide between accepting or rejecting a lot. This lot-by-lot framework is rather inappropriate if, in a strict type B view, a manufacturing process successively generating items shall be controlled. In this case, two reasons oppose the building of lots from the produced items: (a) Arranging the lots would be costly and time consuming. (b) The decision about the actual process quality would be unduly delayed. For example, if products like personal computers are produced on an assembly line, any disturbance shall be detected quickly. Then it makes no sense to wait for controlling the process until a lot has been built up. In this case, continuous item-by-item sampling plans are more appropriate.

Item-by-item sampling plans for continuous control of production processes were developed by Dodge [1]. The quality of the process is measured by the process nonconformance probability, i.e., by the probability that a manufactured item is nonconforming. A continuous sampling plan (CSP) consists of alternating sequences of screening (100% inspection) and sampling inspection. Details of the operation of CSPs are described subsequently. CSPs are rectifying inspection plans where any nonconforming item is reworked or replaced by a conforming one. Thus, product quality after control is better than before.

Dodge CSP-1 Plan

The CSP-1 plans introduced by Dodge [1] in 1943 are specified by the so-called clearance number i , the sampling frequency f , and the following instruction: start with screening, i.e., inspecting each item produced. As soon as i consecutive items have been found to be conforming, sampling inspection is started, where one out of $k = 1/f$ consecutively produced items is selected and inspected. If a sampled item is found to be nonconforming, revert back to screening. Inspection is rectifying, i.e., all nonconforming items found are either replaced

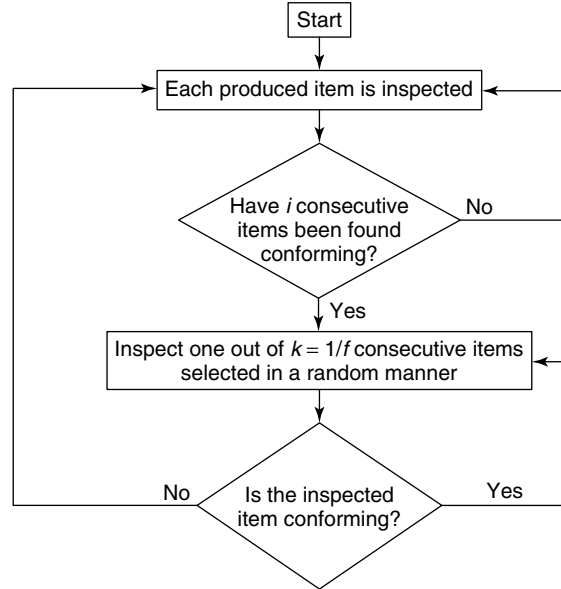


Figure 1 Operating procedure of CSP-1

by conforming ones or reworked. The operating procedure for CSP-1 is shown in Figure 1.

Properties of CSP-1

The choice of the clearance number i and the sampling fraction f is usually based on practical considerations about the manufacturing process and the required statistical properties of the sampling plan. The statistical properties of CSP-1 are derived in [1] under the assumption that the nonconforming probability of the process is constant and given by p , see also [1, 2].

- The average number of items inspected in a 100% screening sequence following the detection of a nonconforming item is equal to

$$u = \frac{1 - (1 - p)^i}{p(1 - p)^i} \quad (1)$$

- The average number of nonconforming items passing, while sampling inspection is performed, is given by

$$v = \frac{1}{fp} \quad (2)$$

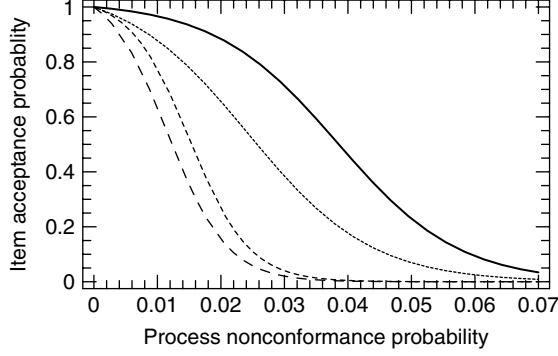


Figure 2 OC curves for the following four cases (from top to bottom): (a) $f = 0.02, i = 100$, (b) $f = 0.08, i = 100$, (c) $f = 0.04, i = 210$, and (d) $f = 0.08, i = 210$

- The average fraction of items (AFI) inspected in the long run is

$$AFI = \frac{u + fv}{u + v} \quad (3)$$

- The average outgoing quality (AOQ) is obtained from equation (3) as

$$AOQ = p(1 - AFI) \quad (4)$$

- The probability of accepting an item produced when using the CSP-1 sampling procedure, is denoted by $L_{i,f}(p)$ and defines the operating characteristic (OC) curve of a continuous sampling plan, CSP-1. It is given by

$$L_{i,f}(p) = \frac{v}{u + v} = \frac{(1 - p)^i}{f + (1 - f)(1 - p)^i} \quad (5)$$

The OC curves for some values of f and i for CSP-1 are displayed in Figure 2. It can be shown that for moderate-to-small values of f , the clearance number i has a larger effect on the shape of the curve than does f .

Derman *et al.* [3] investigates the properties of CSPs without the assumption of a constant nonconformance probability p .

Extensions of CSP-1

Sometimes the detection of a single nonconforming item does not warrant returning to process screening, for example when dealing with minor defects.

To overcome this difficulty, Dodge and Torrey [4], proposed the continuous sampling plans, CSP-2 and CSP-3 as extensions of CSP-1. Under CSP-2, process screening is reinstated only if under sampling inspection two nonconforming items have been found within a sequence of K sample units. It is common practice to choose K equal to the clearance number i . CSP-2 plans are indexed by their average outgoing quality level (AOQL). Note that a given AOQL does not determine the pair (i, f) uniquely, but may be obtained by different combinations of i and f . The continuous sampling plan, CSP-3 is very similar to CSP-2, but it is designed to give some additional protection against spotty production. It requires that after a nonconforming item has been found in sampling inspection, the immediately following next four units in succession should be inspected. If any of these four units is nonconforming, process screening is resumed immediately. If no nonconforming items is found among the inspected four items, the plan continues as under CSP-2. Multilevel CSPs are proposed and analyzed by Liebermann and Solomon [5] and Derman *et al.* [3].

MIL-STD-1253C

MIL-STD-1253C contains five different types of continuous sampling plans: CSP-1, CSP-2, CSP-F (CSP-1 modified for short production runs), CSP-V (CSP-1 with an option to reduce the clearance number), and CSP-T (three level CSP scheme with three levels of the sampling frequency, namely, f , $f/2$, and $f/4$). Tables assist the user in designing a sampling plan.

The sampling plans in MIL-STD-1253C are indexed through the sampling frequency code letter, which is related to the production speed, and the $AOQL$. Additionally, they are also indexed through the acceptable quality level (AQL) of MIL-STD-105E. However, CSP plans are not AQL plans and do not have an AQL value naturally associated with them. CSP plans are designed for situations in which production is continuous and considering lots is not meaningful. Therefore, in MIL-STD-1253C, the sampling plan tables are footnoted and indicate that the AQL values have no meaning relative to the plan, but are only used as an index.

Note that the properties of the plans offered by MIL-STD-1253C assume that the process is in a state of statistical control, i.e., a constant process

nonconformance probability. If this assumption is not met sufficiently well, the sampling plan properties may differ largely from those given in the standard.

References

- [1] Dodge, H.F. (1943). A sampling plan for continuous production, *Annals of Mathematical Statistics* **14**, 264–279.
- [2] Schilling, E.G. (1982). *Acceptance Sampling in Quality Control*, Marcel Dekker, New York.
- [3] Derman, C., Johns, M.V. & Liebermann, G.J. (1959). Continuous sampling procedures without control, *Annals of Mathematical Statistics* **30**, 1175–1191.
- [4] Dodge, H.F. & Torrey, M.N. (1951). Additional “continuous sampling inspection plans”, *Industrial Quality Control* **7**, 7–11.
- [5] Liebermann, G.J. & Solomon, H. (1955). Multi-level “continuous sampling plans”, *Annals of Mathematical Statistics* **26**, 686–704.

Further Reading

Derman, C., Littauer, S. & Solomon, H. (1957). Tightened multi-level continuous sampling plans, *Annals of Mathematical Statistics* **28**, 395–404.

KODAKANALLUR KRISHNASWAMY SURESH
AND ELART VON COLLANI

Skip Lot and Chain Sampling

Skip-Lot Sampling Plans

The **continuous sampling** plan (CSP-1) aims at controlling the outgoing quality of a continuous production process by inspecting individual items. In [1], Dodge proposed to apply the operation rules of CSP-1 to a continuous flow of lots and called the resulting procedure *skip-lot sampling plan* (SkSP-1). The goal is the same as for the CSPs, namely to quit in certain situations 100% control and, thus, reduce inspection costs. Dodge recommended the use of SkSPs in case that the flow of lots is produced in successive batches by a reliable production process or if the tests to be performed are time consuming and expensive.

The operating procedure for SkSP-1 is the same as for CSP-1 with the only differences that the word “item” has to be replaced by the word “lot” and the demand to “inspect an item” by the somewhat vague demand “test the lot”. The basic relations stated for CSP-1 for individual items hold also in case of SkSP-1 for lots. In [1] a nomogram is given for selecting a SkSP-1 plan guaranteeing a given average outgoing quality limit (AOQL), which is interpreted as the maximum fraction of accepted lots that will be nonconforming, where a lot is called *conforming* if the acceptance criterion of the test is met and otherwise it is called *nonconforming*.

Perry proposed in [2] a SkSP of type SkSP-2, where each lot to be inspected is sampled according to some lot **acceptance sampling** plan the so-called reference plan. Then, SkSP-2 is operated as follows:

1. Start with normal sampling inspection (inspecting every lot), using the reference plan.
2. When i consecutive lots are accepted on normal inspection, switch to skipping inspection (inspect a fraction f , $0 < f < 1$, of the lots).
3. When a lot is rejected on skipping inspection, switch to normal inspection.

Perry [2] derived the operating characteristic (OC) function of a SkSP-2 plan and studied the properties of SkSP-2 plan in terms of the average sample number (ASN), the average outgoing quality (AOQ), and the **average run length** (ARL). Finally, Perry [2]

analyzes the AOQ for a SkSP-2 plan. He defines the AOQ as fraction of outgoing lots that are nonconforming, where a nonconforming lot is one “whose sample fails (or would fail, if inspected) to meet the acceptance criterion of the reference plan”. Perry [3] also developed three different types of two-level SkSPs, which are based on the operating procedures of the multilevel plans (MLPs) for continuous sampling of Lieberman and Solomon [4] and of the tightened MLPs of Guthrie and Johns [5]. A detailed description of SkSP-2 plan and the two-level SkSPs is given in Stephens [6].

There are a number of shortcomings with respect to skip-lot sampling as described in literature. A first one refers to the fact that the properties are derived assuming constant nonconformance probability, a condition, which generally is not met in practice. A second severe weakness is the definition of nonconforming lot, which actually characterizes the reference sampling plan, but does not define the nonacceptability of a lot.

Chain Sampling Plans

The concept of chain sampling was developed by Dodge [7] as an alternative to attribute **single sampling** plans with zero acceptance number. Chain sampling plan (ChSP-1) makes not only use of the result of the actual sample, but also of the results concerning a number of preceding lots. The plan avoids rejection of lot on the basis of one single nonconforming item resulting in a relatively poor discrimination **power** between good and bad quality that marks the zero-acceptance number single sampling plan. The operating procedure of a ChSP-1 is as follows.

- For each lot, select a random sample of size n items and test each unit for conformance with the specifications.
- Accept the lot if d , the observed number of nonconforming items among the n sampled items, is zero; if $d > 1$, reject the lot.
- If the sample contains exactly one nonconforming item ($d = 1$), accept the lot provided there have been no nonconforming unit in the previous i samples of size n .

Dodge [7] proposed to apply chain sampling in circumstances where the successive lots are produced under similar conditions and are inspected in the

2 Skip Lot and Chain Sampling

order of production. Such situations may arise in receiving inspection of continuing supply of purchased materials produced within a manufacturing plant. Chain sampling is also recommended, when the probability of producing a nonconforming item is very small and at the same time large samples are practicable. A ChSP might also be of advantage in cases where the inspection is costly requiring a small **sample size**, i.e., the acceptance number zero.

The advantage of chain sampling over single sampling plans with zero-acceptance number is that the acceptance probability does not drop as rapidly for small values of the nonconformance probability p . A detailed review and a discussion on conditions for applications and measures of performance of the ChSP-1 plan are given in Stephens [6], Schilling [8], and Govindaraju [9]. Dodge [7], and Clark [10] studied the properties of ChSP-1 plans based on the OC curves for various combinations of n and i . Note that the OC curve for ChSPs is determined by considering an infinite series of lots produced with constant nonconformance probability p .

A generalized family of “two-stage” ChSPs, which extend the ChSP-1 plans, was developed by Stephens and Dodge [11] and presented in various technical reports and papers. The generalized two-stage ChSP with different sample sizes in the two stages is denoted by ChSP- $(n_1, n_2; k_1, k_2; c_1, c_2)$. The operating procedure and other properties of the general two-stage ChSP is derived in Dodge and Stephens [12] and Stephens and Dodge [11].

Check Inspection and Demerit Rating Sampling

The check inspection and demerit rating plan was introduced by Dodge and Torrey [13] for processes with a tight **process control** resulting in a high and stable quality level, making lot acceptance sampling by the consumer appear unnecessary. In such a situation check inspection is employed for process quality assurance purposes. The inspection shall be carried out by an inspection agency on the manufacturer’s premises from the customer’s point of view. Samples are selected at random from the completed product allowing to control and to maintain process quality at the desired high level. The product to be tested by check inspection refers generally to items which have a number of components with several quality characteristics.

Because of their complexity products are divided into classes. According to some class properties, the monthly sample size necessary for check inspection is taken from some empirically determined curves (see Dodge [13]) as function of the number of items produced per month. The monthly sample is divided into a number of smaller samples, which must be taken at random from the finished product so as to represent the entire output of the month.

Check inspection is particularly applied to complex products and, therefore, the number of nonconforming items or the number of defects in a sample appears to be not refined enough for reflecting quality. Therefore, a classification scheme for the seriousness of defects is used, where the seriousness is measured by so-called demerits per unit. Four classes of seriousness (Class A, B, C, D) are admitted in [13] with demerit weights of 100, 50, 10, and 1, where a higher weight means a higher seriousness of the defect. Based on the sample and some nonconformance criteria, a number of decisions can be made, which include acceptance, rejection, drawing an additional sample, inspect prior and subsequent product, etc.

Besides the sampling procedure described above, a monthly **control chart** is maintained on the basis of the cumulative number of demerits of the considered classes. This control chart reflects the development of the overall process quality and can be used for making decisions with respect to the process. The control limits are determined by empirical and practical considerations.

The possibilities of rating the seriousness of defects by demerits is investigated in many papers. Dodge [14] was one of the first to address the problem. A more recent investigation of the probability distribution of the considered demerits is given in [15].

The check inspection and demerit rating plans as presented in [13] are a purely heuristic inspection scheme without a theoretical basis. Therefore, it is impossible to make a reliable judgment of its performance in a given situation.

Deming’s (k_1, k_2) “All-or-None” Inspection Policy

For lot inspection Deming [16] stresses the necessity of determining the so-called breakeven quality, which distinguishes lots as being acceptable without inspection from lots which should be accepted only after

100% inspection. For defining the breakeven quality Deming introduces the cost to inspect one item with value k_1 , the cost caused by an accepted nonconforming item with value k_2 and the cost to replace a nonconforming item found during inspection with value k . Assuming $k \approx k_1$ Deming obtains the breakeven quality as the ratio (k_1/k_2) .

Clearly, if the production process has a nonconformance probability smaller than the breakeven quality, it makes sense, at least in the long run, to accept incoming lots without inspection. Similar, if the production process has a nonconformance probability larger than the breakeven quality, 100% inspection of all incoming lots is indicated from the long-run perspective. These two rules are known as *Deming's all-or-none inspection policy*. It assumes that the production process is in a state of statistical control and, moreover, the value of the nonconformance probability is known.

For achieving these two conditions, Deming demands that the process must be continuously observed using control charts on the basis of small samples, preferably in cooperation between supplier and purchaser. This is a similar situation as assumed for Dodge's check inspection and demerit rating plan.

For showing that sampling makes no sense when the process is in a state of statistical control with known quality level p , Deming – following Mood [17] – shows that the random numbers of nonconforming items in a sample and in the remainder, respectively, are uncorrelated, implying that the sample result does not yield any additional knowledge about the quality of the remainders. Thus, drawing a sample makes no sense from an economic point of view. This fact also illustrates the vital importance to reach a state of statistical control for the corresponding process. In order to highlight this demand, Deming [16] characterizes a process not under statistical control as “chaos”.

Deming's (k_1, k_2) inspection rule is important, because it attracts the attention to the questions of breakeven quality, which is widely neglected in the literature on lot acceptance sampling. Moreover, Deming stresses the relation between process distribution and lot distribution and the fact that quality is determined by the process and, therefore, statistical **quality control** should focus on the process.

Deming's all-or-none inspection rule has been investigated in many papers under various circumstances, see for example [18].

References

- [1] Dodge, H.F. (1955). Skip-lot sampling plans, *Industrial Quality Control* **11**, 3–5.
- [2] Perry, R.L. (1973). Skip-lot sampling plans, *Journal of Quality Technology* **5**, 123–130.
- [3] Perry, R.L. (1973). Two level skip-lot sampling plans – operating characteristic properties, *Journal of Quality Technology* **5**, 160–166.
- [4] Lieberman, N.J. & Solomon, H. (1955). Multi-level continuous sampling plans, *Annals of Mathematical Statistics* **26**, 686–704.
- [5] Guthrie, D. & Johns, M. (1958). Alternative sequences of sampling rates for tightened multi-level continuous sampling plans, Technical Report 36, Applied Mathematical and Statistics Laboratory, Stanford University.
- [6] Stephens, K.S. (1982). How to perform skip-lot and chain sampling, *ASQC Basic References in Quality Control*, American Society for Quality Control, Milwaukee, Vol.4.
- [7] Dodge, H.F. (1955). Chain sampling inspection plan, *Industrial Quality Control* **11**, 10–13. (Also in *Journal of Quality Technology*, (1977) 9:139–142.)
- [8] Schilling, E.G. (1982). *Acceptance Sampling in Quality Control*, Marcel Dekker, New York.
- [9] Govindaraju, K. (2006). Chain sampling, *Handbook of Engineering Statistics*, Springer, Chapter 15, ISBN:1-85233-806-7.
- [10] Clark, C.R. (1960). OC curves for ChSP-1 chain sampling plans, *Industrial Quality Control* **17**, 10–12.
- [11] Stephens, K.S. & Dodge, H.F. (1976). Two-stage chain sampling inspection plans with different sample sizes in the two stages, *Journal of Quality Technology* **8**, 209–224.
- [12] Dodge, H.F. & Stephens, K.S. (1966). Some new chain sampling inspection plans, *Industrial Quality Control* **23**, 61–67.
- [13] Dodge, H.F. & Torrey, M.N. (1956). A check inspection and demerit rating plan, *Industrial Quality Control* **13**, 5–12.
- [14] Dodge, H.F. (1928). A method of rating manufactured product, *Bell System Technical Journal* **VIII**, 350–358.
- [15] Chimka, J.R. & Cabrera Arispe, P.V. (2006). New demerit control limits for poisson distributed defects, *Quality Engineering* **18**, 461–467.
- [16] Deming, W.E. (1986). *Out of the Crisis*, Massachusetts Institute of Technology, Center for Advanced Engineering Study, Cambridge.
- [17] Mood, A.M. (1943). On the dependence of sampling inspection plans upon population distributions, *The Annals of Mathematical Statistics* **14**, 415–425.
- [18] Vander Wiel, S.A. & Vardeman, S.B. (1994). A discussion of all-or-none inspection policies, *Technometrics* **36**, 102–109.

SAMINATHAN BALAMURALI, CHI-HYUCK JUN
AND ELART VON COLLANI

Economic Sampling Schemes

Purpose of Sampling

In all areas of the human society, decisions must be made to cope with the uncertain future development. Whether a decision is appropriate or not depends on how well it takes account of the relevant initial conditions at hand, which affect the future development strongly. The main reason why there is uncertainty about the future development is the very fact that the relevant conditions are generally not fully, but only partially known. The overall aim of any industrial organization is to make profit. Therefore, any effort aiming at determining relevant initial conditions must be looked at from an economic viewpoint.

At this point **sampling schemes** come into play. They are used within a decision-making process for improving the knowledge about the relevant initial conditions, e.g., the quality of a lot of products. Such knowledge is indispensable for the decision process to avoid an economic loss by making a decision that is inappropriate in the given situation.

Problems with Sampling

There are two interrelated problems with sampling. The first one results from the fact that sampling costs money and constitutes, therefore, itself an economic loss. The second one is even worse. If sampling leads to false beliefs about the initial conditions, i.e., about the economic situation, then the decision will be wrong, for sure. In such a case the wrong decision incurs a loss, which is increased by the sampling cost. The relation between these two problems lies in the relation between sampling error and sampling cost: decreasing the probability of a wrong sampling result requires an increase in the **sample size** which leads to an increase in the sampling cost.

These considerations show that sampling is closely related to the cost situation and that any sampling plan not resulting in an increased economic yield should be called inappropriate. Therefore, selecting a sampling scheme for application without considering the given economic situation is hazardous and leads almost necessarily to an avoidable economic loss.

State of Economic Sampling Schemes

Economic sampling schemes, i.e., sampling schemes based on a thorough analysis of the economic impact of decisions and sampling cost, are more or less unknown in industry. There are several reasons for this surprising fact.

- There are no national or international standards for economic sampling schemes.
- Most of the existing economic sampling schemes are unattractive, because they are too complicated and are often based on unrealistic assumptions.
- Finally, selecting an appropriate economic sampling plan would necessitate a close cooperation including the exchange of relevant data by different departments of an organization. Such a cooperation is often not possible.

The first-mentioned point is explained by the origin of the existing national and international standards for sampling schemes in U.S. military standards, which were developed independently of any economic considerations. By adopting the military approach, the national and international organizations for standardization abandoned any opportunity for developing more target oriented, i.e., economic sampling schemes.

In statistical **quality control**, economic sampling schemes were first discussed and investigated in the realm of Bayes statistics, which required the specification of a prior probability distribution of the quality characteristic of the considered lot. A comprehensive overview on the resulting sampling schemes is given in [1], *see also Prior Information in Sampling Schemes*.

In industrial implementation, Bayesian sampling schemes face two serious obstacles related to knowledge. (a) Often the specification of a prior distribution fails because of a lack of *practical* knowledge. The prior distribution should characterize the process of assorting the lots, but this process is often completely unknown to the consumer wishing to select a sampling plan for incoming inspection. (b) A further obstacle is a lack of *theoretical* knowledge. Most practitioners do not understand the concept of probability sufficiently well and, therefore, the requirement of specifying one puts them off.

Bayesian economic sampling schemes represent one extreme with respect to the available knowledge.

2 Economic Sampling Schemes

The opposite extreme is represented by the so-called minimax regret sampling schemes, which are developed and comprehensively described in [2]. The disadvantage of these plans is the often unrealistic assumption that nothing is known about the underlying production process and the actual quality level of the lot to be inspected. Ignoring any available information leads to unnecessarily high sampling cost representing that part of the overall loss which will occur with certainty.

Finally, there is a more general economic approach: The so-called Π -minimax regret approach which includes both the Bayes and the minimax approaches as special cases. In the following the Π -minimax approach is briefly outlined for the case that the lot quality indicator is given by the proportion of nonconforming items P .

The Model

We consider the quality model of sampling by attributes introduced by **Sampling Inspection of Products**, i.e., item quality indicators of the conforming/nonconforming type, and proportion nonconforming p as the lot quality indicator. The methods described can also be applied to other quality models, see [3]. The aim of sampling is to decide about acceptance or rejection of a lot to increase the economic yield. Hence, the economic consequences must be modeled meeting the following natural conditions:

1. The average loss (profit = negative loss) in case of accepting a lot with fraction nonconforming p is given by a monotonously decreasing function $C_A(p)$.
2. The average loss in case of rejecting a lot with fraction nonconforming p is given by a monotonously increasing function $C_R(p)$.
3. It is more economic to accept a faultless lot $C_A(0) < C_R(0)$, and it is more economic to reject a lot consisting of only nonconforming items $C_A(1) > C_R(1)$.
4. The loss incurred by sampling with sample size n is given by an where a denotes the sampling cost per item.

The following considerations assume a **single sampling** plan (n, c) , where n denotes the sample size and c the acceptance number. If there are not more than c nonconforming items in a sample of

size n taken from the lot, then the lot is accepted, otherwise rejected. The sample is assumed to be taken at random without replacement. Hence, the probability of accepting a lot of size N with fraction nonconforming p is given by the hypergeometric operating characteristic (see **Single Sampling by Attributes and by Variables**).

$$L_{N,n,c}(p) = \sum_{m=0}^c \frac{\binom{Np}{m} \binom{N-Np}{n-m}}{\binom{N}{n}} \quad (1)$$

On the basis of the above quantities, the initial economic conditions are given by the following characteristics:

- The indifference (break even) quality level p_0 defined by

$$C_A(p_0) = C_R(p_0) \quad (2)$$

In case of $p < p_0$ it is more economic to accept the lot. In case of $p > p_0$ it is more economic to reject the lot. In case of $p = p_0$ acceptance and rejection are economically equivalent.

- The avoidable loss is defined as the loss incurred by wrong actions performed due to not knowing the true initial condition. The avoidable loss in case $p \leq p_0$ is given by:

$$e_A(p) = \begin{cases} C_R(p_0) - C_A(p_0) & \text{if the lot is rejected} \\ 0 & \text{if the lot is accepted} \end{cases} \quad (3)$$

- The avoidable loss in case $p > p_0$ amounts to

$$e_R(p) = \begin{cases} 0 & \text{if the lot is rejected} \\ C_A(p_0) - C_R(p_0) & \text{if the lot is accepted} \end{cases} \quad (4)$$

- The expected avoidable loss for acceptable lots ($p \leq p_0$) is given by

$$E_p[e_A(p)] = [C_R(p) - C_A(p)][1 - L_{N,n,c}(p)] \quad (5)$$

- The expected avoidable loss for unacceptable lots ($p > p_0$) is given by

$$E_p[e_R(p)] = [C_A(p) - C_R(p)]L_{N,n,c}(p) \quad (6)$$

- Finally the total avoidable loss denoted $C_{N,n,c}$ is

$$C_{N,n,c}(p) = \begin{cases} E_p[e_A(p)] + an & \text{for } 0 \leq p \leq p_0 \\ E_p[e_R(p)] + an & \text{for } p_0 \leq p \leq 1 \end{cases} \quad (7)$$

The aim is to determine a sampling plan (n, c) , which minimizes in some sense the total avoidable loss.

Π -Minimax Sampling Plans

The subsequently described approach was proposed in [4] and is based on [5] and [6]. Lot proportion nonforming is considered as a random variable P resulting from a process of manufacture, storage, transport etc. It is assumed that a set Π of possible probability distributions of P is given. A sampling plan (n^*, c^*) is called a Π -minimax plan, if

$$\sup_{\pi \in \Pi} R_\pi(n^*, c^*) = \inf_{n,c} \sup_{\pi \in \Pi} R_\pi(n, c) \quad (8)$$

holds, where

$$\begin{aligned} R_\pi(n, c) &= E_\pi[C_{N,n,c}(P)] \\ &= \int_0^{p_0} E[e_A(p)] d\pi(p) \\ &\quad + \int_{p_0}^1 E[e_R(p)] d\pi(p) + an \end{aligned} \quad (9)$$

From the point of view of stochastics, Π represents the prior knowledge on lot quality, see **Prior Information in Sampling Schemes**. Three important instances may occur.

Bayes assumption (complete knowledge): If $\Pi = \{\pi_B\}$, i.e., Π contains only one element, then the probability distribution π_B represents the prior distribution and the resulting sampling plan (n_0, c_0) is a Bayes plan.

No knowledge: If Π contains all possible probability distributions of the fraction nonconforming P , then the sampling plan (n_0, c_0) is a minimax plan.

Partial knowledge: Π is a proper subset of the subset of all probability distributions. Particularly interesting

is the Γ -minimax approach where

$$\Pi = \Gamma = \left\{ \pi \mid \int_0^{p_0} d\pi(p) = \alpha \right\} \quad (10)$$

The assumption (10) means that the probability of obtaining an acceptable lot is known. This assumption is often, at least approximately fulfilled, if there is a long-term business connection between producer and consumer. Define the α -minimax cost G^* by

$$\begin{aligned} \sup_{\pi \in \Gamma} R_\pi(n, c) &= \alpha \max_{0 \leq p \leq p_0} E[e_A(p)] \\ &\quad + (1 - \alpha) \max_{p_0 \leq p \leq 1} E[e_R(p)] + an \\ &= G^*(N, p_0, \alpha; n, c) \end{aligned} \quad (11)$$

A sampling (n_0, c_0) is called a Γ -minimax plan, if it minimizes the α -minimax cost, i.e., if

$$G^*(N, p_0, \alpha; n^*, c^*) = \min_{n,c} G^*(N, p_0, \alpha; n, c) \quad (12)$$

The α -Optimal Sampling Scheme

An interesting modification of the Γ -minimax scheme is the α -optimal sampling scheme proposed in [7]. The starting position is the same as in the case of the Γ -minimax scheme. However, the α -optimal sampling scheme takes account of the fact that often the consumer has to pay for the erroneously accepted lots, whereas the producer pays for the erroneously rejected lots. Therefore, a sampling plan should be selected for which the two costs are balanced. Thus, only sample plans (n, c) are admissible for the α -optimal sampling scheme, for which

$$\alpha \max_{0 \leq p \leq p_0} E[e_A(p)] \approx (1 - \alpha) E[e_R(p)] \quad (13)$$

holds.

The side condition (13) is met by admitting only that sample size n_c for given acceptance number c , which minimizes the difference

$$\begin{aligned} D(n) &= \left| \alpha \max_{0 \leq p \leq p_0} E[e_A(p)] \right. \\ &\quad \left. - (1 - \alpha) \max_{p_0 \leq p \leq 1} E[e_R(p)] \right| \end{aligned} \quad (14)$$

4 Economic Sampling Schemes

Sampling plans (n_c, c) minimizing equation (14) are called α -balanced sampling plans. An α -balanced sampling plan (n_c^*, c^*) is called an α -optimal sampling plan if it minimizes the α -minimax cost G^* :

$$G^*(N, p_0, \alpha; n_c^*, c^*) = \min_c G^*(N, p_0, \alpha; n_c, c) \quad (15)$$

In [7] an algorithm for determining approximate α -optimal sampling plans in the case of linear loss functions C_A and C_R is given.

Conclusion

As a matter of fact most of the research articles concerning economic sampling plans were written at least 20 years ago; for a review see [8]. Nowadays, articles investigating the economic aspect of sampling are published only occasionally. Motivated by the great popularity of Taguchi's "quadratic loss to society", more recently there have been a few papers related to **acceptance sampling** as, for instance, [9] or [10]. However, assuming a quadratic loss function and instead of justifying it, just calling it a "loss to society" does not make much sense. In fact, neither a producer nor a consumer represents society and each of them should take care of their own loss and not of that of society.

Increased globalization has led to the situation in which many companies operate production locations distributed all over the world with ever-increasing concentration in the so-called low-cost countries. This trend has increased significantly the importance of outgoing and incoming inspection. A considerable part of the overall costs is represented by inspection and related costs in many industries.

A reduction in these costs would considerably increase their competitiveness. Therefore, what is needed are user-friendly economic sampling schemes, which could and should replace the inefficient and costly AQL sampling plans still used in industries worldwide.

References

- [1] Hald, A. (1981). *Statistical Theory of Sampling Inspection by Attributes*, Academic Press, London.
- [2] Uhlmann, W. (1969). *Kostenoptimale Prüfpläne*, Physica-Verlag, Würzburg-Wien.
- [3] Göb, R. (1992). α -optimal sampling plans for lot-by-lot defects inspection, *Metrika* **39**, 269–316.
- [4] Krumbholz, W. (1982). Die Bestimmung einfacher Attributprüfpläne unter Berücksichtigung unvollständiger Vorinformation, *Allgemeines Statistisches Archiv* **66**, 240–253.
- [5] Hodges, J.L. & Lehmann, E.L. (1952). The use of previous experience in reaching statistical decisions, *Annals of Mathematical Statistics* **23**, 396–407.
- [6] Robbins, H. (1964). The empirical Bayes approach to statistical decision problems, *Annals of Mathematical Statistics* **35**, 1–20.
- [7] Collani, E. (1986). The α -optimal sampling scheme, *Journal of Quality Technology* **18**, 63–66.
- [8] Collani, E. (1990). Wirtschaftliche Qualitätskontrolle – eine Übersicht über einige neue Ergebnisse, *OR Spektrum* **12**, 1–23.
- [9] Moskowitz, H. & Tang, K. (1992). Bayesian variables acceptance sampling plans: quadratic loss function and step loss function, *Technometrics* **34**, 340–347.
- [10] Ferrell, W.G. & Chhoker, A. (2002). Design of economically optimal acceptance sampling plans with inspection error, *Computer and Operations Research* **29**, 1283–1300.

ELART VON COLLANI

Prior Information in Sampling Schemes

Introduction

The choice of a sampling scheme and the amount of inspection is guided, at least informally, by prior information on quality. Formal rules for adapting the rigor of inspection to past inspection records are, for example, the switching rules contained in many standard sampling schemes like MIL-STD-105E or ISO 2859, *see* **Attributes Sampling Schemes in International Standards** and **Variables Sampling Schemes in International Standards**. These are often based on heuristic arguments, but there are also rules derived from a strict mathematical treatment. An interesting recent example is an article by Bailie and Klaassen [1]. However, these authors do not model prior information explicitly either. Here we are concerned with the formal models of prior information. Together with a loss function, in a very broad sense, these may be used to design optimal sampling strategies or to analyze and compare the performance of existing sampling schemes.

In the “Bayesian approach”, prior knowledge is modeled in terms of a prior distribution. According to Hald [2], a complete statistical model of sampling inspection should contain the expected distribution of submitted lots according to quality (prior distribution) and the cost of sampling inspection, acceptance, and rejection. The best sampling plan is defined as the one within the class minimizing the average costs. This is essentially the paradigm of Bayesian decision theory, *see* [3]. However, the latter is not restricted to costs in a monetary sense. It only requires a numerical goal function that measures the consequences of possible decisions. Bayesian sampling plans will be introduced in the first section of this chapter. In practice, basic assumptions are often not satisfied. Therefore, modifications may be necessary and these will be addressed in the second section.

Bayesian Sampling Plans

A useful guide to practitioners on how and when to use Bayesian sampling plans has been described by Calvin [4]. The basic concepts are introduced in

the following example. We use the framework established by chapters **Sampling Inspection of Products** and **Single Sampling by Attributes and by Variables**: A lot of items $1, \dots, N$ is described by the sequence $z_1, \dots, z_N$ of item quality indicators. Lot quality q is a function of the item quality indicators. Let $z_{i1}, \dots, z_{in}$ be the item quality indicators in a sample of size n from the lot. A sample-based decision rule is represented by a function $\Delta(z_{i1}, \dots, z_{in}) = d$. A **single sampling** plan is given by (n, Δ) where the values d of Δ range over $d = a$ “accept” and $d = r$ “reject”. There are essentially two types of cost models: “type A” means that the costs of items in the sample do not depend on the decision about the remainder of the lot, whereas in “type B” the costs of defective items in the sample and those in the remainder of the lot depend in the same way on the final decision.

Introductory Example

In single sampling by attributes, the item quality indicators are binary variables where $z_i = 1$ means that the i -th item is defective (nonconforming) and $z_i = 0$ means that it is nondefective (conforming). The lot quality indicator is the number $m = z_1 + \dots + z_N$ of defective items in the lot, or the fraction defective $p = mN^{-1}$.

Linear Cost Model of Type B. As a simple example, we introduce a linear cost model of type B. For the interpretation of the parameters, *see* [2]. The cost of decision d per item with characteristic z is given by $k_d(z) = a_d + b_d z$, i.e., it is a_d if the item is conforming and $a_d + b_d$ if the item is nonconforming. The cost K_d of acceptance and rejection of the whole lot is the sum of the individual costs, it depends obviously on p :

$$\begin{aligned} K_d &= \sum_{i=1}^N k_d(z_i) = Na_d + mb_d = N(a_d + b_d p) \\ &= N k_d(p) = K_d(p) \end{aligned} \quad (1)$$

Moreover, there are sampling costs $n \cdot s$, where s is the sampling cost per item.

Prior Distribution and Bayesian Decision Rule.

The decision d is based on the number X of defective items in a random sample without replacement.

2 Prior Information in Sampling Schemes

For fixed m , its probability distribution $P[X = x|m]$ is hypergeometric. In addition, there is a prior distribution $\Psi(m)$ for m (prior to sampling), i.e., the number of nonconforming items in the lot is a random variable M with distribution Ψ . In a frequency interpretation, m varies from lot to lot according to Ψ .

Given a particular observation $X = x$, Bayes theorem allows us to calculate the posterior probability of $M = m$ as $P[M = m|X = x] = P[X = x|m]\Psi(m) \left(\sum_{l=0}^N P[X = x|l]\Psi(l) \right)^{-1}$. From this, the expected posterior loss of decision d can be derived as

$$E(K_d|X = x) = \sum_{m=0}^N K_d \left(\frac{m}{N} \right) P[M = m|X = x] \quad (2)$$

One should accept if $E(K_a|X = x) \leq E(K_r|X = x)$ and reject otherwise. It can be shown that there exists c such that $E(K_a|X = x) \leq E(K_r|X = x)$ if $x \leq c$, where $c = -1$ is admitted with the interpretation “rejection regardless of the outcome of the sampling”. This justifies using a sampling plan (n, c) . This is the first way of looking at Bayes: we have several possible decisions; given an observation, a decision is chosen, which has the smallest posterior loss. However, this does not answer the question which **sample size** n should be used. For this purpose, one has to compare the expected costs for different sampling plans, including sampling costs. For fixed m , the average loss of (n, c) is given by

$$R\left(n, c \middle| \frac{m}{N}\right) = n \cdot s + k_a \left(\frac{m}{N} \right) P[X \leq c|m] + k_r \left(\frac{m}{N} \right) P[X > c|m] \quad (3)$$

Considered as a function of p , $R(n, c|p)$ is called the risk function of (n, c) . Its expected value with respect to Ψ ,

$$R_\Psi(n, c) = E_\Psi \left[R\left(n, c \middle| \frac{m}{N}\right) \right] = \sum_{m=0}^N \Psi(m) R\left(n, c \middle| \frac{m}{N}\right) \quad (4)$$

is called the Bayes risk of using (n, c) . A Bayesian sampling plan (n_Ψ, c_Ψ) minimizes $R_\Psi(n, c)$. This constitutes the second way of looking at Bayes: a decision rule tells us which action to take depending

on the observation. We choose the “Bayesian rule” which minimizes the expected costs.

Priors Related to the Production Process, Additional Examples

Frequently, the distribution of the quality of the lot is related to the distribution of quality in the production process: if the items of lots of size N are drawn independently at random from a production process with constant probability π of producing a defective item, the number m of defective items in the lot varies according to the binomial frequency function $b(m|N, \pi)$. If π itself varies at random from lot to lot according to a distribution $W(\pi)$, the distribution of lot quality becomes a mixed binomial distribution

$$\Psi(m) = \int_0^1 b(m|N, \pi) dW(\pi) \quad (5)$$

This model is frequently mentioned in the literature: it allows analysis of various situations where sampling inspection is advantageous, and derivation of adequate strategies, see [2]. For example, W may represent a production process in statistical control with occasional out-of-control periods. The lots produced during these periods are considered as outliers, and if the product in the in-control state is satisfactory, the problem of sampling inspection is to detect the outliers. Formally, W can be taken as a beta distribution with outliers.

Cost Model of Type A under Single Sampling. In such environments, cost models of type A have a simple structure. Thyregod [5] provides a very elegant treatment, which is restricted to neither the attributes sampling nor the linear cost models. The item quality indicators $Z_1 \dots Z_N$ are assumed to be independent and distributed according to a density $f(z|\omega)$, where ω is the process parameter. This includes various models, in particular: (a) attributes sampling where $Z \in \{0, 1\}$, $\omega = P[Z = 1]$ and (b) sampling by variables where Z is normally distributed with expectation μ and standard deviation σ and $\omega = (\mu, \sigma)$. The parameter ω varies in the production process from lot to lot according to a distribution $W(\omega)$. The decision $d \in \{a, r\}$ is based on a sample of size n , which can be identified with $Z_1 \dots Z_n$. The cost of the decision depends only on the remaining $N -$

n item characteristics, it is the sum of the individual costs $k_d(z_j)$ of the items $j = n + 1, \dots, N$. The function $k_d(z)$ is not restricted to a special type; it includes the above cost model, slight generalizations (e.g., sampling by variables with upper specification limit $U : k_d(z) = a_d$, if $z \leq U$ (item good) and $= a_d + b_d$, if $z > U$ (item defective)) as well as nonlinear cost functions. The cost of sampling is the sum of individual sampling costs $k_s(z_i)$ of the items in the sample.

For $i = a, r, s$, let $k_i^*(\omega) = \int k_i(z) f(z|\omega) dz$ be the expectations of the costs with respect to $f(z|\omega)$. Moreover, let $k_s^* = \int k_s^*(\omega) dW(\omega)$ be the overall expected sampling costs. Since, for given ω , the n observations in the sample are independent from the remaining $N - n$ variables, the expected cost when using the sampling plan (n, Δ) is $R_W(n, \Delta) = n \cdot k_s^* + (N - n) r_W(n, \Delta)$, where

$$r_W(n, \Delta) = \int [P[\Delta = a|\omega] k_a^*(\omega) + P[\Delta = r|\omega] k_r^*(\omega)] dW(\omega) \quad (6)$$

Equation (6) has the formal structure of equations (3) and (4), with Ψ replaced by W and $R(n, c|p)$ replaced by the integrand in brackets in equation (3). This means that the original decision problem regarding the lot is formally equivalent to a decision problem regarding the process parameter ω with prior distribution W . A Bayesian sampling plan minimizes $R_W(n, \Delta)$. Thyregod [5] shows that it has usually a familiar structure: there is a test statistic t , and the lot is accepted if $t \leq c$. For example, in the case of a normally distributed quality characteristic with parameters μ and σ and upper specification limit U ,

$$t = \sqrt{n}(\bar{X} - U) S^{-1} \quad (7)$$

where S is the sample standard deviation.

Extensions and More Recent Results

The Bayesian approach is neither restricted to decision rules of this simple structure nor to monetary costs. For example, Bayesian multistage or **sequential sampling** plans may be designed to minimize the average sample size, which now plays the role of a risk function. Here the class of decision rules is much richer than above. A considerable amount of literature exists on Bayesian sampling,

and with increasing computer **power** more sophisticated models can be studied. For example, Gonzalez *et al.* [6] consider the number of defects per unit of product, whereas Chen *et al.* [7] introduce a Bayesian variable **acceptance sampling** scheme for life testing with mixed censoring. In both the articles, different types of nonlinear loss functions are used.

Modifications of Full Bayes

Robustness with Respect to Prior and Economic Model

Correct assignment of the prior distribution and a monetary cost model, as well as the stability of the prior are problematic issues of the Bayesian approach. General concepts of **Bayesian robustness** are discussed by Berger [3]. Bayesian sampling plans often have a poor discriminating power, which is a problem if the assumed prior distribution is misspecified or unstable. Therefore, Hald [2] recommends using restricted Bayesian sampling plans, where the Bayes risk is minimized within a smaller class of sampling plans with sufficient discriminating power. If full information about costs is not available, a partial cost function or the expected amount of inspection can be minimized under conventional constraints on the consumer's or the producer's risk, see [2]. The latter approach avoids the problem of estimating cost parameters completely. Bayes theorem also allows us to calculate posterior risks. In this spirit, Schafer [8] suggests to minimize the sample size under constraints on the posterior risk to the producer and the consumer risk.

Γ -Minimax

A concept which automatically yields more robust strategies is partial prior information, also called Γ -minimax. For example, Krumbholz and Pflaumer [9] derive sampling strategies which minimize expected costs subject to prior information on the probability that the fraction nonconforming does not exceed a certain upper bound. In a similar spirit, Collani [10] considers the so-called α -optimal sampling schemes, see **Economic Sampling Schemes**. These conditions are examples of "generalized moment conditions" of the prior. Other generalized or ordinary

moments can be estimated more easily from past inspections (see below), but the calculation of a corresponding sampling strategy is usually a complex “ Γ -minimax” problem. This can be solved in a numerically efficient way by semi-infinite programming, see [11].

Availability of Information on the Prior

If the fraction nonconforming p or the process parameter ω varies from lot to lot according to a prior Q , it is possible to gain information about Q from past inspections. However, one does not observe p or ω directly. The observable distribution is the mixture of the conditional distribution of the test statistic, given p or ω , with mixing distribution Q . Methods which use information about the observable distribution to derive a Bayesian decision rule are called **empirical Bayes**; for a general outline we refer again to [3], in particular, the section titled “Nonparametric Empirical Bayes Analysis”. Empirical Bayes has been applied to sampling inspection several times. Another way is to estimate ordinary or generalized moments of the prior indirectly from the observable distribution [12], and to use a Γ -minimax sampling plan. This method is able to incorporate information about the uncertainty of the estimator, for example, a confidence band around the empirical estimate of the marginal distribution.

Adaptive Sampling Strategies

A final issue is stability of the prior. Rendtel and Lenz [13] have developed a sampling system based on sequential tests, which is able to adapt to changing priors. Here it is not necessary to model these changes explicitly. There have also been attempts to model prior information by stochastic processes more general than independent and identically distributed (i.i.d.), but it seems that these have not gained much popularity.

References

- [1] Baillie, D. & Klaassen, C.A.J. (2006). Credit to and in acceptance sampling, *Statistica Neerlandica* **60**, 283–291.
- [2] Hald, A. (1981). *Statistical Theory of Sampling Inspection by Attributes*, Academic Press, London.
- [3] Berger, J.O. (1993). *Statistical Decision Theory and Bayesian Analysis*, 3rd Edition, Springer, New York.
- [4] Calvin, T.M. (1990). *How and When to Perform Bayesian Sampling*, Revised Edition, American Society for Quality.
- [5] Thyregod, P. (1974). Bayesian single sampling acceptance plans for finite lot sizes, *Journal of the Royal Statistical Society. Series B* **36**, 305–319.
- [6] González, C. & Palomo, G. (2003). Bayesian acceptance sampling plans following economic criteria: an application to paper pulp manufacturing, *Journal of Applied Statistics* **30**, 319–334.
- [7] Chen, J., Chou, W., Wu, H. & Zhou, H. (2004). Designing acceptance sampling schemes for life testing with mixed censoring, *Naval Research Logistics* **51**, 597–612.
- [8] Schafer, R.E. (1967). Bayes single sampling plans by attributes based on the posterior risk, *Naval Research Logistics Quarterly* **14**, 81–88.
- [9] Krumbholz, W. & Pflaumer, P. (1982). Möglichkeiten der Kosteneinsparung bei der Qualitätskontrolle durch Berücksichtigung von unvollständigen Vorinformationen, *Zeitschrift für Betriebswirtschaftslehre* **52**, 1088–1102.
- [10] von Collani, E. (1986). The α -optimal sampling scheme, *Journal of Quality Technology* **18**, 63–66.
- [11] Fandom Noubiap, R. & Seidel, W. (2001). An algorithm for calculating Γ -minimax decision rules under generalized moment conditions, *Annals of Statistics* **29**, 1094–1116.
- [12] Seidel, W. (1997). *Unbiased Estimation of Generalized Moments of Process Curves*, In *Frontiers in Statistical Quality Control, Vol. 5*, H.-J. Lenz & P.-Th. Wilrich, eds, Physica-Verlag, Heidelberg, pp. 21–37.
- [13] Rendtel, U. & Lenz, H.-J. (1990). *Adaptive Bayes'sche Stichprobensysteme für die Gut-Schlecht-Prüfung*, Physica-Verlag, Heidelberg.

WILFRIED SEIDEL

Reliability Sampling

Introduction

Reliability sampling is a part of **acceptance sampling** (see **Sampling Inspection of Products**) that is dedicated to the control of those quality characteristics of a product, which are related to its lifetime. In contrast with classical acceptance sampling, in reliability sampling procedures, such as lifetime tests, the quality of sampled items is evaluated using tests that last in time. A further distinction between classical acceptance sampling and reliability sampling stems from the fact that the latter is rather used for the control of process parameters than for the control of finite lots or batches of products. Therefore, popular methods of sampling by attributes are seldom used in reliability sampling. On the other hand, methods of sampling by variables, whose applicability for the acceptance sampling of finite lots is questionable, are very useful in the context of reliability sampling.

Statistical procedures of reliability sampling can be roughly divided into two general groups according to the quality model considered (see **Sampling Inspection of Products** for quality models in acceptance sampling). (a) Sampling procedures used for the control of a fraction (probability) nonconforming $\theta = p$ (probability of failing before prespecified time, probability of incorrectly performing a certain mission, probability of not passing a special test, etc.). The underlying item quality indicators X_i are binary variables where $X_i = 0$ indicates conformance and $X_i = 1$ indicates nonconformance. This case is considered in the section titled “Reliability Sampling by Attributes”. (b) Sampling procedures used for the control of reliability parameters related to time (mean time to failure (MTTF), mean time between failures (MTBF), constant hazard, or failure rate, etc.). The underlying item quality indicators X_i are real-valued variables expressing time to failure, time between failures, and so on. This case is considered in the section titled “Reliability Sampling for **Lifetime Distributions**”.

In reliability sampling, it is common practice to define two characteristic values of the quality parameter θ : an acceptable value θ_0 (often called *producer’s quality*) and an unacceptable value θ_1 . Reliability sampling plans are methods for testing the null hypothesis $H_0 : \theta \leq \theta_0$ against the alternative

$H_1 : \theta \geq \theta_1$. Let α be the probability of wrong rejection of H_0 (producer’s risk), β the probability of wrong acceptance of H_1 (consumer’s risk), and $L(\theta)$ the probability of accepting the null hypothesis as a function of the true value θ of the considered parameter. Then, the reliability sampling plan is designed in such a way that the two requirements $L(\theta_0) \geq 1 - \alpha$ and $L(\theta_1) \leq \beta$ are fulfilled. This design is well known in classical acceptance sampling (see **Single Sampling by Attributes and by Variables**). However, the difference between classical and reliability acceptance sampling plans lies in the way the statistical data are collected. We discuss these differences in detail in the next section.

Reliability Sampling by Attributes

The item quality model of acceptance sampling by attributes distinguishes conforming and nonconforming items (see **Single Sampling by Attributes and by Variables**). In the context of reliability, conformance is usually understood as the correct functioning during a certain period, correct performing of a certain action, or positively passing a special trial (e.g., an environmental trial). If conformance and nonconformance are clearly defined, all well-known sampling plans by attributes, like those described in the international standard ISO 2859-1 [1], may be applied. Specific attributes plans for reliability sampling can be found in the international standard IEC 61123 [2]. The basic reliability sampling plans presented in [2] are the curtailed sequential sampling plans by attributes (see **Multiple Sampling Plans** for information on sequential sampling). In contrast with the original sequential sampling procedures, a sample-size curtailment is introduced in the plans given in [2]. If the cumulative number of tested items reaches the curtailment value n_t and yet no decision on acceptance or rejection has been achieved, inspection terminates and the decision is then taken by comparing the observed number of nonconforming items, e.g., the number of failures, with specific curtailment values of the acceptance number. Parameters of the plans given in [2] are calculated using simple approximate formulas proposed by Wald [3] for non-curtailed sampling plans. Thus, the exact statistical characteristics of the curtailed sequential sampling plans given in this standard are different from the designed ones. Curtailed sequential sampling plans

by attributes based on exactly computed statistical properties can be found in the latest version of the international standard ISO 8422 [4]. In addition to curtailed sequential sampling plans, IEC 61123 [2] gives also reliability sampling plans where the number of trials (**sample size**) or the number of observed failures are predetermined. In all reliability sampling plans mentioned in this section, the design principles are those described in the previous section, i.e., they are based on two prespecified points on the operating characteristic (OC) function of the sampling plan.

Reliability Sampling for Lifetime Distributions

General Description

When reliability requirements are related to lifetimes (survival times) of items, reliability sampling procedures are different from those of classical sampling by variables. In classical sampling by variables the values of the quality characteristic of interest are observed for all items in the sample. In reliability tests, due to unavoidable limitations imposed on the possible time of the reliability test, not all lifetimes of tested items are observed. This feature of lifetime tests is known as *censoring*. Censoring arises in various ways. In the case of type II censoring, the test is terminated after the occurrence of the r th failure. Thus, the number of observed failures is fixed, but the length of the lifetime test is random. In such a case, the lifetimes of the remaining $n - r$ tested items are known only to exceed the time of the r th failure. In the case of type I censoring, the test is terminated after a predetermined time t_B . For this type of censoring the test time is fixed, but the number of observed failures is random. There exist also more complex censoring schemes like progressive censoring (see the section titled “Reliability Sampling Plans for Complex Censoring Schemes”).

The presence of censoring makes the design of reliability sampling plans very complicated. Only for exponentially distributed lifetimes, sampling plans are relatively easily designed. Moreover, the efforts of many researchers resulted in some practical procedures for the cases of the Weibull and the lognormal distributions of lifetimes. For other probability distributions of lifetimes, reliability sampling procedures are practically nonexistent.

Reliability Sampling for the Exponential Distribution

The exponential distribution, defined by the density function

$$f(t) = \frac{1}{\mu} \exp\left(-\frac{t}{\mu}\right) = \lambda \exp(-\lambda t), \quad t > 0 \quad (1)$$

is the most popular statistical model that has been used for the analysis of lifetimes. The reason for its popularity stems from the fact that it is completely defined by only one parameter: the expected value μ , or its reciprocal – the constant hazard rate $\lambda = 1/\mu$. This allows reliability analysts to estimate basic reliability characteristics when only few reliability data are available.

In the statistical analysis of lifetime data described by the exponential distribution, we use the fact that all information from a sample is subsumed by two statistics: the observed number of failures r , and the **total time on test** (total observed lifetimes) T . From a practical point of view this feature is very important: the same information comes from lengthy lifetime tests of few items and shorter tests of many items.

The most popular reliability sampling procedure is a sequential one, originally proposed by Epstein and Sobel [5]. Its curtailed version is described in the international standard IEC 61124 [6]. For this test the observed number of failures r is plotted against the cumulative (total) time on test. If the cumulative time on test T , for the observed number of failures r , is greater than or equal to an acceptance critical value T_a , the test is terminated, and the specified reliability requirement is regarded as being complied with. If, on the other hand, this time is smaller than or equal to a rejection critical value T_r , the test is terminated, and the specified reliability requirement is regarded as not being complied with. If time T , for the observed number of failures r , lies between these two critical values, no decision can be taken, and the test is continued. When the number of observed failures exceeds the curtailment value, the test is terminated, and the specified reliability requirement is regarded as not being complied with. The international standard [6] also provides fixed time/failure terminated reliability sampling plans. For these sampling plans decisions rules of accepting or rejecting compliance are made when the test time for termination has been reached, or the acceptable number of failures has been exceeded. The general guidance on

how to use these reliability sampling plans is given in the standard MIL-STD-781D [7]. The sampling plans of this American military standard are given in the handbook *MIL-HDBK-781A* [8]. The sampling plans given in [6] and [8] are practically the same.

Reliability Sampling for the Weibull and Other Lifetime Distributions

The Weibull distribution is described by the following density function

$$\frac{\beta}{a} \left(\frac{t}{a}\right)^{\beta-1} \exp \left[-\left(\frac{t}{a}\right)^{\beta} \right] t > 0$$

where $\beta > 0$ is the shape parameter, and $a > 0$ is the scale parameter of the distribution. When the value of β is known, we can transform the Weibull distribution to the exponential distribution using a simple transformation $Z = T^{\beta}$. In such a case, we can use simple reliability sampling plans designed for the exponential distribution. However, when the value of the shape parameter is not known, the design of reliability sampling plans is not so simple. Sequential sampling plans, which are the most popular reliability sampling plans, have not been proposed yet for the Weibull case. However, several authors proposed reliability sampling plans with type II censoring by the number of observed failures for the transformed lifetime variable $X = \ln T$, which is distributed according to the extreme value probability distribution. First plans of this type were proposed by Fertig and Mann [9] who used best linear unbiased estimators (BLUEs) for the **estimation** of the parameters of the distribution. The sample statistic is the standardized logarithm of the p th quantile of the Weibull distribution. The sampling plan is given by the sample size n , the censoring number of failures r , and the critical value k of the test statistic. Calculations are based on an approximation of the probability distribution of the sample statistic. Another reliability sampling plan was proposed by Schneider [10], who used the maximum-likelihood estimators (MLEs) of the parameters of the distribution. The decision criterion is typical for sampling plans by variables. The result of the reliability test is positive (acceptance) if $\hat{\mu} - k\hat{\sigma} \geq \ln L$, where $\hat{\mu}$ and $\hat{\sigma}$ are the estimators of the parameters $\mu = \ln a$ and $\sigma = 1/\beta$, respectively, k is the acceptability constant, and L is the lower specification limit for the lifetime. The parameters of

the sampling plan are computed from the corrected asymptotic distribution of the MLEs. Similar sampling plans have also been proposed for the lognormal and the gamma probability distributions.

Reliability Sampling Plans for Complex Censoring Schemes

In real reliability sampling experiments, certain items are removed from the test at its different stages, for instance, to investigate physical **degradation** processes that lead to future failures. When a prespecified number of tested items are withdrawn from the test at the **moments** of failures of some other items, the censoring scheme is known as *progressive censoring of type II*. The general methodology for the design of such reliability sampling plans was proposed by Balasooriya and coauthors. The case of the exponential distribution is considered in [11], and the case of the Weibull distribution is described in [12]. The results presented in these papers, and in other recently published papers on similar topics, are rather far from being applicable in practice.

Reliability Sampling Plans Based on the Results of Accelerated Life Tests

Reliability acceptance sampling plans for highly reliable components may be impractical due to the very long testing times required. In accelerated life tests (ALTs), introduced in the 1960s, items are tested at two or more stress levels that are higher than the normal one. In ALT failures are expected to occur earlier than in normal life tests, and thus the test time may be significantly reduced. There exist two general approaches in ALT. In overstress (or constant-stress) life tests a sample is divided in two or more subsamples tested at different stress levels. In step-stress life tests the whole sample is tested at one (usually lower) stress level for a certain time, and then the test is continued at the next (usually higher) stress level. The number of stress levels in test of both types may be greater than two, but reliability sampling plans based on ALT are usually designed for two levels of stress. Moreover, it is assumed that the functional form of the relationship between the expected lifetime and the stress level is known. The first acceptance sampling plans based on ALT were proposed in the 1990s, and since then many different procedures have been proposed in literature, mainly

for the exponential and the Weibull distributions. For example, constant-stress acceptance sampling plans were proposed by Yum and Kim [13], for the exponential distribution, and by Bai *et al.* [14], for the Weibull and lognormal distributions. For the exponential distribution, step-stress sampling plans were proposed by Bai *et al.* [15]. Chung *et al.* [16] have proposed such plans for the Weibull distribution.

Bayes Reliability Sampling Plans

The reliability sampling plans described in the previous sections of this article are based on purely statistical requirements. All relevant information comes from tested samples. Two extensions are possible: (a) inclusion of economic parameters, e.g., sampling costs, losses induced by wrong decisions, and so on; and (b) inclusion of certain prior information about the reliability of items, especially useful in case of highly reliable objects. Statistical procedures that use information of types (a) and (b) are named *Bayesian procedures*. In the area of reliability sampling an interesting general framework for the design of optimal Bayesian sampling plans in the case of exponential distribution was proposed by Dunsmore and Wright [17]. Bayesian **life test** plans for the Weibull distribution with given shape parameter are also given by Zhang and Meeker [18]. In both instances, however, ready-to-use tables of sampling plans do not exist, and therefore particular sampling plans are to be determined by a user using general algorithms given in the aforementioned papers.

References

- [1] ISO 2859-1 (1999). *Sampling Schemes Indexed by Acceptance Quality Limit (AQL) for Lot-By-Lot Inspection*.
- [2] IEC 61123 (1991). *Reliability Testing – Compliance Test Plans for Success Ratio*.
- [3] Wald, A. (1947). *Sequential Analysis*, John Wiley & Sons, New York.
- [4] ISO 8422 (2006). *Sequential Sampling Plans for Inspection by Attributes*.
- [5] Epstein, B. & Sobel, M. (1966). Sequential life tests in the exponential case, *Annals of Mathematical Statistics* **26**, 82–93.
- [6] IEC 61124 (1997). *Reliability Testing – Compliance Tests for Constant Failure Rate and Constant Failure Intensity*.
- [7] MIL-STD-781D (1986). *Reliability Testing for Engineering Development, Qualification, and Production*.
- [8] MIL-HDBK-781A (1996). *Handbook for Reliability Test Methods, Plans, and Environments for Engineering, Development Qualification, and Production*.
- [9] Fertig, K.W. & Mann, N.R. (1980). Life-test sampling plans for two-parameter Weibull populations, *Technometrics* **22**, 165–177.
- [10] Schneider, H. (1989). Failure-censored variables-sampling plans for lognormal and Weibull distributions, *Technometrics* **31**, 199–206.
- [11] Balasooriya, U. (1995). Failure-censored reliability sampling plans for the exponential distribution, *Journal of Statistical Computation and Simulation* **52**, 337–349.
- [12] Balasooriya, U., Saw, S.L.C. & Gadag, V. (2000). Progressively censored reliability sampling plans for the Weibull distribution, *Technometrics* **42**, 160–167.
- [13] Yum, B.-J. & Kim, S.-H. (1990). Development of life-test sampling plans for exponential distributions based on accelerated life testing, *Communications in Statistics: Theory and Methods* **19**, 2735–2743.
- [14] Bai, D.S., Chun, Y.R. & Kim, J.G. (1995). Failure-censored accelerated life tests sampling plans for Weibull distribution under expected test time constraint, *Reliability Engineering and System Safety* **50**, 61–68.
- [15] Bai, D.S., Kim, M.S. & Lee, S.H. (1989). Optimum simple step-stress accelerated life tests with censoring, *IEEE Transactions on Reliability* **38**, 528–532.
- [16] Chung, S.W., Seo, Y.S. & Yun, W.Y. (2006). Acceptance sampling plans based on failure-censored step-stress accelerated tests for Weibull distributions, *Journal of Quality in Maintenance Engineering* **12**, 373–396.
- [17] Dunsmore, I.R. & Wright, D.E. (1985). A decisive predictive approach to the construction of sequential acceptance sampling plans for lifetimes, *Applied Statistics* **34**, 1–13.
- [18] Zhang, Y. & Meeker, W.Q. (2005). Bayesian life test planning for the Weibull distribution with given shape parameter, *Metrika* **61**, 237–249.

OLGIERD HRYNIEWICZ

Variables Sampling under Measurement Error

Introduction

Variables sampling (*see Single Sampling by Attributes and by Variables*) assumes that the measurement error in measuring the characteristic under investigation is negligible, i.e., the standard deviation of measurement is no more than 10% of the process (or within-lot) standard deviation of the quality characteristic. However, especially in materials testing and in chemical analysis this assumption very often does not hold. In such cases the random variable Y describing the measurement result obtained at a sampled item can be modeled as the sum

$$Y = X + X_{\text{meas}} \quad (1)$$

where $X \sim N(\mu; \sigma^2)$ is the random variable describing the variation of the quality characteristic between items and $X_{\text{meas}} \sim N(0; \sigma_r^2)$ is the random variable, independent of X , describing measurement variation; its expectation is zero since it is assumed that the measurement process is unbiased. In the special case of no measurement error, we have $X_{\text{meas}} \equiv 0$ and hence, $Y = X$, i.e., the measured value y obtained at an item is equal to its “true” value x .

In what follows, only the *type B procedure* of **Single Sampling by Attributes and by Variables** is considered with the *one-sided* case where an upper limit (UL) is specified.^a The fraction of nonconforming items above UL in the process behind the lot is

$$\pi = P(X > UL) = 1 - \Phi\left(\frac{UL - \mu}{\sigma}\right) \quad (2)$$

where μ is the process mean, σ is the process standard deviation, and $\Phi(\cdot)$ is the cumulative distribution function of the standardized **normal distribution**. A **single sampling** plan for inspection by variables is a statistical test of the null hypothesis $H_0 : \pi \leq AQL := \pi_1$ against the alternative hypothesis $H_1 : \pi > \pi_1$ where AQL is specified acceptable quality level.

In [1] and [2] the effect of measurement error on the operating characteristic (OC) function of a single sampling plan for inspection by variables is investigated under the assumption that the ratio of

the measurement standard deviation to the process standard deviation is known. Wilrich [3] deals with the case where each sampled item is measured more than once and an upper bound for the ratio of the measurement standard deviation to the process standard deviation is available.

μ Unknown, σ^2 and σ_r^2 Known

If each of n_σ^* sampled items is measured m times as is often the case in materials testing and in chemical analysis ($m = 1$ is the special case of no repeated measurements), and if

$$\bar{\bar{Y}} = \frac{1}{n_\sigma^* m} \sum_{i=1}^{n_\sigma^*} \sum_{j=1}^m Y_{ij} \quad (3)$$

is the mean of the measurements X_{ij} ; $i = 1, \dots, n_\sigma^*$; $j = 1, \dots, m$ then an estimator of the fraction of nonconforming items, π , is

$$\hat{\Pi} = 1 - \Phi\left(\frac{UL - \bar{\bar{Y}}}{\sigma}\right) \quad (4)$$

If π_c is the critical value of the test, the lot will be accepted for

$$\hat{\Pi} \leq \pi_c \quad (5)$$

or equivalently

$$\frac{UL - \bar{\bar{Y}}}{\sigma} \geq u_{1-\pi_c} := c_\sigma^* \quad (6)$$

$$Z_\sigma := \bar{\bar{Y}} + c_\sigma^* \sigma \leq UL \quad (7)$$

where $u_{1-\pi_c}$ is the $(1 - \pi_c)$ -quantile of the standardized normal distribution.

For given parameters (n_σ^*, c_σ^*) of the single sampling plan the OC function is

$$\begin{aligned} L(\pi) &= P(Z_\sigma \leq UL) \\ &= \Phi\left(\sqrt{\frac{n_\sigma^*}{1 + \gamma_r^2/m}}(u_{1-\pi} - c_\sigma^*)\right) \end{aligned} \quad (8)$$

where

$$\gamma_r = \sigma_r / \sigma \quad (9)$$

is the ratio of the measurement standard deviation to the process standard deviation. A comparison

2 Variables Sampling under Measurement Error

with the case without measurement error (see **Single Sampling by Attributes and by Variables**, Table 3) shows that the OC is equal to the one without measurement error if $c_\sigma^* = c_\sigma$ and

$$n_\sigma^* = n_\sigma(1 + \gamma_r^2/m) \quad (10)$$

measurement error, if it exists, needs the **sample size** to be increased by the factor $(1 + \gamma_r^2/m)$ to keep the OC function of the plan unchanged. If the measurement standard deviation is not larger than 10% of the process standard deviation, i.e., $\gamma_r^2 \leq 0.01$, the influence of measurement error is negligible. As a larger m is chosen, as more vanishes the uncertainty caused by measurement error, and the sample size does not need to be increased very much.

The single sampling plan with the smallest sample size n_σ^* for which

$$L(\pi_1) \geq 1 - \alpha, \quad L(\pi_2) \leq \beta \quad (11)$$

where $P_1(\pi_1; 1 - \alpha)$ and $P_2(\pi_2; \beta)$ with $0 < \pi_1 < \pi_2 < 1$ and $0 < \beta < 1 - \alpha < 1$ are two given points on the **OC curve** which are given by

$$n_\sigma^* = \left(\frac{u_{1-\alpha} + u_{1-\beta}}{u_{1-\pi_1} - u_{1-\pi_2}} \right)^2 \left(1 + \frac{\gamma_r^2}{m} \right) \quad (12)$$

$$c_\sigma^* = \frac{u_{1-\beta}u_{1-\pi_1} + u_{1-\alpha}u_{1-\pi_2}}{u_{1-\alpha} + u_{1-\beta}} \quad (13)$$

u_p is the p -quantile of the standardized normal distribution; n_σ^* has to be rounded off to the next largest integer.

μ, σ^2 and σ_r^2 Unknown

In this section the process standard deviation σ and the measurement standard deviation σ_r are assumed to be unknown. The cases of one measurement per item and of $m > 1$ measurements per item are treated separately.

One Measurement per Item

If each of n_s^* sampled items is measured once ($m = 1$) then the test statistic is

$$Z_s := \bar{Y} + k_s^* S \quad (14)$$

where

$$\bar{Y} = \frac{1}{n_s^*} \sum_{i=1}^{n_s^*} Y_i, \quad S^2 = \frac{1}{n_s^* - 1} \sum_{i=1}^{n_s^*} (Y_i - \bar{Y})^2 \quad (15)$$

are sample mean and sample variance of the sample of size n_s^* , respectively.

For given parameters (n_s^*, c_s^*) of the single sampling plan, the OC function is

$$\begin{aligned} L(\pi) &= P(Z_s \leq UL) = P(T_{n_s^*-1, \delta} \geq c_s^* \sqrt{n_s^*}) \\ &= 1 - F(c_s^* \sqrt{n_s^*}) \end{aligned} \quad (16)$$

where

$$\delta = \sqrt{\frac{n_s^*}{1 + \gamma_r^2}} u_{1-\pi} \quad (17)$$

and $T_{n_s^*-1, \delta}$ is noncentral t-distributed with $\nu = n_s^* - 1$ **degrees of freedom** and noncentrality parameter δ . Unfortunately, the OC function depends on the ratio γ_r of the measurement standard deviation to the process standard deviation, a value which is hardly known, because the production process and the measurement process are completely different processes for which the relationship of the standard deviations would only be known if the standard deviations themselves are known.

If (n_σ, c_σ) are the parameters of the sampling plan for the case of no measurement error and known standard deviation which has the desired properties $L(\pi_1) \geq 1 - \alpha$ and $L(\pi_2) \leq \beta$ the matching plan for the case of unknown standard deviation and no measurement error ($\gamma_r = 0$) has the parameters $c_s = c_\sigma$ and

$$n_s = n_\sigma(1 + c_\sigma^2/2) \quad (18)$$

If the single sampling plan is designed for the case of no measurement error and a measurement error exists when it is applied, the probabilities of acceptance for all fractions π of nonconforming items in the process with $0 < \pi < 1$ will be smaller than without measurement error so that the requirement $L(\pi_1) \geq 1 - \alpha$ is not anymore fulfilled; see Figure 1. On the other hand, if the single sampling plan is designed for the case of measurement error, expressed as a given ratio γ_r , and measurement error is smaller during its application, i.e., the actual γ_r is smaller than the value used for the design, the requirement $L(\pi_2) \leq \beta$ is not fulfilled.

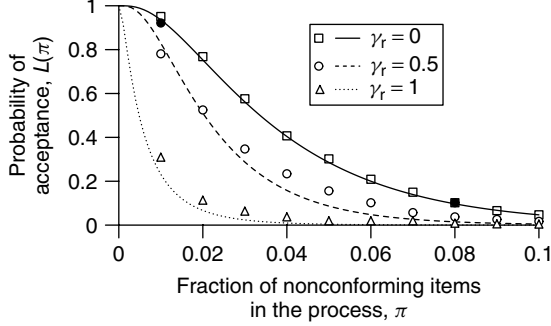


Figure 1 The OC curves of the single sampling plan ($n_s^* = 25$, $c_s^* = 1.83$), $m = 1$ measurement per item, which fulfills the requirements $L(1\%) \geq 1 - \alpha = 93\%$ and $L(8\%) \leq \beta = 10\%$ if $\gamma_r = 0$, but fails if it is used in the case of measurement error, $\gamma_r = 0.5$ and $\gamma_r = 1$. Each symbol represents the acceptance frequency on the basis of 1000 simulation runs

m Measurements per Item

In any case, an existing measurement error, if its standard deviation σ_r is unknown, changes the OC function and, hence, the risks of the single sampling plan, in an uncontrolled manner. However, this difficulty can be partly overcome if each sampled item is measured m times ($m = 2, 3, \dots$) and a one-way analysis of variance is performed, i.e., the total **sum of squares** is broken down into one between sampled items and one within them:

$$\underbrace{\sum_{i=1}^{n_s^*} \sum_{j=1}^m (Y_{ij} - \bar{\bar{Y}})^2}_{Q_{\text{total}}} = m \underbrace{\sum_{i=1}^{n_s^*} (\bar{Y}_i - \bar{\bar{Y}})^2}_{Q_{\text{between}}} + \underbrace{\sum_{i=1}^{n_s^*} \sum_{j=1}^m (Y_{ij} - \bar{Y}_i)^2}_{Q_{\text{within}}} \quad (19)$$

where

$$\bar{Y}_i = \frac{1}{m} \sum_{j=1}^m Y_{ij}; i = 1, \dots, n_s^* \quad (20)$$

and

$$\bar{\bar{Y}} = \frac{1}{n_s^* m} \sum_{i=1}^{n_s^*} \sum_{j=1}^m Y_{ij} \quad (21)$$

is the mean of the measurements $Y_{ij}; i = 1, \dots, n_s^*; j = 1, \dots, m$.

As test statistic for the test we use

$$Z_s := \bar{\bar{Y}} + c_s^* S_w \quad (22)$$

where

$$S_w^2 = \max(0, (S_{II}^2 - S_I^2)/m) \quad (23)$$

and

$$S_I^2 = \frac{Q_{\text{within}}}{n_s^* (m - 1)}, \quad S_{II}^2 = \frac{Q_{\text{between}}}{n_s^* - 1}. \quad (24)$$

The OC function of the single sampling plan (n_s^* , c_s^*) under parameters (γ_r , m) obtained by using a normal approximation of Z_s is

$$L(\pi) = P(Z_s \leq UL) = \Phi\left(\frac{\sqrt{n_s^*}}{A}(u_{1-\pi} - c_s^*)\right) \quad (25)$$

with

$$A = A(c_s^*, \gamma_r, m) = \sqrt{1 + \frac{\gamma_r^2}{m} + \frac{(c_s^*)^2}{2m^2} \left[(m + \gamma_r^2)^2 + \frac{\gamma_r^4}{m - 1} \right]} \quad (26)$$

If (n_σ, c_σ) are the parameters of the sampling plan for the case of no measurement error and known standard deviation which has the desired properties $L(\pi_1) \geq 1 - \alpha$ and $L(\pi_2) \leq \beta$ the matching plan for the case of measurement error, expressed by γ_r , has the parameters $c_s^* = c_\sigma$ and

$$n_s^* = A^2 n_\sigma \quad (27)$$

Unfortunately, as for $m = 1$ the OC function and, hence, n_s^* depend on γ_r . However, since $c_s^* = c_\sigma$ is determined only by $P_1(\pi_1; 1 - \alpha)$ and $P_2(\pi_2; \beta)$ for all values of γ_r , the OC functions of all these sampling plans have the common point $(\pi, L(\pi))$ with π given by $u_{1-\pi} = c_\sigma$ and $L(\pi) = \Phi(0) = 0.5$. Increasing γ_r turns the OC around this 50% point (point of indifferent quality) so that it becomes more flat and does not anymore fulfill both the requirements $L(\pi_1) \geq 1 - \alpha$ and $L(\pi_2) \leq \beta$. On the other hand, if the sampling plan has been designed for a given γ_r and is applied in a situation with a smaller γ_r , the OC becomes steeper.

In most of the practical applications an upper bound for the ratio γ_r is available, e.g., the measurement standard deviation is not larger than the process standard deviation ($\gamma_r \leq 1$), and the single sampling plan can be designed for measurement error with

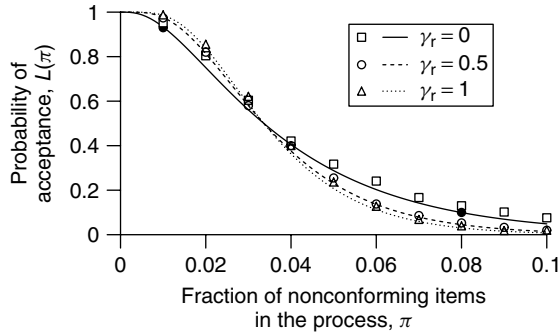


Figure 2 The OC curves of the single sampling plan ($n_s^* = 52$, $c_s^* = 1.83$), $m = 2$ measurements per item, which fulfills the requirements $L(1\%) \geq 1 - \alpha = 93\%$ and $L(8\%) \leq \beta = 10\%$ if $\gamma_r = 1$ and performs better for $\gamma_r = 0.5$ and $\gamma_r = 0$. Each symbol represents the acceptance frequency on the basis of 1000 simulation runs

respect to this upper bound; see Figure 2. This would have the advantage that for all situations for which the measurement error is smaller than this upper bound the actual risks α of the producer and β of the consumer would both be smaller than the values used for the design of the plan. Smaller measurement error would pay for both sides!

Example

A single sampling plan for inspection by variables shall be determined so that for $\pi = 1\%$ the probability of acceptance is $L(\pi_1) \geq 1 - \alpha = 93\%$, and for $\pi_2 = 8\%$ it is $L(\pi_2) \leq \beta = 10\%$.

For the case of σ known, no measurement error ($\gamma_r = 0$) and $m = 1$ we find $c_\sigma = c_s^* = 1.83$ and $n_\sigma = n_s^* = 9$ from equations (13) and (12). For unknown σ we get $c_s = c_s^* = c_\sigma = 1.83$ and $n_s = n_s^* = 25$ from equation (18). Figure 1 shows the design criteria for the plan and the OC curve according to equations (8) and (16) which coincide. If σ is unknown and this sampling plan (where each sampled item is measured once) is used in the presence of measurement error, the OC function depends on $\gamma_r = \sigma_r/\sigma$. Figure 2 shows the OC curves of the plan ($n_s^* = 25$, $c_s^* = 1.83$) for $\gamma_r = 0; 0.5$ and 1.0 . To

check the accuracy of the approximations, results of the simulations are added.

If a measurement error exists and the variances σ^2 and σ_r^2 are known, the sample sizes for $\gamma_r = 1$ are, according to equation (12), $n_\sigma^* = 18, 14, 12, 12, 11$ for $m = 1, 2, 3, 4, 5$ respectively. From $m = 9$ to $m = 212$ the sample size is $n_\sigma^* = 10$ and for $m \geq 213$ it is $n_\sigma^* = 9$, i.e., equal to the sample size for the case of no measurement error.

If a measurement error exists and the variances are unknown the sample sizes for $\gamma_r = 1$ and $m = 2, 3, 4$ are, according to equations (26) and (27), $n_s^* = 52, 40, 35$, respectively. Figure 2 shows the OC curves of the plan ($n_s^* = 52$, $c_s^* = 1.83$) for $m = 2$ if the actual measurement error is $\gamma_r = 1$ (the design value), 0.5 and 0 (i.e., smaller than the design value). Both the producer risk α and the consumer risk β are smaller than the required values. Again, to check the accuracy of the approximations, results of the simulations are added.

Endnotes

^aIn the case of a specified lower limit (LL) the fraction of nonconforming items is $\pi = P(X < LL) = \Phi\left(\frac{LL - \mu}{\sigma}\right)$. In equations (7), (14) and (22) the negative sign has to be used and the lot is accepted if the test statistic is greater than or equal to LL .

References

- [1] Owen, D.B., & Chou, Y.-M. (1983). Effect of measurement error and instrument bias on operating characteristics for variables sampling plans, *Journal of Quality Technology* **15**, 107–117.
- [2] Melgaard, H., & Thyregod, P. (2001). Acceptance sampling by variables under measurement uncertainty, in *Frontiers in Statistical Quality Control*, H.-J. Lenz & P.-Th. Wilrich, eds, Physica-Verlag, Heidelberg, New York, Vol. 6, pp. 47–57.
- [3] Wilrich, P.-Th. (2000). Single sampling plans for inspection by variables in the presence of measurement error, *Allgemeines Statistisches Archiv* **84**, 239–250.

PETER-TH. WILRICH

Attributes Sampling under Classification Error

Introduction

The effect of inspection error in **acceptance sampling** has been studied extensively. For example, see Bennett *et al.* [1], Dorris and Foot [2], Ghosh [3], Hong *et al.* [4], Johnson *et al.* [5–7], Maghsoodloo [8], and Moskowitz and Fink [9]. The effect of inspection error on group testing has also been considered at length by Graff and Roeloffs [10], Huang *et al.* [11], and Johnson *et al.* [12–14]. The monograph by Johnson *et al.* [15] provides a detailed consolidation of the work in this field in a way that this brief article cannot do. Instead, this article serves as an introduction.

Single Sampling

We consider a lot of size N which has D nonconforming items. We will select a single random sample of n units for inspection. Each unit in the sample will be classified as either conforming (C) or nonconforming (NC). We let p_1 be defined as the probability that an NC item is correctly identified as NC and p_2 be the probability that a C item is incorrectly identified as NC. We let the random variable, Y be the number of NC items in the sample and Z be the number of items identified as NC from the sample. The probability that there are $Y = y$ NC items in the sample has a hypergeometric distribution. That is

$$\Pr[Y = y] = \frac{\binom{n}{y} \binom{N-n}{D-y}}{\binom{N}{D}} = \frac{\binom{D}{y} \binom{N-D}{n-y}}{\binom{N}{n}}, \quad \max(0, n - N + D) \leq y \leq \min(n, D) \quad (1)$$

When $p_1 < 1$, only some of the NC items in the sample will be identified as NC. The conditional probability for $Z = z | Y = y$ is binomial,

$$P[Z = z | y] = \binom{y}{z} p_1^z (1 - p_1)^{y-z} \quad (2)$$

Therefore, the probability that $Z = z$ items are identified as NC is

$$\Pr[Z = z] = \frac{\sum_{y \leq z} \binom{D}{y} \binom{N-D}{n-y} \binom{y}{z} p_1^z (1 - p_1)^{y-z}}{\binom{N}{n}} \quad (3)$$

$$E[Z] = np_1 \frac{D}{N} \quad (4)$$

and

$$\text{Var}[Z] = n \frac{(N-n)D}{N-1} \left(1 - \frac{D}{N}\right) + p_1(1 - p_1) \frac{nD}{N} \quad (5)$$

When $p_2 > 0$, some of the C items in the sample will also be identified as NC. The conditional probability that $z - w$ of these $n - y$ C items will be identified as NC is

$$\binom{n-y}{z-w} p_2^{z-w} (1 - p_2)^{n-y-z+w} \quad (6)$$

Therefore, when $p_2 > 0$ and $p_1 < 1$,

$$\begin{aligned} \Pr[Z = z] &= \frac{\sum_{w \leq z} \binom{D}{y} \binom{N-D}{n-y} \binom{y}{w} p_1^w (1 - p_1)^{y-w}}{\binom{N}{n}} \\ &\times \binom{n-y}{z-w} p_2^{z-w} (1 - p_2)^{n-y-z+w} \end{aligned} \quad (7)$$

$$\begin{aligned} E[Z] &= np_1 \left(\frac{D}{N}\right) + np_2 \frac{(N-D)}{N} \\ &= n[p \cdot p_1 + (1 - p)p_2] = n\bar{p} \end{aligned} \quad (8)$$

and

$$\text{Var}[Z] = n\bar{p}(1 - \bar{p}) - n(n-1) \frac{(\bar{p}^2 - \bar{p}^2)}{(N-1)} \quad (9)$$

where $p = D/N$, $\bar{p} = p_1 \cdot p + p_2(1 - p)$ and $\bar{p}^2 = \{Dp_1^2 + (N-D)p_2^2\}/N$. $\bar{p}$ is the probability that a randomly selected item is classified as NC.

When $p_1 = 1$ and $p_2 = 0$, there is no inspection error and $\Pr[Z = z]$ is equal to $\Pr[Y = z]$. When $p < 1$ and/or $p_2 > 0$ the operating characteristic function, $\Pr[Z \leq c] = \sum_{z=0}^c \Pr[Z = z]$ can be substantially different from the values obtained using equation (1), i.e. the values that would be obtained

2 Attributes Sampling under Classification Error

Table 1 $P[Z = 0]$ for acceptance sampling with $N = 100$, $D = 5$, and $n = 5$

	$p_2 = 0$	$p_2 = 0.025$	$p_2 = 0.05$	$p_2 = 0.075$	$p_2 = 0.1$
$p_1 = 0.75$	0.8236	0.7269	0.6395	0.5608	0.4900
$p_1 = 0.8$	0.8126	0.7170	0.6306	0.5527	0.4827
$p_1 = 0.85$	0.8017	0.7071	0.6217	0.5447	0.4755
$p_1 = 0.9$	0.7909	0.6974	0.6129	0.5368	0.4684
$p_1 = 0.95$	0.7802	0.6877	0.6041	0.5289	0.4614
$p_1 = 1$	0.7696	0.6781	0.5955	0.5212	0.4544

without inspection error. For example, the probability of accepting a lot under zero-defect inspection is shown in Table 1 for various values of p_1 and p_2 . The probability decreases if $p_2 > 0$, i.e. if C items may be incorrectly classified as NC. The probability of accepting the lot decreases if $p_1 < 1$, i.e. if some NC items may be classified as C. In Table 1, the effect of $p_2 > 0$ is seen to be more substantial. Similarly, the values for $E[Z]$ and $\text{Var}[Z]$ can differ substantially from those given by $E[Y]$ and $\text{Var}[Y]$. In acceptance sampling, the effect of inspection error should be considered when selecting the sampling plan (n, c) , as described by Ferrell and Chhoker [16].

Multiple Sampling

Instead of taking a single sample of size n , we may take k samples without replacement. The samples may be of different sizes $n = (n_1, n_2, \dots, n_k)$, resulting in random number of NC items $Y = (Y_1, Y_2, \dots, Y_k)$ in each sample as well as random number of items $Z = (Z_1, Z_2, \dots, Z_k)$ that are identified (correctly or incorrectly) as NC in each sample. When there is no inspection error, $Z = Y$. If $k = 2$, this procedure is called *double sampling*. In double sampling, we take a random sample of size n_1 . If Z_1 , the number of NC items (either correctly or incorrectly identified) is less than or equal to a threshold a_1 , accept the lot. If $Z_1 \geq a'_1$, reject the lot. If $a_1 < Z_1 < a'_1$, we take another sample of size n_2 . If $Z_1 + Z_2 \leq a_2$ we accept the lot. Otherwise the lot is rejected. In multiple sampling, additional samples can be selected before deciding whether to accept or reject a lot. For more information on multiple sampling see **Multiple Sampling Plans**.

The probability $p_1 = (p_{1,1}, p_{1,2}, \dots, p_{1,k})$, of correctly identifying NC items as NC and $p_2 = (p_{2,1}, p_{2,2}, \dots, p_{2,k})$ of incorrectly identifying C items as NC may either be the same in each sample (the equal error probabilities case) or the error probabilities may

be different in each sample. Either way, the joint distribution of Y is a multivariate hypergeometric

$$\begin{aligned} \Pr[Y = y] &= \frac{\binom{D}{y} \binom{N-D}{n-y}}{\binom{N}{n}} \\ &= \frac{\prod_{i=1}^k \binom{n_i}{y_i} \binom{n - \sum y_i}{D - \sum y_i}}{\binom{N}{D}} \end{aligned} \quad (10)$$

And, following the development for **single sampling**,

$$\begin{aligned} \Pr[Z = z] &= \frac{\sum_{w \leq z} \binom{D}{y} \binom{N-D}{n-y}}{\binom{N}{n}} \\ &\times \prod_{j=1}^k \binom{y_j}{w_j} p_{1,j}^{w_j} (1 - p_{1,j})^{y_j - w_j} \\ &\times \binom{n_j - y_j}{z_j - w_j} p_{2,j}^{z_j - w_j} \\ &\times (1 - p_{2,j})^{n_j - y_j - z_j + w_j} \end{aligned} \quad (11)$$

Acceptance probabilities can be calculated using this expression. These probabilities are tabulated in [7].

Group Testing (Screening)

Dorfman [17] was concerned with testing large numbers of blood samples for syphilis. To save time and money, he mixed together blood from a group of n_0 samples and tested the entire mixture instead of the individual samples. If the test was negative, testing was completed with only one test. If the test result

was positive, each sample would be tested resulting in a total of $n_0 + 1$ tests. When the inspection procedure is perfect, the savings are maximized by letting n_0 equal $(\text{proportion NC})^{-1}$.

When there are inspection errors, it is possible that the group is incorrectly classified as NC. This can lead to unneeded individual testing. If individual samples are tested, they also may be incorrectly classified. In addition, it is possible that a group contains NC items but it is classified as C. As before, for the individual testing, p_1 is defined as the probability that an NC item is correctly identified as NC and p_2 is the probability that a C item is incorrectly identified as NC. In addition, p_0 is defined as the probability that the group is correctly declared NC when at least one sample in the group is NC and p'_0 is the probability that the group is incorrectly identified as NC when all items in the group are C. The probability that an NC item will be correctly classified is $p_0 p_1 < p_1$. The probability that a C item is correctly classified as $P_0^*(n_0)(1 - p'_0 p_2) + \{1 - P_0^*(n_0)\}(1 - p_0 p_2)$ where $P_0^*(n_0) = (N - D - 1)^{n_0-1} / (N - 1)^{n_0-1}$ is the probability that the group has no NC items given that a particular item is C.

Since the probability of a correct classification of an NC item is decreased by the screening process, the value of screening must come from increased probability of correctly identifying C items and/or a reduction in the number of tests needed. The expected number of tests is $1 + n_0\{p_0 - (p_0 - p'_0)P_0(n_0)\}$ where $P_0(n_0) = (N - D)^{n_0} / N^{n_0}$ is the probability that a group has no NC items.

Acceptance Sampling with Rectification

In acceptance sampling, lots that are not accepted can be either discarded or rectified. Rectification, i.e. replacing or discarding all NC units after 100% inspection of rejected lots is frequently used when manufacturing costs are high. See **Rectification Sampling Schemes** for general information on rectification sampling. Greenberg and Stokes [18] used the information obtained in rectification to estimate the number or rate of NC units, i.e. the average outgoing quality (AOQ) in a series of outgoing lots after zero-defect sampling. The estimator for the number of NC units is

$$\hat{U}_{GS} = \sum_{Y_{i1} > 0} Y_i \frac{1 - P_i}{P_i} \quad (12)$$

Where Y_i is the number of items identified as NC in lot i and P_i is the probability that at least one NC unit will be chosen from lot i in the initial random sample.

$$P_i = \begin{cases} 1 - \frac{\binom{N - Y_i}{n}}{\binom{N}{n}} & \text{if } Y_i \leq N - n \\ 1 & \text{otherwise} \end{cases} \quad (13)$$

Anderson, Greenberg and Stokes [19] show that inspection errors can have a serious biasing effect on this estimator. An alternative estimator allows for c defect sampling and reduces both the **bias** and **mean squared error**.

$$\begin{aligned} \hat{U}_{ca} &= \sum_{I_{ca,i}=1} \frac{Y_i}{p_1 - p_2} \left(\frac{1}{P_{ca,i}} \right) - \sum_{i=1}^T \frac{N p_2}{p_1 - p_2} \\ &\quad - \sum_{i=1}^T \frac{Y_{i1} - n p_2}{p_1 - p_2} p \\ &\quad - \sum_{I_{ca,i}=1} \frac{Y_{i2} - (N - n) p_2}{p_1 - p_2} p_1 \\ &= \frac{1}{p_1 - p_2} \left\{ \sum_{i=1}^T \left[\frac{Y_i I_{ca,i}}{P_{ca,i}} - p_1 (Y_{i1} + Y_{i2} I_{ca,i}) \right] \right. \\ &\quad \left. - T p_2 (N - n p_1) + N_{ca} p_2 p_1 (N - n) \right\} \quad (14) \end{aligned}$$

where $N_{ca} = \sum_{i=1}^T I_{ca,i}$ is the total number of lots rectified from a set of T lots. $I_{ca,i}$ is 1 if lot i is rectified under a modified (n, c) sampling plan or zero otherwise. A small proportion, α , of lots with c or fewer errors in the sample are rectified in addition to all lots with more than c errors. The modification is required to ensure that $P_{c,i}$ is always greater than zero. AOQ is thus estimated to be $\hat{U}_{ca} / NT$. The average total inspections (ATI) will increase slightly because of this modification. ATI increases when $p_2 > 0$, i.e. when C items may be incorrectly classified as NC and decreases when $p_1 < 1$, i.e. when some NC items may be classified as C. These effects may balance out, but more often the effect of $p_2 > 0$ is larger since there are (hopefully) more C than NC items.

Estimating Error Probabilities

Generally it is assumed that the misclassification probabilities p_1 and p_2 are known. In practice this is rarely the case, so the interesting quantities are often tabulated to see how sensitive they are to classification error. If $p = D/N$ is the proportion of devices that are C, and as usual p_1 is the probability that an NC item is correctly identified, and p_2 is the probability that a C item is incorrectly identified as NC, then Z , the number of items identified as NC allows us to estimate the parameter $p \cdot p_1 + (1 - p)p_2$ and does not help us to estimate p , p_1 , or p_2 . If items are tested only once, we will never know whether errors occur in the testing process. It is only when we see items that are identified by one test as C and as NC by another test that we realize that $p_1 < 1$ and/or $p_2 > 0$, i.e., there are classification errors. To estimate these parameters some items must be tested more than once.

The problem of estimating p_1 and p_2 has also been studied extensively (see [20] and [12]) and is acknowledged to be among the most challenging in **attribute sampling**. When neither p_1 nor p_2 can be considered negligible, it is necessary to inspect some items at least three times to estimate p , p_1 , and p_2 . Here, we consider the case in which a number of items, say N , are inspected exactly three times. The simplest estimators of p_1 , p_2 , and p assume that items that fail at least twice are NC and the items that fail fewer times are C. Then p_1 , p_2 , and p are estimated using the number of correct outcomes from the tests. That is

$$\hat{p}_{1,EZ} = \frac{2N_2 + 3N_3}{3(N_2 + N_1)} \quad (15)$$

$$\hat{p}_{2,EZ} = \frac{N_1}{3(N_0 + N_1)} \quad (16)$$

and

$$\hat{p}_{EZ} = \frac{N_2 + N_3}{N} \quad (17)$$

where N_j is the number of items declared nonconforming j times, $j = 0, 1, 2, 3$. If $\hat{p}_{1,EZ}$ is undefined because $N_2 + N_3 = 0$, then we set $\hat{p}_{1,EZ} = 1$. Similarly, if $\hat{p}_{2,EZ}$ is undefined, we set its value equal to zero. Alternatively, maximum likelihood estimators (MLEs) of p_1 and p_2 can be obtained. The MLEs for $\hat{p}_{1,ML}$ and $\hat{p}_{2,ML}$ maximize $L = \prod_{j=0}^3 [P(j)]^{N_j}$,

where

$$P(j) = p \binom{3}{j} p_1^j (1 - p_1)^{3-j} + (1 - p) \binom{3}{j} p_2^j (1 - p_2)^{3-j} \quad (18)$$

Note that it is necessary to simultaneously maximize L with respect to p_1 , p_2 , and p .

Surprisingly, in [19] Anderson, Greenberg, and Stokes found that $\hat{p}_{1,EZ}$ and $\hat{p}_{2,EZ}$ often performed better than $\hat{p}_{1,ML}$ and $\hat{p}_{2,ML}$ and always performed at least as well for values that are typical of this application.

References

- [1] Bennett, G.K., Case, K.E. & Schmidt, J.W. (1974). The economic effects of inspector error on attribute sampling plans, *Naval Research Logistics Quarterly* **21**, 431–433.
- [2] Dorris, A.L. & Foote, B.L. (1978). Effect of human inspector errors on statistical quality control plans, *Proceedings of the AIIE Spring Conference*, American Institute of Electrical Engineer, New York, pp. 487–492.
- [3] Ghosh, D.T. (1985). An investigation of inspection error on AOQ of a single sampling acceptance verification system, *Sankhya, Series B* **47**, 233–241.
- [4] Hong, L.L., Foote, B.L. & Mount-Campbell, C. (1975). The effect of inspector accuracy on the type I and type II error of common sampling techniques, *Journal of Quality Technology* **7**, 157–164.
- [5] Johnson, N.L., Kotz, S. & Rodriguez, R.N. (1984). Inspection errors in link sampling, *Communications in Statistics: Theory and Methods* **13**, 1203–1213.
- [6] Johnson, N.L., Kotz, S. & Rodriguez, R.N. (1985). Statistical effects of imperfect inspection sampling: I. Some basic distributions, *Journal of Quality Technology* **17**, 1–31.
- [7] Johnson, N.L., Kotz, S. & Rodriguez, R.N. (1986). Statistical effect of imperfect inspection sampling: II. Double sampling and link sampling, *Journal of Quality Technology* **20**, 98–124.
- [8] Maghsoodloo, S. (1987). Inspection error effects on performance measures of a multistage sampling plan, *AIIE Transactions* **19**, 340–347.
- [9] Moskowitz, H. & Fink, R.K. (1977). A Bayesian algorithm incorporating inspection errors for quality control and auditing, *TIMS Studies in the Management Science* **7**, 79–104.
- [10] Graff, L.E. & Roeloffs, R. (1974). A group testing procedure in the presence of test error, *Journal of the American Statistical Association* **69**, 159–163.
- [11] Huang, Q., Johnson, N.L. & Kotz, S. (1989). Modified Dorfman-Sterrett screening (group testing) procedures

- and the effects of faulty inspection, *Communications in Statistics: Theory and Methods* **18**, 1485–1495.
- [12] Johnson, N.L., Kotz, S. & Rodriguez, R.N. (1988). Statistical effect of imperfect inspection sampling: III. Screening (group testing), *Journal of Quality Technology* **22**, 128–137.
- [13] Johnson, N.L., Kotz, S. & Rodriguez, R.N. (1989). Dorfman-Sterrett (group testing) schemes and the effects of faulty inspections, *Communications in Statistics: Theory and Methods* **18**, 1469–1484.
- [14] Johnson, N.L., Kotz, S. & Rodriguez, R.N. (1990). Statistical effects of imperfect inspection sampling: IV. Modified Dorfman screening procedures, *Journal of Quality Technology* **22**, 128–137.
- [15] Johnson, N.L., Kotz, S. & Wu, X. (1991). *Inspection Errors for Attributes in Quality Control*, Chapman & Hall, London.
- [16] Ferrell Jr, W.G. & Chhoker, A. (2002). Design of economically optimal acceptance sampling plans with inspection error, *Computers and Operations Research* **29**, 1283–1300.
- [17] Dorfman, R. (1943). The detection of defective members of large populations, *Annals of Mathematical Statistics*. **14**, 436–446.
- [18] Greenberg, B.S. & Stokes, S.L. (1992). Estimating nonconformance rates after zero defect sampling with rectification, *Technometrics* **34**, 203–213.
- [19] Anderson, M.T., Greenberg, B.S. & Stokes, S.L. (2001). Acceptance sampling with rectification when inspection errors are present, *Journal of Quality Technology* **33**, 493–505.
- [20] Johnson, N.L. & Kotz, S. (1988). Estimation for binomial data with classifiers of known and unknown imperfections, *Naval Research Logistics* **35**, 147–156.

BETSY S. GREENBERG

Database Systems for Acceptance Sampling

Introduction: Database Systems

A database system is a triple of a database **DB**, a set of application programs **AP** and a database management system **DBMS**. A database **DB** over a database schema S consists of a finite set R of relations (tables) $r \subseteq \text{dom}(A)$ and a set of integrity constraints, Σ , where $\text{dom}(A)$ is the value set defined by the cross-product of the ranges of all attributes $A \in A$. The attributes have specific data types like *int*, *real*, *string*, *Boolean*, and so on. The application programs support planning and management of batch inspection. The integrity constraints assure technical, statistical, and semantic coherency of the data stored in any $r \in R$ at any time [1].

Database Systems for Acceptance Sampling

With respect to the granularity of data in **DB** and the objective of data retrieval we differentiate between on-line transaction processing (OLTP) and on-line analytical processing (OLAP) databases. OLTP databases used for “on-line transaction processing” have been in use since 1970 and are specially tuned to manage operative mass data from the shop floor [2]. They are used for on-line control like **acceptance sampling**, monitoring or **process control**, and in off-line control, for instance, for storing data of designed experiments. A typical query of **quality control** (QC) administration is “What is the inspection state (‘in control’, ‘accepted’, or ‘rejected’) of a batch submitted by company C ‘last week’ with part-id = 4711?”

Contrarily, OLAP databases (“analytical transaction processing”) are designed to support the middle QC manager or the chief executives of a company with respect to planning, decisions and controlling. They were introduced in business from 1990 on [3]. They store business indicators as aggregated and grouped data by using statistical functions like *avg* (average), *sum*, *min*, *max*, and *counts* (frequency), combined with the grouping operator *group-by*. A typical indicator might be “Percentage of rejected

items in incoming batches grouped by supplier-id, part-id, and quarter”.

A major advantage of a database system is to hide details of the storage, the control of concurrent access and the maintenance of the data from the users. On the *physical level*, the database designer describes in a physical schema $S_p \subseteq S$ how the database is technically organized, stored and accessed. For example, in a distributed database system, data with lowest granularity is recorded on the shop floor. Daily grouping and summarizing of such measurements generates low-level aggregated data which is entered into a departmental data warehouse. Aggregation over departments produces an enterprise-wide view. Further details are skipped here, cf. [4]. The external schema $S_u \subseteq S$ provides a view for each authorized user u according to his individual (“personalized”) profile. Views are important for enforcing data security and usage comfort. Access – as read or write – to the database is enabled by the (declarative) structured query language called *SQL* which is a *de facto* standard [4]. A conceptual schema $S_c \subseteq S$ is an abstract model of the real world focusing on the frame of discernment. It consists of (R, Σ) defining relations (tables) and related constraints. The DBMS furthermore supplies a data definition language to generate the tables and to provide a mapping from the conceptual level to the physical implementation.

We present now, an example of a conceptual schema of a QC database system for an adaptive acceptance sampling system which combines a skip-lot system and the Mil-Std. 105D [5]. This OLTP database system combines two principles of adaptability used in industry: lot-by-lot inspection with varying sample sizes or skipping lots according to the history of quality. We note that attributes which play the role of an identifier of objects (primary key) are underlined. The most elementary tables are those for suppliers, parts, and quality characteristics (attributes) which are an intrinsic part of any enterprise information system like SAP, Oracle Business Suite, and so on.

```
Supplier (supplier_id, name, info),  
Part (part_id, name, description),  
Attribute (attribute_id, name, description).
```

Furthermore a “supply list” is needed as a subset $L \subseteq \text{Supplier} \times \text{Part}$:

```
SupplyPart (supplier_id, part_id,  
conditions).
```

2 Database Systems for Acceptance Sampling

The QC department is responsible for the design of the basis sampling plans of **attribute sampling** for reduced, normal and tightened inspection indexed by batch size, class and average quality level (AQL):

Basis Sampling Plans (lower_batch_size, upper_batch_size, AQL_level, (sample_sizes, acceptance_numbers)).

The same is true for switching rules which are an inherent part of adaptive systems. In the MISL (mixed MIL-Std 105D and Skip-Lot sampling system) system the states 0 to 2 correspond to the states in Mil-Std 105D, i.e., tightened, normal, and reduced inspection levels, while the inspection states 3 to 5 are skip-lot states with varying skip sizes or batch inspection frequencies:

Switching_Rules (AQL_level, relaxation_number, skip_size3, skip_size4, skip_size5).

As the AQL is assumed to be invariant with respect to suppliers, we have to know the AQL value for each quality characteristic and each part :

Part_Attribute (part_id, attribute_id, AQL_level).

Note that the test department, the test device, as well as the measurement method used in the inspection phase can easily be added as further attributes to that relation.

Any system, especially a dynamic one, needs information about the state of inspection of incoming batches. Assuming that incoming batches are strictly sequenced, we get:

Inspection-state (supplier_id, part_id, attribute_id, batch_id, date, inspection_level, inspection_no, skip_no).

As the decision to inspect a batch or not and the **sample size** of batch inspection depends on the so-called history of quality of an attribute, i.e., on the sequence of former inspection states, we need to store the history of inspection for each attribute of every batch supplied by every supplier:

History(date, batch_id, attribute_id, supplier_id, part_id, update_indicator, sample_size, acceptance_no, no_of_defects, batch_size).

We summarize by representing the corresponding entity-relationship (ER) model in Figure 1.

We illustrate the ease of querying of OLTP databases by checking whether or not batch b_1 of

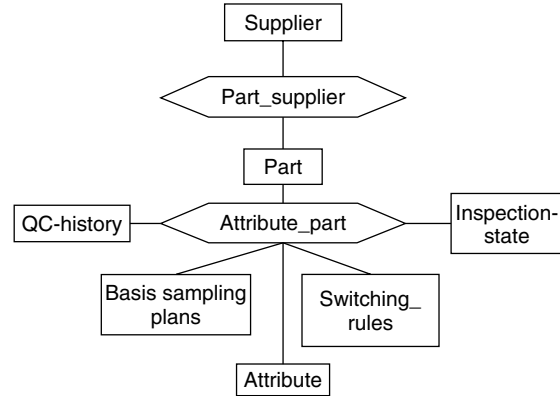


Figure 1 ER model for (adaptive) acceptance sampling

electronic switches supplied on October 27, 2007, passed inspection with respect to the tolerance of electric resistance (TER):

```

select no_of_defects ≤ acceptance_no
from History
where date = 2007-10-27 and batch_id
= 'b 1' and attribute_id = 'TER'.
    
```

Data Warehousing for Acceptance Sampling

While querying of an operative database is user-friendly and is efficiently performed owing to tree-indexed structures and hashing [6], retrieval of aggregated data with range queries becomes more troublesome. Such deficiencies caused the development of OLAP databases called *data warehouses*. The main data structure is the data cube which is called *multiway table* in statistics [3]. It can be defined as a tuple (statistical function f , [summary attribute A], set of dimensions D). The symbol $[]$ means that this term is optional. Facts are composed of a statistical function and – optionally – a summary attribute like “number of defects” or “diameter”. Dimensions are attributes which are not pairwise (strong or weak) functionally dependent [7] cf. Figure 2. For example, the data cube *Supplier-Report* ($100 * \text{acceptance rate}, \emptyset, \{\text{part_id}, \text{supplier_id}, \text{month/quarter/year}\}$) represents the percentage of accepted incoming batches grouped by part, supplier, and time. Note that attributes like part or time may have domains which are partially ordered,

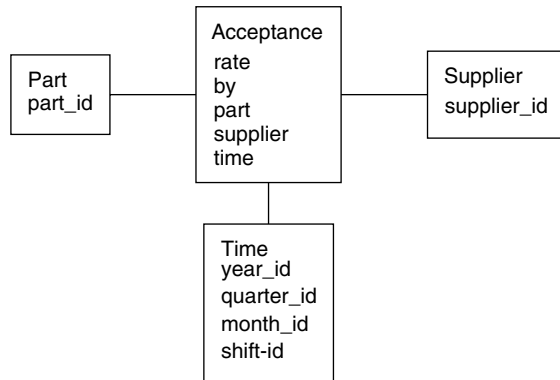


Figure 2 Data cube with one fact table and three-dimensional tables

i.e., form hierarchies. The cube operators needed are cube projection (dicing), cube selection (slicing) and hierarchy traversal (roll-up, drill-down). A cube, which can again be viewed as a table, can be materialized (permanently stored and updated) or virtual, i.e., computed “on the fly”. The database administrator has to find the best trade-off between retrieval time and storage size.

Evidently, the mapping from OLTP to OLAP databases is not simple [1]. OLTP–OLAP data transformations and retrieval of data combine elements from the set of cube operators O , the set of aggregation functions F , and the set of transformations T . Admissible operations should guarantee their statistical and semantic coherency. For instance, consider the transformation T of a consumption rate measured in l/100km to miles/gallon, which is a nonlinear transformation. Testing cars with respect to consumption or mileage is usually done within and outside a city, and the measured rates of consumption or mileage are averaged. Of course, as a linear operator (average) and a nonlinear transformation T are not generally commutative, contradicting decision may be caused when either the consumption rate or the mileage rate is used [8].

Outlook

SAP Netweaver, IBM’s database management system DB2, Microsoft’s SQL server and Oracle’s Business Suite offer a broad spectrum of automated **data collection** facilities and retrieval functionality for mass data. Although the border between querying and statistical data analysis is becoming vaguely more and more, there still remains a demand in industry for explorative data analysis, which can be handled by added-on **data mining** or “intelligent” data analysis tools [9, 10], for the underlying methodology.

References

- [1] Lenz, H.-J. & Thalheim, B. (2005). OLAP schemata for correct applications, *Trends in Enterprise Architectures and Applications (TEAA)*, Springer, Heidelberg.
- [2] Codd, E.F. (1970). A relational model for large shared databases, *Communications of The ACM* **13**(6), 377–387.
- [3] Inmon, W.H. (1992). *Building the Data Warehouse*, Springer, New York.
- [4] Elmasri, R. & Navathe, S.B. (1999). *Fundamentals of Database Systems*, 3rd Edition, Addison-Wesley.
- [5] Lenz, H.-J. (1987). Design and implementation of a sampling inspection system for incoming batches based on relational databases, in *Frontiers in Statistical Quality Control*, H.-J. Lenz, G.B. Wetherill & P.-Th. Wilrich, eds, Physica, Heidelberg, Vol. 3, pp. 116–127.
- [6] Gaede, V. & Günther, O. (1998). Multidimensional access methods, *ACM Computing Surveys* **30**(2), 170–231.
- [7] Lehner, W., Albrecht, J. & Wedekind, H. (1998). Normal forms for multidimensional databases, *Proceedings of the 10th International Conference On Scientific and Statistical Data Management (SSDBM’98)*, IEEE, pp. 63–72.
- [8] Hand, D. (1994). Deconstructing of statistical questions (with discussion), *Journal of the Royal Statistical Society. series A*, **157**, 317–356.
- [9] Hand, D., Mannila, H. & Smyth, P. (2001). *Principles of Data Mining*, MIT Press, Cambridge.
- [10] Berthold, M. & Hand, D. (eds) (1999). *Intelligent Data Analysis*, Springer, Berlin, Heidelberg.

HANS-J. LENZ

Acceptance Sampling in Modern Industrial Environments

Introduction

In the mass production environments of the first half of the twentieth century, **quality control** rested on the triple *sorting* principle of *inspecting* product, *detecting* inappropriate product, and *sorting out* inappropriate product. Accordingly, statistical **product inspection** (SPI) based on sampling was the first strong branch of modern statistical quality control (SQC). Two prominent SPI schemes are the Dodge–Romig tables developed at Bell Telephone Laboratories in the 1920s and 1930s, and the Military Standard (MIL-STD) series developed in the 1940s and 1950s for the U.S. armament industry. The sorting principle is clearly expressed by the definition of the purpose of sampling inspection given by H. F. Dodge and H. G. Romig (see [1, p. 11]): “The purpose of the inspection is to determine the acceptability of individual lots submitted for inspection, i.e., sorting good lots from bad.”

Industry changed drastically in the second half of the twentieth century. Quality control adopted a new principle: *provision* for appropriate product, instead of sorting out inappropriate product. Consequently, emphasis shifted from product inspection to issues like design control, **process control** maintenance, and service. The same change of paradigms occurred in SQC. However, SPI can play a reasonable role in modern industrial quality control, if its status, purpose, and methods are defined adequately. Some essential challenges for modern, useful, and promising SPI are discussed in the subsequent sections.

The Adjunct Role of SPI in Modern Product Control

SPI is a tool of product quality control, where the term *product* is used in a broad sense to include any completed identifiable entity like materials, parts, products, services, or data. Product control is located at the interface of two parties: a *supplier* or *vendor* or *producer* as the party providing product and a *customer* or *consumer* as the party receiving product.

Traditional product control relied on the sorting principle at the interface of supplier and consumer. Modern product control primarily relies on quality provision and prevention at the interface. The particular task of the interface is to integrate supplier and consumer into joint and concerted efforts to create and maintain favorable conditions for quality in design, manufacture, handling, transport, and service. In this scheme, inspection may be useful, however not as the primary tool of product control, but in an adjunct role, subordinate to the guiding principles of provision and prevention.

The adjunct role of SPI is expressed in [2] in marked contrast to the sorting view of Dodge and Romig [1]: “The shipper must understand the acceptance sampling plan in order to be sure that the quality control of the production process provides lots of a quality sufficient to meet his objective. This is the basic principle of acceptance sampling – to assure the production of lots of acceptable quality. The primary purpose is not to discriminate between acceptable and nonacceptable lots.” The subordinate role of SPI is connected with the following purposes:

1. **Information and quality measurement purpose:**
SPI serves as an economical supplementary source of information on the product quality level. In particular, SPI is used to learn about the product quality level if the supplier–consumer interface is new and sufficient quality records do not exist.
2. **Incentive purpose:**
Continuous SPI puts an incentive on the supplier’s quality efforts.
3. **Protection purpose:**
SPI protects against sudden and unexpected shifts in the quality level.
4. **Precaution purpose:**
SPI serves as a precaution for the consumer, if the supplier–consumer interface is new or if the quality efforts of the supplier are not completely reliable.

Adaptiveness and Prior Knowledge on Product Quality

The interest in inspection varies along the history of the supplier–consumer interface. Continuing satisfactory quality records enhance mutual confidence

and reduce the interest in precautionary inspection. Confronted with deteriorating quality, the supplier will feel the need for tighter inspection. Accordingly, the design and the implementation of SPI should provide an adaptive learning mechanism, which methodically varies the amount of inspection (**sample size**) along with changes in the supplier–consumer relations. In particular, the adaptation mechanism should be able to react on continuing satisfactory quality records by not only changing the sample size but also by redefining the purpose of inspection, e.g., by switching to the no-sampling alternative.

Most customary acceptance **sampling schemes** are mechanisms to vary the sample size as a function of product quality records: the switching rules in the MIL-STD series (*see* **Attributes Sampling Schemes in International Standards**); sample size varying with the process proportion nonconforming in the Dodge–Romig rectification scheme (*see* [1] and **Rectification Sampling Schemes**); sample size varying according to some prior information on the distribution function of product quality (*see* **Economic Sampling Schemes; Prior Information in Sampling Schemes**). However, all these approaches give no guidelines on when to switch to the no-sampling alternative. Feigenbaum [3] and Deming [4] give all-or-none rules to choose between rectifying 100% inspection and acceptance without inspection on grounds of a linear cost model with the process proportion nonconforming as an input parameter; see the presentation of Deming’s $k_1 - k_2$ inspection policy in **Skip Lot and Chain Sampling**. An additional rule vaguely recommends sampling inspection only if the process proportion nonconforming is “close” to some breakeven point p_0 . Joyce Orsini’s scheme, as presented by Deming [4], specifies the meaning of “close to breakeven point” by an exact interval I around p_0 , and recommends sampling inspection with sample size $n = 200$, if the process proportion nonconforming is in I . Such rules are simple to administer, but their basis is heuristic.

Obviously, a suitable mechanism of representing and exploiting prior knowledge and forecasts on product quality is a necessary condition for adaptiveness of an SPI scheme. The absence of prior knowledge parameters is one of the weaknesses of the popular MIL-STD scheme. Most other inspection schemes are based on some more or less refined assumption on prior knowledge on product quality. The analysis of Feigenbaum [3] and Deming [4]

stipulates at least some interval of variation of the process proportion nonconforming to be known. The Dodge–Romig rectification scheme (*see* [1] and **Rectification Sampling Schemes**) is based on a similar assumption. Bayes schemes, as discussed by **Prior Information in Sampling Schemes**, assume the complete distribution function of product quality to be known – an assumption, however, which is hardly ever met in industrial practice. A reasonable alternative is to assume some incomplete knowledge of the distribution function of product quality. A simple and practicable instance of this approach is the α -optimal sampling scheme (*see* **Economic Sampling Schemes**).

Among industrial SPI standards, only ISO 13448 provides a systematic adaptive mechanism that depends on information on product quality and covers the entire range between no sampling and 100% inspection. ISO 18414 at least allows for successive reduction of sample size in response to a good quality history (*see* **Sampling in Industrial Standards** for more detailed descriptions). Both standards are recent, issued in 2004, 2005, and 2006, and broad industrial experience is still missing.

The importance of collecting and using information on product quality has been propagated by quality assurance literature and the ISO 9000 standard. Modern industrial environments are characterized by highly automated and intensively monitored processes of manufacture, handling, and shipping. Product quality data are often exchanged at the supplier–consumer interface. Both in theory and in practice of SPI, a satisfactory and generally accepted framework for exploiting such prior information is still missing. The issues of adaptiveness and prior information continue to be imminent challenges for SPI.

The Economic Concern of SPI

Product control is a business activity. As such, it is subject to an economic cost–benefit analysis: costs incurred by control activities have to be compared with the economic benefit obtained from maintaining a specified product quality level. The following economic parameters are relevant for product **inspection policies**: cost of inspection, profit or loss incurred from accepting or rejecting product of a specified quality level. See [3] for a detailed economic analysis.

The economic concern of SPI is ideally reflected by an *economic SPI scheme* with an economic model, an economic objective function, and an economic design built on optimizing the objective function (see **Economic Sampling Schemes** for details). Sometimes rigorous economic designs may be impractical because the parameters of the economic model are not fully known. In such cases, other designs may be used (see **Sampling Inspection of Products** for a survey). In any case, however, the values of the central design parameters, in particular the values of critical quality limits, should be chosen on grounds of economic reasoning.

The analysis and due consideration of quality economics, in particular the economic cost–benefit analysis of quality control activities, has been propagated by quality management literature (see [3] or [4]) and by industrial standards like the ISO 9000-9004 series (see [5]). However, this propaganda had little impact on the practice of SPI. The MIL-STD series and its national and international derivatives continue to be the most widely used SPI schemes in industry. These standards contain neither an economic model nor economic guidelines on how to choose their central design parameter, the acceptable quality level (AQL). As stated by Deming [4], p. 430, it is no astonishment to learn that MIL-STD “can lead to a plan whose total cost is double the cost of 100 per cent inspection”. The situation is not better for the Dodge–Romig scheme: the choice of central design parameters like LTPD (lot tolerance percent defective) and AOQL (average outgoing quality limit) is not supported by economic guidelines. None of the national or international SPI standards has adopted a precise economic framework. Economic SPI schemes, even of simple type as the scheme considered by **Economic Sampling Schemes** or of convenience type as Joyce Orsini’s rules, have not been accepted by industrial practice.

The current practice of SPI does not reflect the economic concern of SPI. It is to be feared that many applications rather involve redundant costs than contribute to the economical control of product quality.

Definition of Quality Levels

The purpose of product control is to assure an acceptable product quality level. An adequate and precise

definition of quality levels refers to a quantitative product quality indicator γ (see **Sampling Inspection of Products**). Quality levels are distinguished by means of threshold values of γ . In case of a simple binary quality model, an acceptable quality level is distinguished from an unacceptable (rejectable) quality level by means of a *breakeven quality* or *breakeven point*. **Sampling Inspection of Products** and **Economic Sampling Schemes** discuss the derivation of the breakeven quality from an economic model.

A precise distinction between acceptable and rejectable product quality by a breakeven point is postulated by **quality management** literature, e.g., [3] or [4]. Among SPI schemes, only economic schemes support the definition of a break even quality. Customary SPI schemes like MIL-STD or Dodge–Romig scheme use a variety of parameters seemingly related to quality levels, like AQL, LTPD, CQ (consumer’s quality), PQ (producer’s quality), LQL (limiting quality level), and RQL (rejectable quality level) (see **Sampling Inspection of Products**). The usefulness of these parameters is restricted: (a) they are not defined uniquely and precisely and most definitions are purely operational; (b) the role as quality thresholds is never made explicit; and (c) there are no guidelines on how to assign values to the parameters.

Most problematic is the AQL in MIL-STD. MIL-STD-105E [6] vaguely defines as follows: “When a continuous series of lots is considered, the AQL is the quality level which, for the purposes of sampling inspection, is the limit of a satisfactory process average.” Diverging interpretations of this definition have been given. Some authors consider the AQL directly as a breakeven quality. Others, like the authors of [7] or [3], establish functional relationships between Breakeven qualities calculated on the grounds of economic models and the AQL in MIL-STD. Although the cost models are the same, the relations to the AQL differ drastically. The disastrous consequence is that totally different sampling plans result from the same economic constellation. Confronted with such confusion, industrial practice has widely refrained from determining specific AQL values in specific business situations. As a pragmatic solution, AQL values are negotiated by contractual parties, or dictated by large companies, or recommended as standards in specific industrial segments. The sampling plans chosen under such precepts make no reasonable contribution to product

control. On the contrary, they incur additional costs, and lead to barely comprehensible and potentially dangerous conclusions.

SPI under Modern Quality Requirements

Modern industrial environments prescribe tight quality requirements and achieve outstanding quality levels. The Six Sigma scheme aspires to a proportion nonconforming level below 3.4 parts per million (ppm). [8] reports on plants and companies achieving extremely small nonconforming rates between 0 and 2 ppm over long periods. In contrast, customary SPI schemes consider the smallest proportion nonconforming thresholds of 0.01% AQL in MIL-STD-105E or 0.5% LTPD in the Dodge–Romig scheme. MIL-STD-105E [6] overtly admits that it is useless under strict quality requirements: “The sampling plans described in this standard are applicable to AQL’s of 0.01 percent or higher and are therefore not suitable for applications where quality levels in the defective parts per million range can be realized.” Strangely enough, MIL-STD sampling plans with AQLs up to 4% are used in industrial segments that in other respects are striving for extremely tight quality levels, like the electronic or nuclear industry. Such applications are merely pro forma, without any real impact on product quality. Their only effect is to cause additional costs. Among industrial SPI standards, only ISO 14560 accounts for relatively small proportion nonconforming thresholds, with the smallest so-called LQL at 500 ppm (*see Sampling in Industrial Standards* for more information). Broad industrial experience with this recent standard issued in 2004 is still missing.

The formation of sound inspection schemes for high-quality product at economically tolerable sample sizes is an imminent challenge for SQC. This task is closely related to two other aspects discussed in this article: (a) accounting for previous information on product quality is essential for reducing sample size; (b) the problem of measuring the product quality level is particularly serious in case of extreme values of product quality indicators.

Quality Measurement as a Purpose of SPI

Montgomery [9, p. 646], states: “It is the purpose of acceptance sampling to sentence lots, not to estimate the lot quality. Most acceptance sampling plans

are not designed for **estimation** purposes.” This is a correct description of the present state of acceptance sampling. However, requirements for modern acceptance sampling as discussed above suggest that the purpose of acquiring information by measuring an unknown quality parameter, e.g., a proportion nonconforming, should play a much more important role than it presently does. This purpose requires sufficiently reliable and sufficiently accurate measurement procedures, taking into account any available prior knowledge on product quality. According to the specific situation and interest, the measurement targets may be lot or process parameters.

Measurement of Lot Quality Parameters

Consider the measurement of an unknown proportion nonconforming in a lot by sampling. In most customary approaches (*see Prior Information in Sampling Schemes*), prior knowledge is considered in the form of a nondegenerate probability distribution of lot proportion nonconforming. This approach is difficult to understand for practitioners, and is theoretically controversial. In fact, for a given lot the actual proportion nonconforming p is a fixed, but unknown, number. Therefore, it is more appropriate to represent prior knowledge on p by an interval $I = \{p \mid p_1 \leq p \leq p_2\} \subset [0; 1]$, which includes the fixed, but unknown, actual value of the proportion nonconforming. If the interval contains only one element, i.e., $I = \{p^*\}$, then the actual proportion is known, while the case $I = \{p \mid 0 \leq p \leq 1\}$ stands for the case that nothing is known about the actual proportion nonconforming.

Sampling-based measurement procedures to determine the unknown value of a proportion nonconforming, under prior knowledge represented by an interval $I \subset [0; 1]$, have been developed in [10]. The procedures are made available on a CD-ROM and can be used straightforwardly without the need for sophisticated calculations.

Measurement of Process Parameters

In modern industrial environments, plenty of data on the lot generating process are available. Such process data should be used for developing a stochastic process model. A stochastic model as introduced by [11] establishes a relation between the indeterminate process output on the one hand and variables representing the physical nature of the process on the other. Because the physical conditions

of the lot generating process are generally not known exactly, the stochastic model includes not only one probability distribution, but a set of potential probability distributions of the process output. If by means of the stochastic model it can be proved that the proportion nonconforming in a generated lot is satisfactorily small, then it is more efficient not to control lot quality directly by measuring the proportion nonconforming in the lot. Instead, sampling should be used regularly for measuring the actual values of the relevant process parameters to exclude process disturbances, which may spoil the lot quality. The necessary stochastic measurement procedures are introduced in [12].

Selection of the Measurement Target

The measurement target of a sampling procedure is either a lot quality parameter, e.g., lot proportion nonconforming, or parameters of the lot generating process.

Lot parameters should be chosen as the measurement target in the following cases: (a) there is insufficient knowledge on the lot generating process and no stochastic process model is available; (b) an available stochastic model for the lot generating process indicates unsatisfactory process quality.

The choice of process parameters as measurement targets requires the availability of a stochastic model for the lot generating process. Under this provision, process parameters should be focused on in the following cases: (a) the process model indicates satisfactory process quality; (b) the required degree of measurement accuracy is very high. In this case, the sample size in pure lot quality measurement procedures as described above is uneconomical.

Future Principles for SPI

The analysis of the preceding sections implies that the following principles in the theory and practice of SPI will be useful in modern industrial environments:

1. Define SPI as a subordinate tool in a product control scheme based on provision and prevention.
2. Make use of the quality relevant information that is available in modern product control. In particular, make use of information on the product generating process.
3. Develop and implement SPI schemes that are adaptive to varying degrees of quality and that are able to switch to the no-sampling alternative.
4. Integrate precise and substantial economic reasoning into the design of SPI schemes.
5. Develop and implement SPI schemes for the control of high-quality products.
6. Consider SPI as tool for acquiring information on quality. Develop procedures to measure the actual values of lot and process parameters reliably and accurately.
7. Abandon the use of MIL-STD and its national and international derivatives. These standards are obsolete, since they have no sound theoretical foundation and do not conform to any of the above principles (1) through (6).

References

- [1] Dodge, H.F. & Romig, H.G. (1959). *Sampling Inspection Tables*, 2nd Edition, John Wiley & Sons, New York.
- [2] Godfrey, A.B. & Mundel, A.B. (1984). Guide for selection of an acceptance sampling plan, *Journal of Quality Technology* **16**, 50–55.
- [3] Feigenbaum, A.V. (1961). *Total Quality Control*, McGraw-Hill Book Company, New York.
- [4] Deming, W.E. (1986). *Out of the Crisis*, MIT Press, Cambridge.
- [5] International Organization for Standardization ISO 9004. (2000). Quality management systems – guidelines for performance improvements, International Organization for Standardization, Geneva.
- [6] *United States Department of Defense Military Standard, Sampling Procedures and Tables for Inspection by Attributes (MIL-STD-105E)*, U.S. Government Printing Office: Washington, DC, 1989.
- [7] Enell, J.W. (1984). Which Sampling Plan should I Choose? *Journal of Quality Technology* **16**, 168–171.
- [8] Kinni, T.B. (1996). *America's Best: Industry Week's Guide to World-Class Manufacturing Plants*, John Wiley & Sons, New York.
- [9] Montgomery, D.C. (2005). *Introduction to Statistical Quality Control*, 5th Edition, John Wiley & Sons, New York.
- [10] Collani, E. & Dräger, K. (2001). *Binomial Distribution Handbook for Applied Scientists and Engineers*, Birkhäuser, Boston.

6 Acceptance Sampling in Modern Industrial Environments

- [11] Collani, E. (2004). Theoretical stochastics, in *Defining the Science Stochastics*, E., Collani, ed, Heldermann Verlag, Lemgo.
- [12] Collani, E. (2004). Empirical stochastics, in *Defining the Science Stochastics*, E., Collani, ed, Heldermann Verlag, Lemgo.

ELART VON COLLANI AND RAINER GÖB

Sampling Inspection of Products and Statistical Process Control

Introduction

Acceptance sampling and statistical process control (SPC) are the most frequently used methodologies of statistical **quality control** (SQC). In acceptance sampling, the result of evaluation of a random sample taken from a group of products (a *lot*) or services is used for making decisions about these products or services. Therefore, acceptance sampling is aimed at product or service disposition decisions, based on statistical testing of stated quality requirements. As opposed to classical acceptance sampling, **control charts**, the best-known tools of SPC, are used to find whether monitored processes are in a “state of statistical control”. In other words, control charts are primarily used as a management tool to verify process stability. For this application of control charts any quality requirements need not be taken into account.

For industrial practice, the aspect of satisfying quality requirements is at least as important as guaranteeing process stability and eliminating sources of process variation. *Acceptance control charts*, introduced by Freund [1, 2], combine features of acceptance sampling and control charts, i.e., they can be used both for monitoring process stability and for verifying compliance with product or service quality requirements. Compliance with quality requirements is a separate issue, which has to be considered on equal footing with the interest of eliminating all sources of process variation. Quality requirements may be satisfied in the presence of small shifts of the process level around its target values. For example, if uniform large batches of raw material are used during a production process, the differences between the process levels related to between-batches variability may be tolerable. Similarly, small infrequent and irregular shifts of the process levels that cannot be eliminated because of economic or engineering (possibility of “overcontrol” of the process) considerations may be tolerated if their impact on the quality of the product is negligible.

Acceptance Control Charts: Design and Operation

The design principles, and the operation of acceptance control charts, have been described in the international standard ISO 7966 [3]. The standard considers the acceptance control of a process characterized by observations X_i , normally considered as continuous measurements, where the procedures can also be used for discrete variables like conforming/nonconforming indicators. It is assumed that there exist two characteristic process levels. The acceptable process level (APL) represents a process that should be accepted almost always, with probability $1 - \alpha$. On the other hand, the rejectable process level (RPL) represents a process that should almost never be accepted. When the process operates at the RPL level, or worse, the probability of its acceptance should be not greater than β . Thus, any process centered between the target value and the APL will have a risk smaller than α of not being accepted. On the other hand, the risk of acceptance of processes located further away from the target value than the RPL is smaller than β . The process levels lying in an “indifference zone” between the APL and RPL yield products of borderline quality. From the point of view of quality control, this “indifference zone” should be as narrow as possible. Unfortunately, the narrower this zone, the larger the **sample size** will have to be. Therefore, some economic considerations have to be present in the design process of acceptance control charts.

To operate an acceptance control chart it is necessary to define two additional parameters: the sample size n , and the action criterion or acceptance control limit (ACL). If the sample average value of the quality characteristic $\bar{X}$, plotted on an acceptance control chart, falls beyond ACL (i.e., above the upper acceptance control limit ACL_U or below the lower acceptance limit ACL_L , in case of double specification limits) the process shall be considered nonacceptable. The location of the APL and RPL levels, and the ACL limits is presented in Figure 1.

In the case of two-sided specifications the process levels APL and RPL are located above and below the target value (see Figure 1). When the APL and α , and RPL and β are defined, the upper and lower

2 Sampling Inspection of Products and SPC

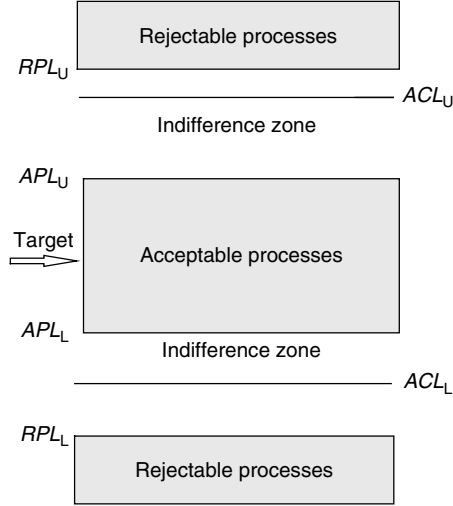


Figure 1 Basic elements of the acceptance control chart

ACLs are calculated as

$$ACL_U = APL_U + \left(\frac{y_{1-\alpha}}{y_{1-\alpha} + y_{1-\beta}} \right) \times (RPL_U - APL_U) \quad (1)$$

and

$$ACL_L = APL_L - \left(\frac{y_{1-\alpha}}{y_{1-\alpha} + y_{1-\beta}} \right) \times (APL_L - RPL_L) \quad (2)$$

respectively. The values of $y_{1-\alpha}$ and $y_{1-\beta}$ are the **quantiles** of the standard **normal distribution** of order $1 - \alpha$ and $1 - \beta$, respectively. The formulas for the calculation of the upper and lower ACLs remain the same in the cases of one-sided specifications or asymmetrical quality requirements. For the calculation of the sample size n , it is necessary to know the value of the standard deviation σ_w within the sample (rational subgroup), i.e. the standard deviation of the measured quality characteristic X , which describes its inherent and natural variability. The value of n is obtained by rounding-up to the nearest integer the value calculated from the following expression

$$n = \max \left\{ \left[\frac{(y_{1-\alpha} + y_{1-\beta})\sigma_w}{(RPL_U - APL_U)} \right]^2, \left[\frac{(y_{1-\alpha} + y_{1-\beta})\sigma_w}{(APL_L - RPL_L)} \right]^2 \right\} \quad (3)$$

Let p_0 be the acceptable fraction of nonconforming items, at the output of the process, and p_1 be the rejectable fraction of nonconforming items. If these two values are known from economic considerations, and the values of the lower specification limit (LSL) and the upper specification limit (USL) are given, we can calculate the values of APL and RPL from the following formulas:

$$APL_U = USL - y_{1-p_0}\sigma_w \quad (4)$$

$$APL_L = LSL + y_{1-p_0}\sigma_w \quad (5)$$

$$RPL_U = USL - y_{1-p_1}\sigma_w \quad (6)$$

$$RPL_L = LSL + y_{1-p_1}\sigma_w \quad (7)$$

The formulas given above are valid when the quality characteristic is normally distributed. In cases of other probability distributions the quantiles of the standardized normal distribution should be replaced with the respective quantiles of the standardized versions of those distributions.

Example

Bottles of nominal value of 1000 cm^3 are filled with sunflower oil. It has been verified that the amount of oil in a bottle is described by the normal distribution. The filling process is considered acceptable if less than 1% of the bottles does not meet the requirement described by the LSL of 995 cm^3 and USL of 1005 cm^3 . The process is considered nonacceptable if the percentage of the bottles that does not meet these requirements exceeds 5%. The standard deviation σ_w , estimated from historical data, is equal to 1.5 cm^3 . The risks α and β are set at the same value, 0.05.

To design the appropriate acceptance control chart we have to do the following computations:

Step 1. For $p_0 = 0.01$ we have $y_{0.99} = 2.326$ and $y_{0.96} = 1.751$.

Step 2.

$$APL_U = 1005 - 2.326 \times 1.5 = 1001.51 \quad (8)$$

$$APL_L = 995 + 2.326 \times 1.5 = 998.49 \quad (9)$$

$$RPL_U = 1005 - 1.751 \times 1.5 = 1002.37 \quad (10)$$

$$RPL_L = 995 - 1.751 \times 1.5 = 997.63 \quad (11)$$

Step 3. The ACLs are calculated as follows:

$$ACL_U = 1001.51 + \left(\frac{1.645}{1.645 + 1.645} \right) \times (1002.37 - 1001.51) = 1001.94 \quad (12)$$

$$ACL_L = 998.49 + \left(\frac{1.645}{1.645 + 1.645} \right) \times (998.49 - 997.63) = 998.06 \quad (13)$$

Step 4. The initial calculation of the sample size gives

$$n = \left[\frac{(1.645 + 1.645) \times 1.5}{1002.37 - 1001.51} \right]^2 = 3.128 \quad (14)$$

Thus, the required sample size is $n = 4$.

Discussion

The acceptance control charts presented in this article are practically the same as acceptance sampling plans by variables. As a matter of fact, popular acceptance sampling plans by variables, presented in the international standards from the series ISO 3951 (see, for example, ISO 3951-1 [4]) should be rather

used to accept (or reject) process levels expressed in terms of fraction nonconforming and labeled by acceptance quality limits (AQLs), than to accept (or reject) individual lots of products. Accordingly, individual sampling plans from these standards could be used for the purpose of SPC. For example, the economically optimal control chart, based on the sampling plans taken from [4] (for unknown values of the standard deviation σ) is proposed by Hryniewicz [5].

References

- [1] Freund, R.A. (1957). Acceptance control charts, *Industrial Quality Control* **14**(4), 13–23.
- [2] Freund, R.A. (1960). A reconsideration of the variables control chart, *Industrial Quality Control* **16**(11), 35–41.
- [3] ISO 7966 (1993). *Acceptance Control Charts*.
- [4] ISO 3951-1 (2005). *Sampling Procedures for Inspection by Variables – Part 1: Specification for Single Sampling Plans Indexed by Acceptance Quality Limit (AQL) for Lot-By-Lot Inspection for a Single Quality Characteristic and a Single AQL*.
- [5] Hryniewicz, O. (2001). Application of ISO 3951 acceptance sampling plans to the inspection by variables in statistical process control (SPC), in *Frontiers in Statistical Quality Control*, H.-J. Lenz & P.-Th. Wilrich, eds, Physica-Verlag, Heidelberg, Vol. 6, pp. 80–92.

OLGIERD HRYNIEWICZ

Sampling in Pharmaceutical and Chemical Industries

Introduction

There are some characteristics of the pharmaceutical and chemical industries that make a difference on how sampling inspections are performed. The product characteristics are given by biological, chemical, and physical measurements. It is often related to batch production. At least for the pharmaceutical industry they are often highly regulated by government authorities with high demands on quality standards and level of documentation.

The measurement methods being used include both slow and expensive off-line methods and fast and cheap on-line measurements. The trend in these industries is that more of the fast and cheap on-line methods become available and to some extent replace the off-line methods. This also means that larger sampling and data evaluation becomes an integrated part of sampling inspection and batch release.

The US Food and Drug Administration (FDA) has recently launched a “process analytical technology (PAT) initiative” supporting this development, see [1]. A desired goal of the PAT framework is to design and develop processes that can consistently ensure a predefined quality at the end of the manufacturing process. Such procedures would be consistent with the basic tenet of quality by design and could reduce risks to quality and regulatory concerns while improving efficiency. Gains in quality, safety, and/or efficiency vary depending on the product and are likely to come from the following actions:

- reducing production cycle times by using on-line measurements and controls;
 - preventing rejects, scrap, and reprocessing;
 - considering the possibility of real-time release;
 - increasing automation to improve operator safety and reduce human error;
 - facilitating continuous processing to improve efficiency and manage variability
- using small-scale equipment (to eliminate certain scale-up issues) and dedicated manufacturing facilities,

- improving energy and material use and increasing capacity.

There are obviously large benefits to gain from this development for the customers in terms of quality, safety, and cost. There will still be a need to perform off-line measurements and classical sampling inspection and release testing in the pharmaceutical and chemical industries as there may be no alternative.

Sampling under Measurement Uncertainty

Sampling of discrete items in lots for inspection is mainly used in the medical device part of the industry. Figure 1 is an example of a person with diabetes injecting himself with a prefilled insulin pen. The example illustrates the critical environment for certain medical products, hence the need for documented high quality products. This is why, for this part of the industry, there may be a specific need for taking measurement uncertainty into consideration.

Sampling plans for inspection by attributes are given in the ISO standard series ISO 2859 [2]; and similar plans for inspection by variables are given by the ISO standard series ISO 3951 [3]. It is a general idea – at least among statisticians – that sampling by variables is a more efficient procedure for **acceptance sampling**, than sampling by attributes. This idea is supported by the fact that under a correct model, a parametric estimator of the fraction nonconforming is



Figure 1 Person with diabetes injecting himself with a prefilled insulin pen

more efficient than the nonparametric estimator based upon the crude count of nonconforming items in the sample.

In industrial praxis, an important issue is the presence of measurement errors. In the measurement situation, the ideal situation is of course that measurement errors are negligible compared to the tolerance interval of the item to be measured. This is not always the case, though. For medical devices there are often high requirements on certain critical parameters, e.g. the dosing accuracy. This again leads to high requirements on the individual components in the tolerance stack up. Both the ISO standards mentioned before [2, 3] consider the case of negligible measurement errors. Therefore, in the following, we shall discuss the implications of using acceptance sampling by variables under measurement error.

The theory for acceptance sampling by variables under a **normal distribution** dates back to the papers by Lieberman and Resnikoff [4] and [5]. The current ISO standard with sampling plans, ISO 3951 [3], is based upon this original idea of basing the test upon a minimum variance, unbiased estimate of the fraction nonconforming in the process. A description of the theory may be found in [6].

In [7], Owen and Chou, consider the effect of measurement error and a constant offset error on the operating characteristic (OC) curves (*see Single Sampling by Attributes and by Variables*) of the one-sided plans. In [8], Thyregod and Melgaard, however, extend these considerations to the double-sided plans in the case of measurement error and a random “laboratory” **bias**.

We shall assume that items in a lot are produced from an in-control process. Let X denote the quantity of interest. We shall assume that the distribution of X over the items in a lot may be described by independent and identically distributed (i.i.d.) random variables that follow a normal distribution with mean $E(X) = \mu$ and $V(X) = \sigma^2$. It should be noted, that in pharmaceutical and chemical processes the assumption of independent items in a lot often must be challenged. Simple graphical methods such as plotting the items in the order of production will often reveal correlations that must be taken into account.

Let the specification limits for X be upper limit, U , and lower limit, L , respectively. The fraction of nonconforming items, p , is the fraction of items below L or above U . In the related article **Variables Sampling under Measurement Error** the case of

high-frequency measurement error is discussed in detail.

High-Frequency Measurement Error

Assume that the quantity of interest cannot be measured without measurement error.

Let $Y = Y(X)$ denote the measurement result for an item with value X . Assume that

$$Y_i = X_i + E_i \quad (1)$$

where $E_1, \dots, E_n$ are i.i.d. and normally distributed with $E(E_i) = 0$ and $V(E_i) = \sigma_e^2$. Let, as usual,

$$\bar{Y} = \sum \frac{Y_i}{n} \quad (2)$$

and

$$s_y^2 = \sum \frac{(Y_i - \bar{Y})^2}{(n-1)} \quad (3)$$

Then $\bar{Y}$ follows a normal distribution with $E(\bar{Y}) = \mu$ and $V(\bar{Y}) = (\sigma^2 + \sigma_e^2)/n$, and s_y^2/σ^2 follows a $(1 + \sigma_e^2/\sigma^2)\chi^2(f)/f$ distribution with $f = n-1$, and $\bar{Y}$ and s_y^2 are independent. For simplicity we will use the OC curves, corresponding to the one-sided plans, which are similar to the ones given in ISO 3951. These will be close to the double specification limits case, when $\sigma \ll \sigma_{\max}$. It is shown in [8] that by using the same sampling plan from ISO 3951 – but in this case random measurement error is present – according to the model (1) the **OC curve** is given by

$$L_1(p) = P \left\{ \frac{U - \bar{Y}}{s_y} \geq k \right\} = P\{t_{n-1, \delta_1} \geq k\sqrt{n}\} \quad (4)$$

where $\delta_1 = -\sqrt{n}z_p/\sqrt{1 + \sigma_e^2/\sigma^2}$. One thing to remark is that the OC curve in the case of random measurement noise is always to the left of the noise-free OC curve, which means that in the case of high-frequency measurement noise; there is a higher chance of rejecting the batch. An example from ISO 3951 is given by Figure 2.

Low-Frequency Measurement Error

Assume instead that the major part of the measurement error is a random “laboratory” or instrument bias that affects all measurements in the same manner, i.e.

$$Y_i = X_i + B \quad (5)$$

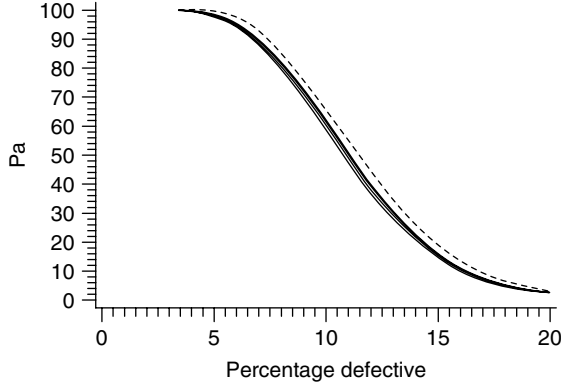


Figure 2 OC curves under high-frequency measurement error, $AQL = 6.5$, code K, $n = 50$, $\sigma_e = 0.2\sigma$, the dashed curve is for the model without measurement error

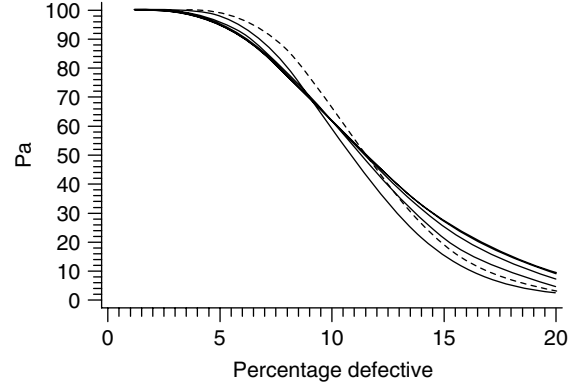


Figure 3 OC curves under low-frequency measurement error, $AQL = 6.5$, code K, $n = 50$, $\sigma_B = 0.2\sigma$, the dashed curve is for the model without measurement error

where B varies from lot to lot as i.i.d. and normally distributed with $E(B) = 0$ and $V(B) = \sigma_B^2$. Under this assumption we have that $\bar{Y}$ follows a normal distribution with $E(\bar{Y}) = \mu$ and $V(\bar{Y}) = \sigma_B^2 + \sigma_e^2/n$, and s_y^2/σ^2 follows a $\chi^2(f)/f$ distribution with $f = n - 1$, and $\bar{Y}$ and s_y^2 are independent. Thus, in this case the uncertainty of the estimate, $\bar{Y}$, of the position, is heavily influenced by the measurement uncertainty, but the estimate, s_y^2 of the process spread is not affected.

The OC curve for this model is then given by

$$\begin{aligned} L_2(p) &= P \left\{ \frac{U - \bar{Y}}{s_y} \geq k \right\} \\ &= P \left\{ t_{n-1, \delta_2} \geq \frac{k\sqrt{n}}{\sqrt{1 + n\sigma_B^2/\sigma^2}} \right\} \end{aligned} \quad (6)$$

where $\delta_2 = -\sqrt{n}z_p/\sqrt{1 + n\sigma_B^2/\sigma^2}$. If we let

$$n_* = \frac{n}{1 + n\sigma_B^2/\sigma^2} \quad (7)$$

in the equation for $L_2(p)$ we see, that the equation is very close to the noise-free one, but with a **sample size** of n_* . The only difference is that the **degrees of freedom** in the t distribution are not changed. This means, that approximately the OC curve in case of a systematic measurement error is similar to the OC curve of the noise-free case, but with a reduced sample size.

In the example given by Figure 3, we have a sample size of $n = 50$ and a measurement uncertainty of $\sigma_B = 0.2\sigma$. Approximately this corresponds to the OC curve of the noise-free case with a sample size of

$$n_* = \frac{50}{1 + 50(0.2)^2} = 16.7 \quad (8)$$

This fact is also revealed from Figure 3. It is important to notice the difference between the noise-free case and the situations with high-frequency as well as low-frequency noise. Especially for testing in the pharmaceutical and chemical industries the measurement uncertainty is not negligible. This requires that the measurement method is well documented, including quantification of the high- and low-frequency measurement error. Since the influence of a low-frequency measurement error can be large, as shown in the previous example, steps will normally be taken to minimize this effect, e.g. by calibrating between sampling of lots.

Bulk Sampling

Purpose and Procedure of Bulk Sampling

Chemical analyses of bulk materials as powder or liquids are often found in the pharmaceutical and chemical industries by its nature. Sampling can be used for the inspection of raw materials, intermediate products, or final product release. In this case a lot is a definite quantity of bulk material. The quality

of the lot is measured by a single suitable quality indicator, e.g. the content of active ingredient. It is assumed that the mean quality of the lot is to be determined and that this is the factor used for determining the acceptability of the lot. In this case, acceptance sampling plans and procedures for the inspection of bulk materials are found in ISO 10725 [9].

A number of increments (smaller volumes), k , of the same size are taken randomly from the lot. These increments are mixed into a gross sample. From this gross sample a laboratory sample is prepared and a number of samples, m , is taken from this laboratory sample and analyzed individually. The results, x_i , $i = 1, \dots, m$, are combined into a single mean value that is representative of the lot. The mean value of the analyzed samples is used as the estimate of the mean quality of the lot. The mean is calculated as

$$\bar{x} = \sum_{i=1}^m \frac{x_i}{m} \quad (9)$$

The uncertainty (variance) of this mean value is given by

$$V(\bar{x}) = \frac{\sigma_A^2}{k} + \sigma_B^2 + \frac{\sigma_C^2}{m} \quad (10)$$

where

- σ_A^2 is the variance of the increments from the bulk due to variations between containers and within containers,
- σ_B^2 is the variance associated with the preparation of the laboratory sample taken from the gross sample,
- σ_C^2 is the variance describing the uncertainty from preparing and analyzing the individual samples for analyses taken from the laboratory sample.

The model above is a simple statistical model describing the bulk sampling situation. To have a “statistical rationale” for the choice of k , the number of containers to sample from, and the choice of m , the number of chemical analyses to perform, some estimates of the **variance components** must be found, e.g. from designed statistical experiments, see [9].

Retesting of Bulk Materials

Retesting is when new laboratory samples are taken from the gross sample of a lot to be analyzed because

of suspicious analytical results. For specific chemical or biochemical laboratory analysis, a large number of steps are involved and the analyses are not always fully automated. In these cases there is a possibility of occasional (human) errors affecting the results. Therefore, there are well-described procedures for good manufacturing practice (GMP) including procedures for handling out-of-specification (OOS) measurement results, see [10]. In the case of discrete items handled in the section titled “Sampling under Measurement Uncertainty”, the ISO standards have built-in procedures on how to handle the situation of OOS results and batch rejection.

The procedure [10], applies to laboratory testing. It applies to the situation described previously where the mean quality of the lot is of interest. It does not apply if the purpose of the analyses is to measure uniformity of a lot, e.g., content uniformity, release profile, and powder blend.

To identify the cause of the OOS result, statistical **hypothesis testing** may be relevant. Hypothesis testing may consist of repetition of the test procedure or part of the procedure, or of experiments designed specifically with the purpose of identifying an analytical problem. Below is an example of retesting.

Use of t -Test to Compare OOS Results and Retest Results.

The t -test is appropriate in a retest situation where an OOS result is compared to a group of results obtained specifically, to compare it to the earlier result, as described by Søren Andersen [11]. The test statistic is given by

$$T = |OOS - m_{\text{retest}}| / (s_{\text{retest}} \sqrt{(1 + 1/m)}) \quad (11)$$

where m_{retest} and s_{retest} are the mean and standard deviation of the retest results and m is the number of retests. The retest samples should be representative of the lot being sampled, i.e., use the gross sample from the previous test or create a new gross sample based on increments randomly taken from the lot. The calculated test statistic, T , is compared to the critical value of the t -test, given e.g., in [11]; if T is greater than the 5% critical value, the OOS is considered a statistical **outlier**. If the OOS result is considered a statistical outlier, we have a strong indication that the result is due to a laboratory error. A qualified person should then be able to evaluate the necessary steps to eventually release the batch.

References

- [1] FDA PAT initiative: <http://www.fda.gov/Cder/OPS/PAT.htm>, 2005.
- [2] ISO 2859, series *Sampling Procedures for Inspection by Attributes*, International Organization for Standardization, 1992–2007.
- [3] ISO 3951. (1989). *Sampling Procedures and Charts for Inspection by Variables for Percent Nonconforming*, International Organization for Standardization, Geneva.
- [4] Resnikoff, G.J. (1952). A new two-sided acceptance region for sampling by variables, Technical Report No. 8, Applied Mathematics and Statistics Laboratory, Stanford University, Stanford.
- [5] Lieberman, G.J. & Resnikoff, G.J. (1955). Sampling plans for inspection by variables, *Journal of the American Statistical Association* **50**, 457–516.
- [6] Schilling, E.G. (1982). *Acceptance Sampling in Quality Control*, Marcel Dekker, New York.
- [7] Owen, D.B. & Chou, Y.-M. (1983). Effect of measurement error and instrument bias on operating characteristics for variables sampling plans, *Journal of Quality Technology* **15**, 107–117.
- [8] Melgaard, H. & Thyregod, P. (2001). Acceptance sampling by variables under measurement uncertainty, in *Frontiers in Statistical Quality Control*, H.-J. Lenz & P.Th. Wilrich, eds, Physica-Verlag, Heidelberg, Vol. 6, pp. 47–57.
- [9] ISO 10725. (2000). *Acceptance Sampling Plans and Procedures for the Inspection of Bulk Materials*, International Organization for Standardization.
- [10] *Guidance for Industry Investigating Out-of-Specification (OOS) Test Results for Pharmaceutical Production*, Center for Drug Evaluation and Research (CDER) and U.S. Food and Drug Administration, 2006.
- [11] Andersen, S. (2004). *An Alternative to the ESD Approach for Determining Sample Size and Test for Out-of-Specification Situations*, Pharmaceutical Technology, May 2004.

HENRIK MELGAARD

Sampling in Semiconductor Manufacturing

Introduction

A state-of-the-art complementary metal oxide semiconductor (CMOS) manufacturing process combines a myriad of transistors connected by several metalization layers on a tremendously small area. This is achieved by several hundred processing operations, i.e., each being a physical treatment on processing equipment. The processing results are verified with certain types of metrology. For an overview refer to [1, 2] and references therein. Typical metrology applications are:

- defect inspection (visual, microscope)
- critical dimension scanning electron microscope (SEM) measurement
- overlay registration measurement
- thickness measurement using ellipsometry
- step height and profile measurement using profilometers, atomic force microscope (AFM), and scatterometry
- electrical tests on test structures (e.g., single resistors or transistors, ring oscillators).

Since semiconductor manufacturing is performed in clean-room environments, the available floor space for process and especially metrology equipment is limited. Furthermore, a high production line throughput with low cycle times enforces tight metrology budgets. Consequently, sampling is about balancing metrology objectives and constraints.

Examples: (a) A very stable and established process step may receive only moderate defect inspection and loose parametric control. This may apply in a proved process to a hydrofluoric (HF) clean or a furnace- poly-silicon deposition. (b) Conversely, a very critical process step, which may even operate on a novel tool platform, must be followed by several metrology steps, with very high sampling rates. This may, for example, apply to the whole gate formation process sequence.

A typical semiconductor manufacturing scenario is grouping several (e.g., 25), silicon wafers into lots and processing them at a processing operation.

Whenever a metrology inspection is performed only a subset of its wafers, and of course, only certain positions on a wafer, receive the metrology inspection.

The structure of the sampling decision depicted by Figure 1 is organized as follows:

(a) Select a lot, (b) pick wafers within the lot, and (c) decide for positions on the wafers. The three levels of a sampling decision are denoted by *lot sampling*, *wafer sampling*, and *site sampling* in the following text.

In modern 300-mm manufacturing, the actual task of sampling must be fully automated in order to practically account for the huge number of influencing variables.

This means that the tasks have to,

1. make the sampling decision;
2. dispatch and transport the lot to a metrology station;
3. select the to-be inspected wafers and the metrology recipe (to-be measured features and their positions);
4. start the inspection process on the metrology equipment;
5. acquire the measured data;
6. provide measured data to attended customers (e.g., **statistical process control** (SPC), advanced process control (APC), data mart, etc.);
7. remove the lot from metrology station.

These tasks are performed completely automatically by state-of-the-art MES (manufacturing execution system), EI (equipment interface), and other FAB (fabrication plant for semiconductor manufacture) automation and IT (information technology) infrastructure. The above tasks and their application

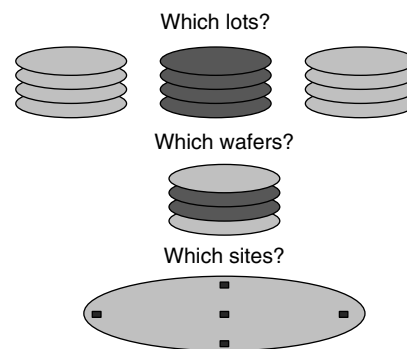


Figure 1 Structure of sampling decision

2 Sampling in Semiconductor Manufacturing

on the manufacturing floor do, nonetheless, allow human intervention and may in certain cases still require human interaction.

The subsequent paragraphs discuss in more detail the task 1 – to make the sampling decision.

Sampling Decision

State-of-the-art semiconductor manufacturing depends on relevant metrology data in order to monitor and maintain the process, thereby minimizing the risk of processing errors/faults. The term *relevant data* is directly linked to the sampling decision, which has to provide a representative part of a population for the purpose of determining parameters or characteristics of the whole population [3, 4]. The decision process itself is restricted by certain expense criteria. In conclusion, sampling has to select the appropriate material for metrology based on manufacturing needs (also known as *metrology objectives*) while fulfilling manufacturing constraints (also known as *metrology cost*) [5].

Metrology Objectives

Typical state-of-the-art semiconductor manufacturing sampling follows the metrology objectives given below, with the goal to minimize the risk for misprocessing:

- A1. Reconfirm correct processing after wafer-external production disturbance, i.e., logistic (e.g., **preventive maintenance**, tool repair), or manufacturing (e.g., batch changes, recipe adjustments).
- A2. Reconfirm correct processing after wafer-intrinsic production disturbance, e.g., detected processing faults (after metrology or sensor signal) should trigger verification metrology for affected lots/wafers to decide about further steps like rework or even wafer scrap.
- B1. Maintain correct processing by detecting unexpected disturbances, i.e., measure regularly, all suspected sources of variation like entity, process recipe, product, processing order, and used chambers.
- B2. Maintain correct processing by detecting expected disturbances, i.e., provide data to quantify control adjustment by (a) measuring

regularly (static), all sources of variation that are known/modeled and (b) adapting (dynamic) to the needs in case of performance issues.

- C. Reconfirm correct processing according to particular request, e.g., external customer request, or engineering request.

The elements defining a specific abstract metrology objective are denoted as *context elements* in the following text. Typical context elements are (a) on lot level: equipment, preventive maintenance, product; (b) on wafer level: used chambers, defect count; and (c) on-site level: wafer zone, site signature.

Metrology Constraints

Besides the aim to minimize the risk, sampling is also about minimizing the cost for metrology. Metrology expense factors are:

- A. Individual lot cycle time limitations, e.g., wafer-carrier transportation, idle state of non-inspected wafers.
- B1. Metrology acquisition expenses, e.g., tool, clean-room floor space, supply/waste facilities.
- B2. Metrology operating expenses, e.g., service, maintenance, **calibration**, **power** supply.
- C. Metrology capability limitations, e.g., wafer contamination with particles/deposits, destruction of test structures by metrology itself.

Sampling Decision Process

Metrology objectives are imposed by the risk of misprocessing while metrology constraints are defined by certain budgets. The key purpose of sampling is to balance these two opposed factors, i.e., to minimize the risk of misprocessing at tolerable metrology expenses. Therefore, the sampling decision should incorporate a weighting of the metrology objectives, e.g., by organizing them in a hierarchical structure.

Furthermore, it cannot be assumed that any subpopulation of context elements is independent and identical; thus a simple arrangement of context elements is not sufficient.

This leads to additional complexity beside the hierarchical structure, because the sampling decision must also incorporate dependencies among context elements.

The sampling decision is the more optimal, the more simultaneous metrology objectives can be met

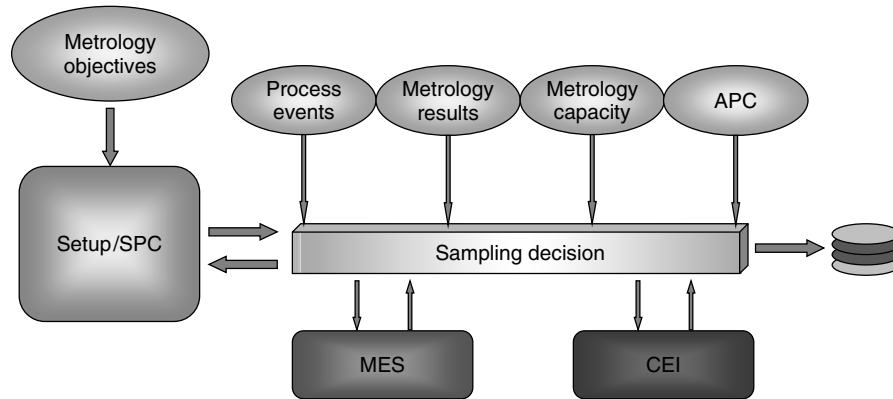


Figure 2 Sampling system overview

at once. The decision system must, therefore, detect and actively create such coinciding situations. This enables the possibility of low sampling rates, which also means low metrology cost, while achieving best risk mitigation.

In principle, the sampling decision lot/wafer/site is achieved by connecting lot/wafer/site context elements to sampling rules, e.g., a fixed minimum sampling rate (objectives B), or forced metrology (objectives A). Thereby, the ability to mix the granularity (lot/wafer/site) of sampling decision and metrology objective is a must, e.g., a site-level metrology objective may trigger the decision on lot level. However, even a coarser granularity, i.e., focusing on a batch of lots can be practical and needed.

The sampling decisions with respect to lots, wafers, and sites differ in nature. The lot sampling decision is binary in its result. A lot can either receive metrology or not, mostly independently, from other lots or affected wafers/sites owing to its associated high metrology cost on lot level, e.g., transport. Since the incurred cost of measuring additional wafers in many semiconductor metrology applications is small, the wafer sampling task is mostly to pick the optimal wafers satisfying many metrology objectives at once. Basically, site sampling specifies which feature should get measured at which position of a wafer. Thereby, the term “feature” denotes a subtask or sub-recipe of a metrology recipe. It should be mentioned that most current semiconductor metrology tools do not currently support automatic position and/or feature selection. Actual state-of-the-art site sampling means selecting predefined metrology recipes which serve specified metrology objectives.

Implementation of Sampling

An example for a state-of-the-art sampling solution is presented schematically in Figure 2 and will be denoted as *sampling system*. The sampling system provides a framework that has the capacity to:

- incorporate weighing of the metrology objectives;
- incorporate dependencies among context elements;
- incorporate complexity in a reasonable way;
- choose the appropriate lots and pick the optimal wafers/sites.

In general, engineering knowledge is crucial in order to specify the metrology objectives. Therefore, an integral part of the sampling system is a kind of content management system providing the ability to easily and transparently specify the high-level metrology objectives. Furthermore, the sampling system is linked to SPC systems in order to verify metrology results [6].

Lot Sampling

Of course, in order to end up with acceptable overall sampling rates, the system requires some degree of freedom. The metrology objectives (context elements) determine the *sampling rules* that are applied to the lots. Among the sampling rules that can be applied the most common ones are (a) Simple sampling rules like “measure 30% according to a deterministic pattern”. (b) Check explicitly for metrology gaps only, e.g., number of lots not measured or a

4 Sampling in Semiconductor Manufacturing

certain time window. (c) Flexible sampling rules targeting a certain sampling rate over a period of time.

The sampling system also incorporates a weighing of metrology objectives. For example, assume that objectives A1, A2, and C are associated with high risks, while B1 and B2 are associated with lower risks. The sampling system can weigh an event of objective B1 or B2 with the possibility to achieve lower sampling rates when waiting for a simultaneous event of objective A1 or A2.

The circumstance of incorporating dependencies among context elements is met by the fact that the sampling system is capable of combining arbitrary context elements. This allows, for example, straightforward setup of rules like “critical products” with “latest technology generation” with “suspected tool set”.

In order to incorporate the complexity in a reasonable way, the context elements representing the metrology objectives are organized into a hierarchically structured finite-state automaton. The automata approaches are known to be very well suited for computer simulations and have therefore become quite popular in various fields [7–9]. Thus, all the above objectives are organized into simple local rules reproducing the complex sampling tasks by their dynamical interaction.

The following is a typical lot sampling example from lithography. It illustrates how the sampling

system reconciles multiple overlapping rules with the fewest possible measured lots. Assume that the metrology objectives are:

- A1. Measure the next lots after certain events, e.g., resets.
- A2. Consider particular lots, e.g., rework, yield learning or development.
- B1. Ensure a minimum sampling rate as follows: monitor (a) the process entity, (b) the combination of process entity and operation, and (c) the product, i.e., AMD Opteron, each with an individual baseline and minimum sampling rate of $X\%$.
- B2. Consider APC application performance.

Figure 3 shows that the sampling system is able to achieve a reduction in the overall sampling rate of more than 20% relative to an additive sampling approach. The sampling rate (simulated) of an additive approach is represented by the tallest bar. However, the effective sampling rates established by the sampling system are represented by the next three bars, in respect to metrology objectives B1. Note that each of these three bars is a combination of individual sampling rates setup, e.g., critical operations or entities equipped with higher rates. The remaining bars reflect how often the rules connected to the above metrology objectives, i.e., A1, A2, B2, and minimum

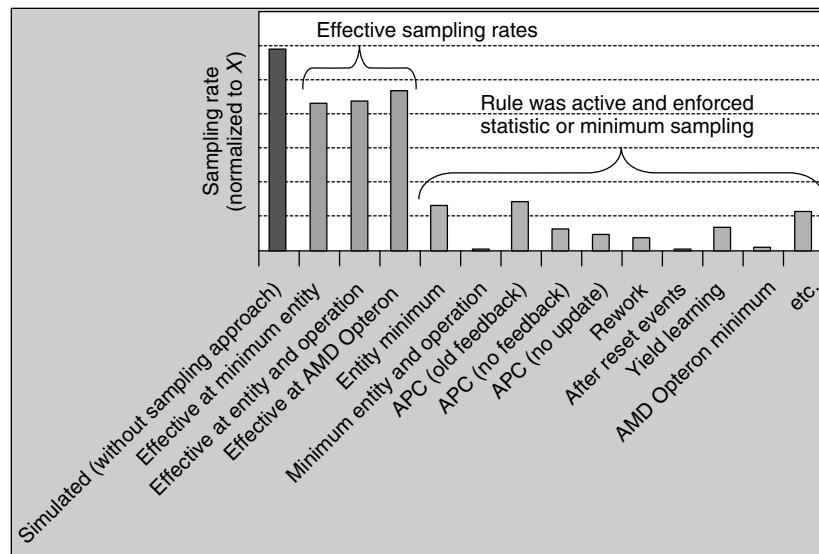


Figure 3 Lot sampling example

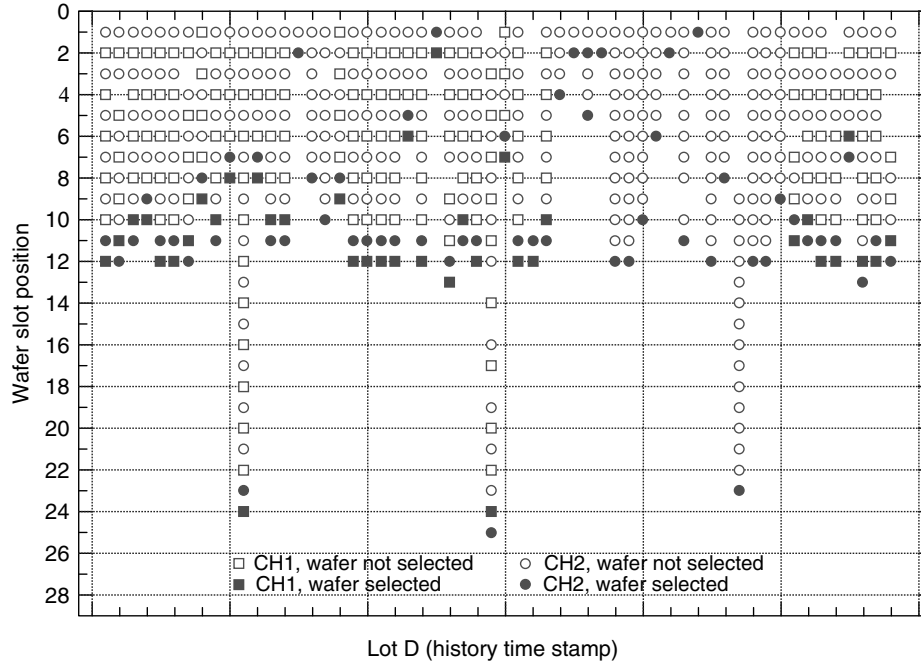


Figure 4 Wafer sampling example

sampling rates of B1, have become active. They further illustrate that the balancing among the sampling objectives works quite well, e.g., the rule requiring a minimum sampling rate on the process entity already ensures that the combination of process entity and operation, or measurement of AMD Opteron product is mostly satisfied.

Wafer Sampling

The wafer sampling task, to pick the best wafers satisfying the metrology objectives, is hard to satisfy, considering the tight metrology budgets. In general, it is not possible to comply with all metrology requirements at every metrology step, in the sense that all wafers can receive metrology. To solve this problem all rules are assigned with a certain penalty. A mixed integer linear program (MILP) is used to minimize the number of selected wafers according to the assigned penalties [10, 11]. The penalties can also be increased after rule violation in order to satisfy all requirements with a finite frequency in course of time.

An example of wafer sampling is given in the following text. Assume that the metrology objectives are B1 to monitor all used chambers and B2 to prefer

the last processed, meaning most likely contaminated wafers (highest slots).

In Figure 4 the results for wafers selected by the sampling system are depicted for the given example. Each column represents a single lot consecutively (x axis) processed on a two-chamber tool (CH1, CH2) at a certain process operation. The symbols within a column represent the wafers within a lot positioned at a certain slot (y axis). A filled symbol corresponds to a selected wafer and the chambers where the wafers were processed are distinguished by the shading. It can be seen that the wafer sampling systems picks both chambers (two wafers), if possible (objective B1), and prefers the last wafer processed (objective B2). The last wafer processed, is always one that is positioned in the highest slot.

Assuming that as a further restriction, only one wafer may be selected per metrology run, the algorithm would decide to pick both chambers alternating, if equipped with the same penalty.

Site Sampling

Finally, the site sampling task has to be performed for the wafers selected. This is organized as follows:

(a) selection of the appropriate metrology recipe for a wafer, e.g., standard short recipe or special long recipe; (b) selection of the metrology features to be performed for the chosen recipe; and (c) selection of the metrology sites (positions) where the features have to be performed. Comparable to lot Sampling (a), (b) and (c) are determined by context elements and may also be connected to sampling rules, e.g., perform standard recipe for 90% of wafers while for 10% the special recipe is performed. As mentioned before, most current semiconductor metrology tools do not currently support automatic position and/or feature selection. The site selection process is, therefore, often reduced to an automated recipe selection.

References

- [1] May, G.S. & Spanos, C.J. (2006). *Fundamentals of Semiconductor Manufacturing and Process Control*, Wiley-IEEE Press.
- [2] International Technology Roadmap for Semiconductors (ITRS), <http://www.itrs.net> Semiconductor Industry Association, 2005.
- [3] William, F. (1968). *An Introduction to Probability Theory and Its Applications*, John Wiley & Sons, New York.
- [4] Leonard, K. (1975). *Queueing Systems*, John Wiley & Sons, New York.
- [5] Hopp, W.J. & Spearman, M.L. (2000). *Factory Physics*, McGraw-Hill, Boston.
- [6] Montgomery, D.C. (1985). *Introduction to Statistical Quality Control*, John Wiley & Sons, New York.
- [7] Gill, A. (1962). *Introduction to the Theory of Finite-State Machines*, McGraw-Hill.
- [8] Chopard, B. & Droz, M. (eds) (1998). *Cellular Automata Modeling of Physical Systems*, Cambridge University Press.
- [9] Lee, H.K., Barlović, R., Schreckenberg, M. & Kim, D. (2004). Mechanical restriction versus human overreaction triggering congested traffic states, *Physical Review Letters* **92**, 38702.
- [10] Dantzig, G.B. (1963). *Linear Programming and Extensions*, Princeton University Press, Princeton.
- [11] Winston, W.L. (2004). *Operations Research Applications and Algorithms*, Duxbury Press, New York.

ROBERT BARLOVIĆ, ANDRE HOLFELD
AND JAN RÄBIGER

Sampling in Software Development

Introduction

Sampling inspection in software development is more theory than practice. Surveys have shown that the use of quantitative **quality management** methods is scarce in industry [1, 2]. However, the theory is quite well established, and some examples of successful use exist. This article first gives a short historical perspective on quality control in software engineering. Then we elaborate on the use of sampling in software inspection and software testing, and focus on **estimation** of fault content and failure behavior.

Statistical Quality Control in Software Development

In the software sector, the focus of sampling and statistical quality control is on the *development* process, while in areas of physical products, the *production* process is the main focus. The production of software is trivial, i.e., copying some electronic files to a storage medium, while software development involves handling the inherent properties of software, defined by Brooks as, (a) it is more complex than any other construct of its size, (b) it must conform to arbitrary interfaces, (c) it is subject to constant change, and (d) it is invisible [3]. These challenges have led to the need for defining and adding support to the software development process.

The use of **statistical process control** in software development has its roots in the late 1960s. At the NATO conference in Garmisch-Partenkirchen 1968, the term *software engineering* was launched, which was later defined as “the application of a systematic, disciplined, quantifiable approach to the development, operation and maintenance of software; that is, the application of engineering to software” [4].

Winston Royce’s seminal paper, presenting the waterfall development process model [5], was an initial step toward software quality management. Royce also proposed the process being iterative, a concept further developed by Basili [6]. Mills’ Cleanroom method, which was developed at IBM during the 1960s is based on iterative processes

and focuses on the statistical quality control [7], applying inspection and sampling-based testing to assess the software quality. The concepts of iterative development have gradually moved into practice in industry, while the statistical quality control is less widespread, although approaches to adopt statistical process control to software development exist [8] and have been revitalized in the context of agile software development [9].

In this article, we define software *failures* as externally observable misbehavior, encountered during execution. Software *faults* (or defects) are the internal phenomena, statically residing in the code or another artifact. *Error* is a human action resulting in a fault. Errors committed during development may cause a fault in the software code, resulting in a failure during operation [4].

Sampling in Software Inspection

Software inspection is “a static technique that relies on visual examination of development products to detect error, violations of development standards, and other problems. Types include code inspection and design inspection” [4]. The inspected objects can hence be different kinds of documents, specifications as well as code. After an inspection, the number of faults found in the document is known. However, from a management perspective it is more important to know how many faults remain to schedule resources for quality improvement activities and to make release decisions.

Capture–Recapture Technique. The capture–recapture technique was first applied to software inspections by Eick *et al.* [10]. Basically, different reviewers inspect the same document – requirements, design, or code – and report the faults found. On the basis of the overlap between the reviewers’ detected faults, an estimate of the number of remaining faults can be achieved. The simplest model is the Lincoln–Peterson’s estimation model for two reviewers, see equation (1)

$$\hat{N} = \frac{n_1 \times n_2}{m_2} \quad (1)$$

where $\hat{N}$ is the estimate of the total fault population, n_1 and n_2 are the number of faults found by reviewers 1 and 2, respectively, and m_2 is the number of faults found by both reviewers.

2 Sampling in Software Development

Table 1 Capture–recapture models with their respective assumptions and estimators used [2]

Model	Fault probability	Reviewer capability	Estimator
M_0	Equal for all faults	Equal for all reviewers	ML
M_t	Equal for all faults	Different for different reviewers	ML, Chao
M_h	Different for different faults	Equal for all reviewers	Jackknife, Chao
M_{th}	Different for different faults	Different for different reviewers	Chao

There are four principal models, which have different assumptions on the probability for a fault to be found and for the capability of the reviewers to find a fault, see Table 1. The models can be combined with different estimators [2].

Most studies find the M_h -model with jackknife estimator providing most robust estimates, which is natural since it is nonparametric. The research on capture–recapture in software inspection is elaborated in depth by Petersson *et al.* [2].

Sample-Driven Inspection. Sample-driven inspection is an approach to improve the efficiency of the inspection process by concentrating the inspecting effort on those documents that have most faults [11]. This inspection process has four main steps: (a) sampling, (b) preinspection, (c) resources scheduling, and (d) main inspection. A portion of each software document is randomly sampled for preinspection, constituting 10–50% its size, which is a trade-off between the acceptable uncertainty in the prediction and the available resources. The sampled portions are preinspected and the fault content is estimated on the basis of the sample. The documents are then ranked according to the estimate and the remaining inspection effort is spent on the documents with highest ranks according to the fault estimate.

The method is evaluated using **Monte Carlo** simulations by Thelin *et al.* [11], indicating that the method is more efficient than random selection of a subset. Twenty to thirty percent of the document should be preinspection to achieve good resources scheduling. The method is particularly efficient when the quality varies significantly across the documents.

Sampling in Software Testing

Software testing is “an activity in which a system or component is executed under specified conditions, the results are observed or recorded, and an

evaluation is made of some aspect of the system or component” [4]. In contrast with inspection, testing is dynamic, i.e. conducted by executing the software. Hence, only the software code can be verified during testing, and not the requirements and design specifications. Similar to inspection, we want to predict the number of remaining faults after testing. Predicting the operational failure behavior in terms of reliability or mean time between failures is also desirable after testing. A test case is a specific combination of inputs to the program, leading to a predetermined output. To run all combinations is impossible for all but very small programs. Hence, testing is implicit or explicit sampling from the future use of the program.

Different strategies exist for sampling the test cases to execute. Strategies may be functional (also called *black box*, i.e., based only on the externally visible interfaces of a program or component) or structural (also called *white box*, i.e., based on the program structure). The input/output domain or the structure-based domain is split into partitions, which are assumed to have similar properties for a given input. Test cases are sampled from each partition. The test cases may be sampled to cover all the partitions, the most prioritized partitions, or randomly sampled. Which sampling strategy to use is based on its effectiveness in detecting faults, while the ability to estimate faults and predict failure behavior is a secondary goal.

Fault Estimation

To *estimate remaining fault* content, Mills proposed *seeding* additional faults into the code [12]. Faults are seeded before testing, and it is assumed that the sampled test cases reveal the same proportion of original faults as of the seeded faults. The number of remaining faults is estimated using e.g., the Lincoln–Peterson model mentioned above, where

n_1 and n_2 represent found original and seeded faults, respectively. The main problem of this approach is to seed representative faults into the code. The seeding also requires extra effort.

Stringfellow *et al.* [13] applied the capture–recapture techniques to testing at different test sites, testing the same components in parallel. Hence, no extra effort is needed for seeding faults.

Failure Prediction

To predict the operational failure behavior, **software reliability testing** can be used (also called *statistical testing* or *operational profile testing*). The principal idea is that test cases are sampled from a model of the future usage of the software, i.e. selected based on the operational usage frequencies. The failure data from these tests are then employed to fit a software **reliability growth** model [14], thus predicting the operational failure behavior in terms of mean time between failures (MTBF) or reliability based on the sample from testing. The future usage is specified in terms of a usage model, e.g. a list of test cases or features and their corresponding operational usage frequencies.

Different types of usage models are proposed, including domain models [15], Markov models [16], and the operational profile model [17]. Sampling from these models can be conducted with or without replacement [16] and also using stratifying techniques [18]. Various approaches are presented to accelerate the testing without sacrificing the predictive capabilities [19]. However, most reliability models require random sampling from the future usage. Despite decades of research, the state of practice on software reliability testing is rather immature, except in domains of extremely high quality requirements.

The failure data recorded during the test case execution is fit into a software reliability prediction model. There are an extensive number of models proposed in the literature. The two most common approaches are time between failure models and failure count models.

The first *time between failure model* was the Jelinski–Moranda model [20], where it is assumed that the times between failures are independently exponentially distributed. Let X_i denote the time between the $(i - 1)$ th and i th failure, then the **probability density function** of X_i is defined as $f_{X_i}(t) = \lambda_i e^{-\lambda_i t}$. λ_i is assumed to be a function of the

remaining number of failures and is derived as $\lambda_i = \phi(N - (i - 1))$, where ϕ is a constant representing the failure intensity per fault.

Goel and Okumoto proposed a *failure count model* [21] in which $N(t)$ is described by a non-homogeneous **Poisson process**, in which the failure intensity decreases for every fault that is removed from the code. The expected number of faults found at time t is described as $m(t) = N(1 - e^{-bt})$, where N is the total number of faults in the program and b is a constant. The probability function for $N(t)$ is expressed as $P(N(t) = n) = (m(t)^n / n!) e^{-m(t)}$.

A comprehensive overview of software reliability models can be found in [14].

Summary

Sampling inspection can be applied to software development during inspection and testing. The purpose is to find faults and to estimate the number of remaining faults and the operational failure behavior. The theory is quite well established, while the practical application is rare, except in industries with very high quality demands.

Acknowledgments

Thanks to Dr Martin Höst for reviewing an earlier version of this article.

References

- [1] Runeson, P., Andersson, C. & Höst, M. (2003). Test processes in software product evolution – a qualitative survey on the state of practice, *Journal of Software Maintenance and Evolution* **15**(1), 41–59.
- [2] Petersson, H., Thelin, T., Runeson, P. & Wohlin, C. (2004). Capture-recapture in software inspections after 10 years research–theory, evaluation and application, *Journal of Systems and Software* **72**(2), 249–264.
- [3] Brooks, F.P. (1987). No silver bullet – essence and accidents of software engineering, *IEEE Computer* **20**(4), 10–19.
- [4] IEEE Std 610.12 (1990). *IEEE Standard Glossary of Software Engineering Terminology*.
- [5] Royce, W.W. (1970). Managing the development of large software systems, in *Proceedings IEEE WESTCON*, pp. 1–9. Reprinted in *Proceedings of the 9th International Conference on Software Engineering*, Monterey, California, USA, 1987, pp. 328–338.

4 Sampling in Software Development

- [6] Basili, V.R. & Turner, A. (1975). Iterative enhancement: a practical technique for software development, *IEEE Transactions on Software Engineering* **SE-1**(4), 390–396.
- [7] Cobb, R.H. & Mills, H.D. (1990). Engineering software under statistical quality control, *IEEE Software* **7**(6), 45–54.
- [8] Florac, W., Carleton, A.D. & Barnard, J.R. (1999). *Measuring the Software Process. Statistical Process Control for Software Process Improvement*, Addison-Wesley, Reading.
- [9] Andersen, D.J., Agile Management for Software Engineering. (2003). *Applying the Theory of Constraints for Business Results*, Prentice Hall.
- [10] Eick, S.G., Loader, C.R., Long, M.D., Votta, L.G. & Vander Wiel, S. (1992). Estimating software fault content before coding: software engineering, in *Proceedings International Conference on Software Engineering*, Melbourne, Australia, pp. 59–65.
- [11] Thelin, T., Petersson, H., Runeson, P. & Wohlin, C. (2004). Applying sampling to improve software inspections, *Journal of Systems and Software* **73**(2), 257–269.
- [12] Mills, H.D. (1972). On the statistical validation of computer programs, Technical Report FSC 72–6015, IBM. Reprinted in Harlan Mills, *Software Productivity*, Dorset, 1983.
- [13] Stringfellow, C., Andrews, A., Wohlin, C. & Petersson, H. (2002). Estimating the number of components with defects post-release that showed no defects in testing, software testing, *Verification and Reliability* **12**(2), 93–122.
- [14] Lyu, M.R. (1996). *Handbook of Software Reliability Engineering*, McGraw Hill.
- [15] Woit, D.M. (1994). A framework for reliability estimation, in *Proceedings 5th International Symposium on Software Reliability Engineering*, Monterey, California, USA, pp. 18–24.
- [16] Whittaker, J.A. & Thomason, M.G. (1994). A Markov Chain model for statistical software testing, *IEEE Transactions on Software Engineering* **20**(10), 812–824.
- [17] Musa, J.D. (1993). Operational profiles in software reliability engineering, *IEEE Software* **10**(2), 14–32.
- [18] Podgurski, A., Masri, W., McCleese, Y. & Yang, C. (1999). Estimation of software reliability by stratified sampling, *ACM Transactions on Software Engineering and Methodology* **8**(3), 263–283.
- [19] Ehrlich, W.K., Nair, V.N., Alam, M.S., Chen, W.H. & Engel, M. (1998). Software reliability assessment using accelerated testing methods, *Applied Statistics* **47**(1), 15–30.
- [20] Jelinski, Z. & Moranda, P. (1972). Software reliability research, *Statistical Computer Performance Evaluation* W. Freiberger, ed, Academic, New York, 465–484.
- [21] Goel, A. & Okumoto, K. (1979). Time-dependent error-detection rate model for software reliability and other performance measures, *IEEE Transactions on Reliability* **28**(3), 206–211.

PER RUNESON

Sampling in Data Mining

Data Mining and KDD

Knowledge discovery in databases (KDD) is commonly defined as the *nontrivial process of identifying valid, novel, potentially useful, and ultimately understandable patterns in data* [1, p. 6]. The first step in this process is **data collection**. The collected data are stored in (operational) databases, possibly together with background information (metadata). Usually, the data collected in this manner are noisy and inconsistent, so they have to be cleaned, integrated, and transformed. To make the data available for later analysis in such a form that the adequate information is provided in adequate form and detail, the data are organized in a *data warehouse*. A *data warehouse is a subject-oriented, integrated, time-variant, nonvolatile collection of data in support of the management's decision making process* [2, p. 31], i.e., the data are not organized in transactional form but according to major subjects, providing information over a longer period of time. The warehouse integrates data from different data sources in a cleaned and unique format. It is separated from the operational databases, updated only at certain times such that existing data remains unchanged.

The next step in the KDD process consists in analyzing the data stored in the warehouse. If fast queries are desired, then methods of *on-line analytical processing (OLAP)* are the right choice. It is a widely accepted guideline that OLAP tools should deliver results within about 5 s and be able to deal with multidimensional data. To get deeper insight into the data and to detect complex patterns and regularities from vast amounts of data, methods of **data mining** should be applied. The term *data mining* in the strict sense refers to this analytic purpose only, in a more global sense it is also used as a synonym for KDD as a whole. Data mining methods are much more complex than those of OLAP, with the only restriction that they must be able to handle large amounts of data, usually too large to fit in main memory. In the last step of the KDD process, the obtained knowledge is presented to the user. For more details on KDD and data mining see [3].

In this article, we will investigate the topic of *sampling in data mining*. We begin with a general description of typical situations of sampling, and

continue by discussing techniques of sampling that are commonly used in data mining. Finally, we review various fields of data mining where sampling techniques are applied.

Sampling from Databases: Overview

Sampling approaches are applied to various areas of KDD and data mining. The motivation for sampling in data mining, however, is different from that of sampling in traditional statistics. In the latter case, data are often considered expensive and difficult to collect, while it is available in huge amounts in the case of data mining. In fact, it is a common strategy to collect as much data as possible, with the highest possible dimensionality. As a consequence, one often has too much data to be processed by data mining procedures, and sampling methods have to be used to reduce these amounts of data.

The situation sketched above can be related to the general sampling scheme of Deming [4]. Three levels in the sampling process are distinguished: the basic level is the *cause system*, i.e., the population behind the collected data. The next level is the *lot*, the available data from the cause system. In many situations of KDD, it is reasonable to identify the lot with the data warehouse. From this lot, a *sample* is drawn, which will be analyzed with methods of data mining. If the results obtained are used to estimate the properties of the lot, the data are used for an *enumerative* analysis, if it is used to estimate properties of the cause system, Deming [4] speaks of an *analytic* study. The distinction between both types of problems is important, leading to different sampling errors, for instance. The distinction is also known to the KDD community, where *descriptive* and *predictive* data mining are distinguished [3], although it is not clear whether it is always considered in practice.

Sampling in Data Mining: Techniques

The sampling techniques used for data mining are closely related to standard approaches of statistics. Depending on the purpose of sampling – for more details see the following section – techniques like simple random sampling, with or without replacement and possibly under stratification, are used, as well as **sequential sampling** strategies and cluster

sampling. For a description of these techniques see **Sampling Designs in Surveys**. The terminology of these approaches, however, is often different from that of KDD. Depending on the evaluation condition, for instance, sequential sampling is referred to as *adaptive*, *progressive*, or *dynamic* sampling.

Especially for estimating model performance, techniques of data partitioning are commonly used in data mining. The *holdout method* divides the available dataset of n records into a training and a test set, possibly considering stratification. In *k-fold cross-validation*, the initial data are split into $k \leq n$ disjoint subsets $S_1, \dots, S_k$ of approximately equal size. In each iteration $i = 1, \dots, k$, S_i serves as a test set, the union of the remaining sets as training data. The final result is obtained as an average of the k individual results. A special case is the *leave-one-out method*, where $k = n$ [3, 7.9.1].

Sampling in Data Mining: Purpose

In this section, we will briefly review areas of KDD and data mining, where sampling is used: sampling for data reduction, for algorithm recommendation, for initialization of algorithms, for query optimization, for privacy protection, for model validation, and for generating multiple models.

Sampling for Data Reduction

Data, although possibly collected in transactional form, is usually analyzed in tabular form. The columns of a data table represent the *attributes*, *features*, or *variables*, the rows represent the individual *records*, *instances*, *examples*, or *cases*. The number of features is the *dimension* of the data. The need for data reduction before performing analyses can be due to both the dimension of the data and the number of records. Although sampling is usually used to reduce the number of records, it can also be helpful to reduce the dimensionality of the data. Filter techniques of *feature selection* aim at removing irrelevant features from the data, where the decision about relevance is based on given training data. The training data can be obtained by simple random sampling, but other approaches are also known. Liu *et al.* [5], for instance, describe a strategy called *active feature selection*, which applies a selective sampling strategy for building the training sample. The procedure aims

at selecting those records that are most informative in determining the relevance of features.

The main purpose of sampling in data mining concerns the reduction of the number n of records of a dataset, compare [6]. The run time of data mining algorithms usually grows at least linearly in n . In addition, the data will not fit into main memory if n is larger than a certain upper bound. The process of making procedures able to deal with datasets of growing size n is called *scaling up* of algorithms. In some cases, scaling-up can be done by optimizing the computational complexity of the algorithm, or by using distributed algorithms. Another approach to scaling-up is to reduce the amount of data by sampling. This can be done for one specific analysis only, using a specially designed sampling strategy. An example is the windowing technique for decision tree building, where an initial random sample is replenished with misclassified records of a test set. Or one tries to construct a sample, which is representative for the whole database, i.e., a sample, the main characteristics (distribution, **moments**) of which match those of the database. Representativity is commonly checked by comparing marginal properties like univariate means. Inmon [2] even suggests building a reduced data warehouse, called *living sample data*, where “living” expresses that the database needs to be refreshed periodically. For more details on scaling-up algorithms see [6].

Sampling for Algorithm Recommendation

Given a specific analytical task and a set of potential algorithms, *algorithm recommendation* aims to select a subset of algorithms, preferably of size 1, that are expected to work best on the problem. Therefore, the candidate algorithms are commonly evaluated based on a random sample of the data. An example concerning classification methods is described by Schaffer [7], who proposes the use of cross-validation to determine the method with highest average accuracy.

Sampling for Initialization of Algorithms

Sampling can be applied to determine starting values for learning algorithms, to carry them out more efficiently. Toivonen [8], for instance, applied sampling to association rule mining, where a database of itemsets (“market baskets”) is analyzed for frequent

subsets, i.e. subsets which are contained in at least $s_{\min}$ of all available records. Toivonen [8] suggests application of the algorithm to a random sample first, small enough to fit into main memory, so as to find a superset of the collection of all frequent sets. The whole database is then used to determine the exact frequencies of these specific sets only, resulting altogether in an approximate solution of the problem of detecting all frequent sets.

Sampling for Query Optimization

Sampling can be used to optimize the design of queries before applying them to the complete database. On the basis of the selected random sample, the execution time of the query or the size of the result returned by the query can be estimated. For more details see [9]. This is also an important purpose of the living data sample proposed by Inmon [2], which was discussed above.

Sampling for Privacy Protection

Users of a database are often allowed to perform only aggregate queries like COUNT, SUM, and so on, to prevent them from obtaining information about individuals contained in the database. Asking highly specific COUNT queries, designed such that the answer will either be 1 or 0, however, enables individual characteristics to be reconstructed. An approach to provide privacy is described by Olken and Rotem [9]: queries are answered against a data sample only, introducing sampling errors and precluding precise computations.

Sampling for Model Validation

To allow **estimation** of the accuracy of a model to be constructed, techniques of data partitioning are widely used. To build a classification tree, for instance, the data are often partitioned into training and test data (holdout method). Sometimes, a third part of the data is necessary to prune the tree. Alternatively, methods like cross-validation or bootstrapping (see **Sampling from Virtual Populations**) are commonly applied [3, 7.9.1].

Sampling for Generating Multiple Models

A way proposed to increase the accuracy of classifiers is that of constructing $k > 1$ models from the same

data first. Given a new data record, each method leads to a certain classification. The class with the most votes (bagging) or the highest-weighted vote (boosting) is then assigned to the new object. Sampling occurs within the construction of the models, where bootstrapping (see **Sampling from Virtual Populations**) is applied: for each $i = 1, \dots, k$, a training set S_i is sampled with replacement, and the i th model is learned from S_i . In the case of boosting, weights are assigned to each training sample and updated during the model-building phase according to previous misclassifications [3, 7.9.2].

Conclusion

A variety of applications of data mining to quality improvement have been reported in literature: data mining to support the development of a metacontroller for the delivery of high-quality aluminum products, data mining to determine the range of control parameters in wafer production, data mining to obtain process parameters minimizing the production of faulty printed-circuit boards, and many more [10]. The above discussion showed that one essential part of data mining in several situations is sampling. The importance of sampling for quality improvement is thus demonstrated for an area outside traditional statistics also.

References

- [1] Fayyad, U.M., Piatetsky-Shapiro, G. & Smyth, P. (1996). From data mining to knowledge discovery: an overview, in *Advances in Knowledge Discovery and Data Mining*, U. M. Fayyad, G. Piatetsky-Shapiro, P. Smyth & R. Uthurusamy, eds, AAAI/MIT Press, Cambridge, pp. 1–34.
- [2] Inmon, W.H. (2002). *Building the Data Warehouse*, 3rd Edition, John Wiley & Sons.
- [3] Han, J. & Kamber, M. (2001). *Data Mining: Concepts and Techniques*, Morgan Kaufmann Publishers, San Francisco.
- [4] Deming, W.E. (1953). On the distinction between enumerative and analytic surveys, *Journal of the American Statistical Association* **48**, 244–255.
- [5] Liu, H., Motoda, H. & Yu, L. (2002). Feature selection with selective sampling, in *Proceedings of the Nineteenth International Conference on Machine Learning*, Sydney, Australia. pp. 395–402.
- [6] Provost, F. & Kolluri, V. (1999). A survey of methods for scaling up inductive algorithms, *Data Mining and Knowledge Discovery* **2**, 131–169.

4 Sampling in Data Mining

- [7] Schaffer, C. (1993). Selecting a classification method by cross-validation, *Machine Learning* **13**(1), 135–143.
- [8] Toivonen, H. (1996). Sampling large databases for association rules, in *Proceedings of the 22nd Conference on Very Large Data Bases*, Bombay, India. pp. 134–145.
- [9] Olken, F. & Rotem, D. (1995). Random sampling from databases – a survey, *Statistics and Computing* **5**(1), 25–42.
- [10] Kusiak, A. (2006). Data mining: manufacturing and service applications, *International Journal of Production Research* **44**(18–19), 4175–4191.

CHRISTIAN H. WEIB

Sampling in Implementation of Statistical Process Control

Introduction

Quality assurance activities in the manufacturing process consist of proactive activities to build in manufacturing quality at each process step (improve manufacturing process and ensure process capability) and defensive activities to prevent nonconformities from continuing in the production flow. Proactive activities are obviously fundamental to quality assurance. Yet, in spite of proactive activities, nonconformities may eventually occur and have to be prevented from escaping into the process. One means of prevention is 100% inspection; however, this is not to control the process but to determine the acceptability of each unit.

On the other hand, **statistical process control** (SPC) clearly plays a major role in proactive quality assurance activities. Unlike 100% inspection, implementation of SPC is based on sampling to build quality into processes.

The life cycle of SPC (SPC life cycle) to build quality into processes consists of three steps (see, for example, Nishina *et al.* [1]):

1. Ensuring machine performance and then ensuring short-term process capability by pilot production in the preparation stage for mass production.
2. Ensuring long-term process capability in early-stage mass production. Early-stage mass production is the stage when process variation is reduced and the process is brought into the state of statistical control.
3. Maintaining process capability in routine mass production. Routine mass production is the operation and maintenance of process capabilities that were achieved in early-stage mass production.

The above SPC life cycle is a typical approach of parts manufacturing. However, the same approach can be expanded to other fields.

In this chapter, we consider the role of **data collection** by *sampling* in the implementation of SPC over all three life cycle steps. Implementation of SPC obviously varies with the industrial field and the form

of production. Since it is neither possible nor efficient to discuss each of these in detail, attention is restricted to discrete parts manufacturing and semiconductor manufacturing. However, discussion will be held as general as possible. We also provide examples that illustrate problems in applying the concept of inherent process variation, and how to deal with them in general since the inherent process variation can be due to the way items are sampled.

SPC Life Cycle and Sampling

Process Control in Preparation for Mass Production and Sampling

The SPC life cycle starts with ensuring machine performance. Before a production line is formed, it is necessary to confirm that the capability of the individual machines can fully satisfy the product specifications.

An important point here is that, with regard to production process factors, variation is assessed by projecting the situation after the line is formed. It is recommended that an experiment be designed with the factors being the parameters of the assumed production process after the line has been formed, and that variation be assessed from the results (Kuzuya, K., private communication, 2006) – for example, an experimental design in which material conditions are intentionally changed, and variation factors such as conditions before and after exchange of cutting tools, abrasive wheel dressing, or electrodes are incorporated. These experimental factors correspond to noise factors in the Taguchi method. Robustness can be evaluated at the same time in this experiment. At this stage, sampling means collecting data according to the rules of the experimental design, for example, the **orthogonal arrays**.

Implementing the experimental design after a line has been formed is generally very difficult. Conducting an evaluation with consideration of the downstream processes during the upstream stage of SPC will lead to more efficient SPC.

There is also the approach that, to the extent possible, machine performance should be evaluated separately after eliminating the influence of other factors in the production process (for example, ISO DIS 13700 [2]). This method allows evaluation of the contribution of that part of the process variation

2 Sampling in Implementation of Statistical Process Control

caused by machines. However, even if, for example, a large amount of data is obtained under fixed machine conditions using a uniform test piece, the downstream reproducibility of that evaluation is poor. Therefore, the perspectives of achieving good process capability by the ISO approach are not satisfactory.

In the experimental mass production preparation stage, variation is evaluated at the level of individual machines separately. Hence the total process variation of the projected production line is underestimated. One strategy to deal with this effect is to modify the capability calculations. In Japanese industry often the specification range ($S_U - S_L$) is not divided by six times the standard deviation s , but rather

$$P_m = \frac{S_U - S_L}{8s} \quad (1)$$

is used as a machine performance index.

After confirming that the performance of each machine is at a satisfactory level, the production line is formed and the process moves on to pilot production stage. The aim in the pilot production stage is to evaluate and ensure short-term process capability. Since this is a process capability study in short term, the amount of data (about 100 items) is roughly equivalent to a few lots.

The short-term process capability is evaluated from this sample. Whether or not to move to mass production is determined by the short-term process capability evaluation. However, the capability includes little variation due to material and man; therefore, it should be noted that the total process variation of the projected production line is underestimated, e.g., with respect to machine performance.

Process Control in the Mass Production Stage and Sampling

Monitoring and control by Shewhart **control charts** are effective from early-stage mass production through routine mass production. Shewhart control charts can be used for two purposes: (a) to bring a process to a state of statistical control; (b) to judge whether the state of statistical control is maintained. The former role corresponds to the early-stage mass production, and the latter role to the routine mass production.

In either case, Shewhart control charts are set up on the basis of an assessment of the variation

due to chance causes and then they are used to detect the occurrence of assignable causes. Therefore, the identification of chance cause is essential in the development phase of a control chart. Chance causes are related to the inherent variation exhibited by the so-called “5M1E”: Machine, Man, Material, Method, Measurement, and Environment, which are uncontrollable from the technical, economic, and organizational viewpoints.

Shewhart proposed a suitable way to determine the amount of variation due to chance causes. The method is based on relatively small subgroup populations, in which the 5M1E conditions are more or less equal, and the variation due to chance causes from variation within the subgroup is assessed.

Rational subgroups are defined as subgroups *within* which variation is due only to random causes, and *between* which variation is due to special causes. They are selected to enable the detection of any special (assignable) cause of variation among subgroups (by ISO 3534-2 [3]). Ideally such a subgroup is a sample in which all of the items are produced under conditions in which only random effects are responsible for the observed variation (see Nelson [4]). However, in practice it is difficult to specify rational subgroup populations. Generally, subgroup populations are made up in terms of a given day, operator shift or treatment lot, or the like. The essential issue is that rational subgroup populations are made up in terms of a time period, which is related to a physical characteristic as mentioned above.

Subgroup composition and sampling are explained by $\bar{X} - R$ control charts because $\bar{X} - R$ control charts are widely used in comparison with $\bar{X} - s$ and $\bar{X} - s^2$ control chart. Sampling is conducted so that within subgroup variation can be efficiently estimated by the range statistic R . Figure 1 shows an example of sampling in an $\bar{X} - R$ control chart assuming a continuous production process. The example in Figure 1 considers 1 day to be a subgroup population, from which the sample (subgroup) is taken, and the subgroup size to be 4. In $\bar{X} - R$ control charts subgroup size (n) of 4–5 is often used. When, for example, measurement is accompanied by destruction, cases requiring man hours or cost are taken to be $n = 1$. The sampling interval is generally of fixed length (for example, every 2 h) as shown in Figure 1. This is the most practical sampling method based on standardized control methods. Batch sampling should

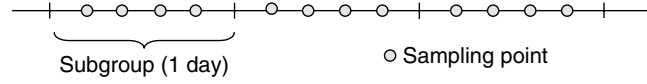


Figure 1 Sampling with $\bar{X} - R$ control charts (example of continuous production process)

be avoided as it underestimates the size of within subgroup variation and produces **bias** in the estimate of process mean.

Other sampling strategies with varying sampling intervals or varying sample sizes have been proposed (see, for example, Reynolds *et al.* [5] and Park *et al.* [6]); however, these approaches have not been yet utilized. They may be applied to monitoring the condition-based maintenance of equipment.

Manufacturing standards are improved in early-stage mass production, and when a sufficiently high process capability is obtained and a state of statistical control is achieved, the process proceeds to routine mass production. At this time, the process capability and the **process capability index** C_p are calculated as

$$\text{Process capability} = 6 \times \frac{\bar{R}}{d_2} \quad (2)$$

$$C_p = \frac{S_U - S_L}{6 \times \bar{R}d_2} \quad (3)$$

from within subgroup variation. Here, $\bar{R}$ is the value of the center line in the R control chart, and d_2 is a coefficient for the unbiased **estimation** of the standard deviation from the range R . In other words, the within subgroup variation on the $\bar{X} - R$ control chart is considered as the inherent process variation, and in routine mass production process, variation is controlled so that it is maintained as the inherent process variation.

In practice, however, Caulcutt [7] and Kuzuya [8] pointed out that the within subgroup variation is not necessarily the inherent process variation. Kuzuya [8] investigated 63 cases of control charts in the stage of routine mass production used for monitoring essential quality characteristics for processes exhibiting a capability index exceeding 1.33. He found that half of the X or $\bar{X}$ charts triggered alarms, thus confirming Caulcutt's results [7]. Nishina *et al.* [1] proposed a measure in which the variation due to chance causes includes some of the variation between subgroups and showed a case study of the measure. It is essential that the process capability in early-stage mass production

is regarded as the standard value of variation due to chance causes in routine mass production. Especially, in batch processing there may be cases where it is not appropriate to view variation from within-batch sampling as variation due to a chance cause.

Sampling to Prevent the Escape of Nonconforming Product

In routine mass production, quality checks by sampling are done to prevent the escape of nonconforming product. At such instances, the quality check is generally determined with the use of a go–no-go gage.

A shift in the 5M1E of production is always followed by a “first item check”. There is a shift in the 5M1E when, for example, a changeover of operator shift is made, a changeover in the material lot is made, or immediately after maintenance.

In addition to the first item check, regular quality checks, in which generally one item is sampled, are conducted at given time intervals. The sampling interval of the regular quality checks is determined with consideration of traceability when a quality problem is discovered (in a parts plant, for example, the time up until the parts are shipped in the truck yard). Some plants use acoustic signals to indicate regular quality checks at periodic intervals, e.g., at 2-h intervals. In machining processes, there are also cases when quality checks are done after a fixed number of pieces, based on the consideration of cutting tool life.

Conditions for this quality check system are (a) the development of a flow production system infrastructure guaranteeing traceability of products; (b) being in routine mass production, in which process capability is high (for example, process capability index $C_p \geq 1.33$); and (c) that the process is in a state of statistical control. This means that, if there are no quality problems in a successive regular quality check, all products in the check interval are assured to be in conformity; in additions, in the event that a nonconforming item is discovered, one allows to trace all products corresponding to the immediately

4 Sampling in Implementation of Statistical Process Control

preceding regular quality check interval. However, it should be kept in mind that sampling to prevent escape of nonconformities into the production flow is not a control activity to lead directly to process improvement.

Practical Cases with Respect to Inherent Process Variation and Sampling

Sampling is closely related to the formation of inherent process variation. This section shows two examples that illustrate how inherent process variation is formed.

Inherent Process Variation in Attribute Control Charts and Sampling

A description of inherent process variation and sampling error will be given for c control charts, considering a semiconductor manufacturing process. In semiconductor manufacturing, particles need to be controlled with respect to various factors. The quality characteristic to be controlled is the number of particles on a single wafer taken from a processed batch. The sampling method is illustrated by Figure 2. In this case, the wafers are considered as items and the batches are subgroup populations. One wafer in a batch is taken as sample since it is known that the variation between the wafers in a batch is relatively small.

SPC literature shows that c control charts are applied to this case. The control limit of the c control chart is

$$\bar{c} \pm 3\sqrt{\bar{c}}$$

Here, $\bar{c}$ is the mean value for the number of particles. c control charts are established on the assumption that the quality characteristic follows a Poisson distribution. This assumption deserves some caution: it is valid only under the condition that the probability of particle attachment is uniform on the wafer and that the number of positions on the

wafer that potentially carry a particle is large. Under these assumptions, for the sampling method shown in Figure 2, only the sampling error is considered as the inherent process variation on c control charts with $\bar{c} \pm 3\sqrt{\bar{c}}$ control limits. That is, the elements in the 5M1E factors of the production process are not included in the inherent process variation.

However, the inherent process variation may include variation due to sources other than the sampling error in the case of the sampling shown in Figure 2. Figure 3 shows a c control chart for particle control in a process in the semiconductor wafer process. Judging from the usually observed data, it is difficult to say that the data in Figure 3 indicate an abnormal state. However, many points are beyond the control limits, and this is not a suitable basis for control charting. The behavior in Figure 3 should be considered as variation due to chance cause, and shows that the variation on the wafer is not sampling error only. The particle attachment rate may be not uniform on the wafer, and the variation due to chance cause needs to be obtained with consideration of the variation in the particle attachment rate on the wafer. One measure that has been proposed for this situation is to assume a negative binomial distribution rather than a Poisson distribution for the particle number distribution [9]. Figure 4 shows a control chart using control limits that assume a negative binomial distribution. The control chart in Figure 4 shows a state of statistical control, and illustrates that it is valid to assume a negative binomial distribution for the particle number.

This example illustrates the relation between sampling error in count data and inherent process variation.

Inherent Process Variation Related to Pattern on Two Dimensions

This section explains the inherent process variation related to patterns in the chemical mechanical polishing (CMP) process, a semiconductor wafer process. The CMP process is a process to polish and plane

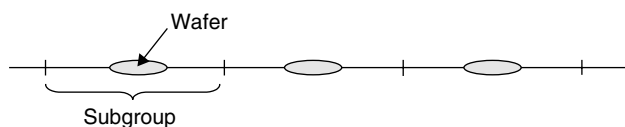


Figure 2 Sampling with an attribute control chart

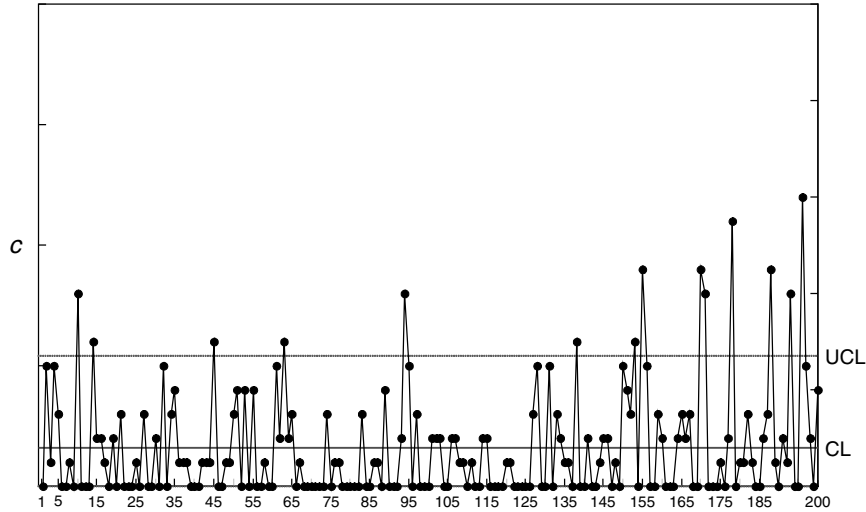


Figure 3 P article number control chart (conventional c control chart)

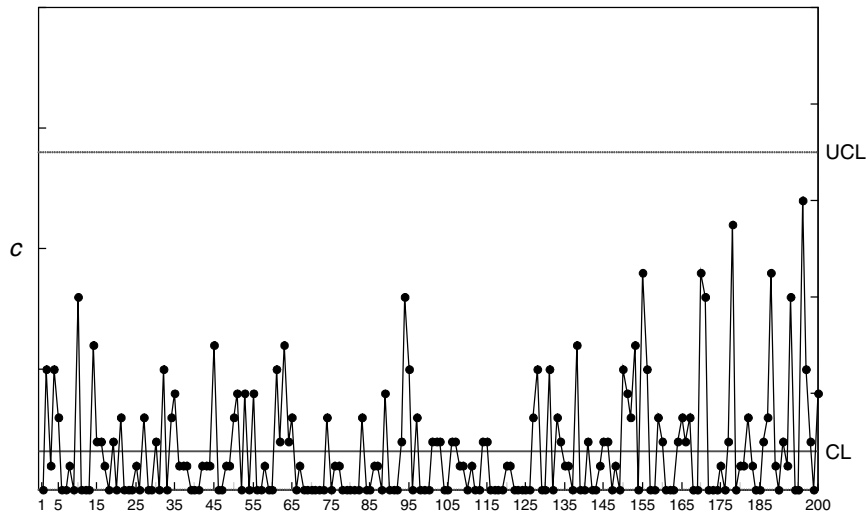


Figure 4 Control chart assuming a negative binomial distribution for particle number

wafer surfaces, excluding irregularities in the film that forms the wafer surface. The quality characteristic of this process is the thickness of the film that remains after polishing.

In the CMP process, one lot consists of 25 wafers. Polishing is done for each wafer unit, but the process parameters are not changed within a lot. Process variation can be broken down into variation between lots, variation within a lot, and variation on a wafer. This process has the following characteristics.

1. Variation on a single wafer is large, and has a pattern.
2. Variation within a lot is small.
3. Variation between lots is large (shows a trend according to abrasive wheel friction).

In this process, therefore, between-lot trends and within-wafer variation must be controlled.

The remaining film in this process is measured on the nine points on the wafer shown in Figure 5.

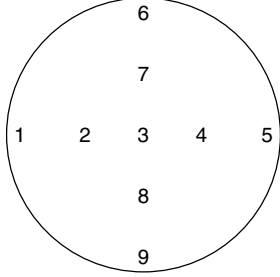


Figure 5 Measurement points of remaining film thickness in a wafer (numerals stand for measurement points)

The between-lot variation and within-wafer variation are controlled using the data from these nine points. Since the within-lot variation is small, it is sufficient to sample one wafer in a lot.

The amount of polishing is adjusted by the polishing time. Between-lot variation can be handled by feedback control as the modulator for polishing time. Automatic process control is often used as shown for this handling in the semiconductor process.

In processes such as the semiconductor production process, equipment maintenance is a key control activity in routine mass production. Since maintenance costs are high, the maintenance period is extended as long as possible. Higashide *et al.* [10] proposed statistics to indicate the maintenance period from data measured on the nine points in Figure 5.

In the target process, there is an inherent process pattern in the within-wafer variation, and there was a change in that pattern before and after maintenance. Higashide *et al.* used principal component analysis for the matrix of the sum

of products following a double centralized transformation, $y_{i \times j}$, shown in equation (4) to identify inherent process pattern in the within-wafer variation.

$$y_{i \times j} = y_{ij} - \bar{y}_{i \cdot} - \bar{y}_{\cdot j} + \bar{\bar{y}} \quad (4)$$

Here, y_{ij} is the remaining film thickness at position j on the i th wafer. $\bar{y}_{i \cdot}$ and $\bar{y}_{\cdot j}$ indicate the row and column averages of the data matrix (sample $\times$ measurement position), respectively, and $\bar{\bar{y}}$ stands for the total average of the data matrix.

As a result of principal component analysis, the contribution rate of the first principal component is about 50%, and the first principal component score may be considered to show the majority of the remaining film thickness pattern. The first principal component score is shown in a **time series** in Figure 6. From Figure 6, the first principal component score shows a downward trend before maintenance. After maintenance a big swing in the positive direction is observed. If a weighting factor to calculate the first principal component score is determined from the early-stage mass production data, the timing of the maintenance period in routine mass production and the maintenance effect can be monitored.

Conclusion

The SPC life cycle consists of the stages of mass production preparation, early-stage mass production, and routine mass production. The article gives a practical description of the kind of data that should be obtained by sampling for the objective of reducing process variation at each cycle step.

Fundamentally, in the mass production preparation stage, data should be acquired using experimental design methods. In early-stage mass production, data

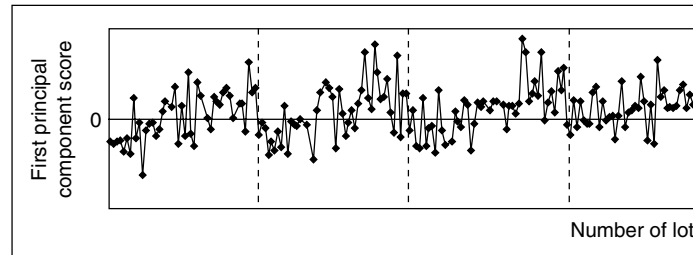


Figure 6 Time series of first principal component scores (the dotted lines stand for maintenance time)

should be acquired through on-line monitoring that will help to achieve process improvements and a state of statistical control. In routine mass production, data obtained should be such that it will help in operation and maintenance of processes such as equipment maintenance, using a standardized monitoring method. In addition, sampling is needed for quality checks to prevent the escape of nonconforming products in the mass production stage.

Sampling is closely related to the assessment of inherent process variation. It is essential to grasp what is considered as inherent process variation.

Acknowledgments

The author would like to thank Mr. K. Kuzuya, formerly of Denso Corporation and Mr. M. Higashide, NEC Electronics Corporation, for their valuable comments, and Prof. R. Goeb and Prof. E. V. Collani of the University of Wurzburg for their written recommendation.

References

- [1] Nishina, K., Kuzuya, K. & Ishii, N. (2000). Reconsideration of control charts in Japan, *Frontiers in Statistical Quality Control* **8**, 136–150.
- [2] ISO DIS 13700 (2006). Machine performance studies – Measured data – discrete parts.
- [3] ISO 3534-2 (2006). *Vocabulary and symbols – Part 2: applied statistics*.
- [4] Nelson, L.S. (1988). Rational subgroups and effective applications, *Journal of Quality Technology* **20**(1), 73–75.
- [5] Reynolds, M.R., Amin, M.W., Arnold, J.C. & Nachlas, J.A. (1988). $\bar{X}$ charts with variable sampling intervals. *Technometrics* **30**(1), 181–192.
- [6] Park, C. & Reynolds, M.R. (1999). Economic design of a variable sampling rate $\bar{X}$ -chart. *Journal of Quality Technology* **31**(4), 427–443.
- [7] Caulcutt, R. (1995). The rights and wrongs of control charts, *Applied Statistics* **44**(3), 279–288.
- [8] Kuzuya, K. (2000). Control charts are renewed by adaptation of new JIS, in *Proceeding of the 65th Conference of Japanese Society of Quality Control (in Japanese)*, Nagoya, Japan, pp. 72–75.
- [9] Friedman, D.J. (1993). Some considerations in the use of quality control techniques in integrated circuit fabrication, *International Statistical Review* **61**, 97–107.
- [10] Higashide, M., Nishina, K., Kawamura, H. & Ishii, N. (2007). Practice of statistical process control in semiconductor manufacturing process, in *Proceedings of the 9th International Workshop on ISQC'07*.

KEN NISHINA

Sampling from Virtual Populations

Introduction

As the complexity of problems and methods for performing statistical inference in quality and reliability increases, the availability of simple solutions becomes rare. For example, the computation of the **average run length** for a **control chart** may be simple when the process is normally distributed. However, the average run length under other distributions, or when the process distribution is unknown, may be difficult, if not impossible, to compute in a closed form. In such cases simulation methods such as the **bootstrap** can be used to estimate the average run length of the control chart. Another example occurs when using process-capability indices to evaluate the capability of a process. Even when the process distribution is normal, many of the standard process capability indices in use today do not have simple closed-form procedures for computing **confidence intervals** or performing hypothesis tests. Asymptotic approximations have been developed for many of the problems discussed above. However, these approximations may not provide the accuracy required for industrial use. In such cases methods based on computer simulations and the bootstrap can often be used to obtain approximate confidence intervals and tests.

Let $\mathcal{X} = \{X_1, \dots, X_n\}$ be a set of independent and identically distributed random variables from a distribution $F \in \mathcal{F}$, where $\mathcal{F}$ is a collection of distribution functions (see **Cumulative Distribution Function (CDF)**). This relationship will be denoted as $\mathcal{X} \sim F$. Let θ be a functional parameter of F . That is, $\theta = T(F)$ for some $T : \mathcal{F} \rightarrow \mathbb{R}$. Let $\hat{\theta}$ be an estimator of θ (see **Estimation**). Of interest is performing statistical inference to learn about the true value of θ based on the sample $\mathcal{X}$. Such methods generally require the sampling distribution of $\hat{\theta}$ or of a related function $R_n(\theta, \mathcal{X})$. Often the function $R_n(\theta, \mathcal{X})$ is a test statistic for hypothesis testing or a pivotal quantity for generating confidence regions (see **Confidence Intervals; Hypothesis Testing**). Define $H_n(R, t) = P[R_n(\theta, \mathcal{X}) \leq t | \mathcal{X} \sim F]$ as the sampling distribution of interest. Let $h_\xi = H_n^{-1}(R, \xi) = \inf\{t : H_n(R, t) \geq$

$\xi\}$ be the ξ th quantile of $H_n(R, t)$ for any $\xi \in (0; 1)$ (see **Quantiles**).

As an example, suppose $\mathcal{X} \sim F \in \mathcal{F}$, where $\mathcal{F}$ is the collection of all normal distributions with mean θ and variance $\sigma^2 < \infty$ (see **Normal Distribution**). Let $R_n(\theta, \mathcal{X}) = n^{1/2}(\hat{\theta} - \theta)/\hat{\sigma}$ where $\hat{\theta}$ and $\hat{\sigma}$ are the usual sample mean and standard deviation computed on the sample $\mathcal{X}$. The $R_n(\theta_0, \mathcal{X})$ is the well-known t statistic for testing the null hypothesis $H_0 : \theta = \theta_0$. This particular choice of $R_n(\theta, \mathcal{X})$ can also be used to derive a confidence interval for θ .

There are several basic situations that can occur in statistical inference:

1. If F is known then $H_n(R, t)$ and h_ξ will be known, though they may not have a closed form. When $H_n(R, t)$ does not have a closed form, empirical methods based on simulations can be used to approximate $H_n(R, t)$ to any specified degree of accuracy.
2. If F is unknown, the structure of $R_n(\theta, \mathcal{X})$ and $\mathcal{F}$ may allow the form of $H_n(R, t)$ to be known, or be approximated using simulation.
3. If F and $H_n(R, t)$ are both unknown, then $H_n(R, t)$ must be estimated, usually using the bootstrap.

This article will address the issue of the sampling mechanism used for each of these cases when simulation must be used to approximate, and perhaps estimate, the sampling distribution $H_n(R, t)$. In each of these cases samples are generated from a *virtual population*, one which is used as a proxy for the actual population from which the original data were sampled. This article will also address the issue of sampling from virtual populations that do not fit within the framework of independent and identically distributed sampling developed above, specifically regression and **time series** applications.

Virtual Populations

In each case studied in this section the virtual population used with the sampling mechanism in simulations is chosen so that the virtual population and sampling mechanism mimic the true population and sampling mechanism as close as possible. Several cases are studied below.

2 Sampling from Virtual Populations

F Known

Consider the case where F is completely known but $H_n(R, t)$, and hence h_ξ does not have a closed form. In this case the distribution $H_n(R, t)$ can be approximated empirically as follows.

1. Generate b samples of size n from F . That is, sample $\mathcal{X}_i^* \sim F$ for $i = 1, \dots, b$. Each sample, and the components of the sample should be mutually independent.
2. Calculate $R_i^* = R_n(\theta, \mathcal{X}_i^*)$ for $i = 1, \dots, b$.
3. Approximate the distribution function $H_n(R, t)$ with the empirical distribution function computed on $R_1^*, \dots, R_b^*$,

$$H_n(R, t) \simeq \tilde{H}_n(R, t) = b^{-1} \sum_{i=1}^b \Omega(R_i^* \leq t) \quad (1)$$

where $\Omega(B)$ is the indicator function of a set B .

This method of approximating $H_n(R, t)$ is not new. Gossett implemented this method in his investigations of what became the t distribution. These results were published in 1908 [1].

To generate the samples required in step 1, usually uniform (pseudo) random numbers are generated on a computer, using a random number table, or by some other electronic or physical device. Observations following the distribution F are then generated by transforming the uniform observations through F^{-1} . Other methods, such as acceptance–rejection methods, can be used as well. For example in [1] the device used to generate the samples was a well-shuffled deck of cards which had the values in the population written on them. Many software packages, such as The R Computing Environment which is freely available for most platforms at www.r-project.org, also offer convenient methods for generating random observations from most of the common statistical distributions. An overview of these methods can be found in [2].

The empirical distribution function used in step 3 provides a convenient nonparametric method for approximating $H_n(R, t)$ based on the simulated samples. With respect to the sampling frame within the simulation when n and F are fixed, $\tilde{H}_n(R, t)$ is pointwise unbiased and pointwise consistent as $b \rightarrow \infty$, and has the property

$$\lim_{b \rightarrow \infty} \sup_{t \in \mathbb{R}} |\tilde{H}_n(R, t) - H_n(R, t)| = 0 \quad (2)$$

which is implied by the Glivenko–Cantelli theorem [3]. This property implies that the approximation $\tilde{H}_n(R, t)$ can be made arbitrarily accurate by taking b to be large enough. The choice of b is crucial to the accuracy of the approximation. No simple rules exist for choosing b due to the variety of problems that can be solved by this method. For example, if the tails of $H_n(R, t)$ are of interest then larger values of b would be required compared to the case where the mean of $H_n(R, t)$ is of interest. An overview of some basic methods for determining how long a simulation should run can be found in [4].

The percentile h_ξ can be approximated as

$$h_\xi \simeq \tilde{h}_\xi = \inf_{t \in \mathbb{R}} \{t : \tilde{H}_n(R, t) \geq \xi\} = R_{(i)}^* \quad (3)$$

when $i/b < \xi \leq (i+1)/b$ for $i = 1, \dots, b-1$, where $R_{(i)}^*$ denotes the i th largest value in $R_1^*, \dots, R_b^*$. Other versions of the sample quantile computed on $R_1^*, \dots, R_b^*$ can be used as well.

F Unknown, R Distribution Free

When F is unknown it is helpful to determine whether $R_n(\theta, \mathcal{X})$ is *distribution free* over the collection $\mathcal{F}$, which implies that $R_n(\theta, \mathcal{X})$ has the same distribution regardless of which $F \in \mathcal{F}$ in the parent population. In this case the algorithm described above can be implemented without any changes where F can be taken to be *any* member of $\mathcal{F}$.

Returning to the earlier example, suppose $\mathcal{X} \sim F \in \mathcal{F}$, where $\mathcal{F}$ is the collection of all normal distributions and let $R_n(\theta, \mathcal{X}) = n^{1/2}(\hat{\theta} - \theta)/\hat{\sigma}$. It is well known that $R_n(\theta, \mathcal{X})$ follows a t distribution with $n-1$ **degrees of freedom**, regardless of which member of $\mathcal{F}$ the population follows (see **Probability Density Functions**). This implies that $R_n(\theta, \mathcal{X})$ is distribution free on the collection of normal distributions. This means that if the form of the t distribution were unknown, but it is still known that $R_n(\theta, \mathcal{X})$ is distribution free, one could simulate samples from *any normal distribution* to approximate the form of the t distribution.

F Unknown, R Not Distribution Free

The case of broadest application occurs when F is an unknown member of $\mathcal{F}$ and either it is not known whether $R_n(\theta, \mathcal{X})$ is distribution free over $\mathcal{F}$, or it is known that $R_n(\theta, \mathcal{X})$ is not distribution free over

$\mathcal{F}$. In this case, it is not possible to find $H_n(R, t)$ exactly, but it can be estimated using the *bootstrap* developed by [5]. The bootstrap estimate of $H_n(R, t)$ is computed by finding the distribution of $R_n(\theta, \mathcal{X})$ under the assumption that the population F is really equal to an estimate of F , conditional of the observed sample $\mathcal{X}$. That is, the bootstrap estimate of $H_n(R, t)$ is given by Efron

$$\hat{H}_n(R, t) = P^*[R_n(\hat{\theta}, \mathcal{X}^*) \leq t | \mathcal{X}^* \sim \hat{F}] \quad (4)$$

where $P^*(A)$ represents the conditional probability measure $P(A|\mathcal{X})$ and $\hat{F}$ is an estimate of F . Note that the calculation of the bootstrap estimate in equation (4) does not necessarily imply the need for a simulation to carry out the required calculations. In fact [5] presents an example based on the binomial distribution where the bootstrap estimate in equation (4) can be computed explicitly. For other examples where functions of $\hat{H}_n(R, t)$ can be computed explicitly, see [6, 7]. In other cases, analytic approximations such as the saddle point approximation and the delta method can also be used. See [5], [8–11] for examples of these types of applications. The normal approximation $H_n(R, t) \simeq \Phi(t)$ is another type of approximation usually applicable when the central limit theorem applies to the statistic $H_n(R, t)$. However, in many cases a simulation approximation such as the one described above is implemented to approximate the bootstrap estimate as it rarely exists in closed form and either many of the approximations require complex computations or the approximations are not accurate enough for the problem being studied. The only change required in the algorithm described above is that the sampling in step 1 should be from $\hat{F}$, conditional on $\mathcal{X}$, instead of F itself.

The most common estimator used for F is the empirical distribution function defined in equation (1) computed on $\mathcal{X}$. Note that conditional on $\mathcal{X}$, the empirical distribution function computed on $\mathcal{X}$ corresponds to a discrete mass function that places a mass of n^{-1} on each of the observed sample points in $\mathcal{X}$. Therefore, generating samples from the empirical distribution function of $\mathcal{X}$ is equivalent to sampling n times, with replacement, from the observations in the original sample $\mathcal{X}$. This type of sampling is usually called **resampling**. The key to the success of the application of the bootstrap in many problems is that under specific conditions on F and $R_n(\theta, \mathcal{X})$ the

bootstrap provides a consistent estimate of $H_n(R, t)$. For a summary of the major results see Chapter 3 of [12]. In addition, when the parameter of interest θ is a smooth function of a mean vector, it can be shown that the bootstrap provides an asymptotically more accurate estimate of $H_n(R, t)$ than is given by the central limit theorem.

More efficient methods, that is methods that require smaller values of b to obtain the same accuracy, for performing resampling have been developed using algorithms from the statistical computing literature. Examples of these algorithms include balanced resampling [13], centering [14], and antithetic resampling [15].

In some cases the use of the discrete empirical distribution as the estimator of F , and hence the virtual sampling population, becomes problematic. This is particularly true when the parameter θ is not a smooth function of F , for example when θ is a percentile of F such as the median. See [16] and [17]. The usual alternative estimate used for F in this case is a smoothed estimate such as a kernel distribution function estimate, or equivalently, a kernel density estimate. See [18] for a detailed theory of these estimates. Other methods for computing smooth estimates of F have also been developed (see **Density and Failure Rate Estimation**). See [19] and [20] for an exposition of some of these estimates.

Virtual Error Populations

Sampling from virtual populations is also used for data that are not independent and identically distributed. Consider the standard general linear model where an n -dimensional random vector $\mathbf{Y}$ is observed conditional on an observed and fixed $n \times p$ matrix $\mathbf{X}$ with the property that $E(\mathbf{Y}|\mathbf{X}) = \mathbf{X}\boldsymbol{\theta}$ where $\boldsymbol{\theta}$ is a p -dimensional vector of unknown constants. The usual model used to generate this type of data is of the form $\mathbf{Y} = \mathbf{X}\boldsymbol{\theta} + \boldsymbol{\epsilon}$ where $\boldsymbol{\epsilon}$ is an n -dimensional random error vector with mean vector $E(\boldsymbol{\epsilon}) = \mathbf{0}$ and **covariance** matrix $V(\boldsymbol{\epsilon}) = \sigma^2\mathbf{I}$, where $\mathbf{0}$ is an n -dimensional vector whose entries are all 0 and $\mathbf{I}$ is the identity matrix. It is usually assumed in this case that the components of $\boldsymbol{\epsilon}$ are mutually independent. Using this structure, it will be assumed that the components of the error vector can be represented as a set of independent and identically distributed random variables from a distribution F , called the *error distribution*.

Note that while the components of the error vector are independent and identically distributed, the components of $\mathbf{Y}$ are not identically distributed except in special cases. Many common statistical techniques fall within the framework of the linear model. These include linear regression analysis and analysis of variance techniques.

The parameter vector θ is usually estimated by minimizing $\|\mathbf{Y} - \mathbf{X}\theta\|_q$ where $\|\mathbf{w}\|_q$ is the L_q norm. Let $\hat{\theta}$ be an estimator of θ , $R_n(\theta, \mathbf{X}, \mathbf{Y})$ be a function of θ , $\mathbf{X}$, and $\mathbf{Y}$, and $H_n(R, t) = P[R_n(\theta, \mathbf{X}, \mathbf{Y}) \leq t | \epsilon_1, \dots, \epsilon_n \sim F]$. Inequalities between vectors will be interpreted elementwise. When $\mathbf{X}$ is full rank then the estimator exists uniquely and the model is usually called a *regression model*. When $\mathbf{X}$ is not full rank then $\hat{\theta}$ is not unique and only estimable functions of θ can be uniquely estimated. It will be assumed that if $R_n(\theta, \mathbf{X}, \mathbf{Y})$ is a function of $\hat{\theta}$ in this case, then it is an estimable function.

As an example consider the case of linear regression fitted by **least squares** ($q = 2$). Assuming that $\mathbf{X}$ has full rank and that F is normal, it follows that the function

$$R_n(\theta, \mathbf{X}, \mathbf{Y}) = \frac{(n-p)^{1/2} \mathbf{e}_j' [(\mathbf{X}'\mathbf{X})^{-1} \mathbf{X}'\mathbf{Y} - \theta]}{[(\mathbf{Y}'\mathbf{Y} - \hat{\theta}'\mathbf{X}'\mathbf{Y}) \mathbf{e}_j' (\mathbf{X}'\mathbf{X})^{-1} \mathbf{e}_j]^{1/2}} \quad (5)$$

follows a t distribution with $n-p$ degrees of freedom where $\mathbf{e}_j$ is a vector whose j th element is equal to 1 and the remaining elements are equal to 0. This function can be used to derive confidence intervals and tests for the j th element of θ .

In the case of least-squares fitting ($q = 2$) with a normal error distribution, there are many analytical results about the distribution of useful functions $R_n(\theta, \mathbf{X}, \mathbf{Y})$ (see **Least-Squares Estimation**). When the error distribution is not normal or it is unknown, or the least-squares criteria are not used for estimating θ , fewer analytical results are known and simulation methods and approximations become more important. The most common approximation in this case is given by the central limit theorem for regression models [21]. When analytic results are not available and approximations are either not available or not practical, the simulation method presented in the previous section can be adapted to these types of models as well.

To perform simulations for regression models, new $\mathbf{Y}$ vectors need to be generated conditional on $\mathbf{X}$ when

θ is unknown. The most common solution to this problem begins by obtaining the estimate $\hat{\theta}$ from $\mathbf{X}$ and $\mathbf{Y}$. If F is known, or $R_n(\theta, \mathbf{X}, \mathbf{Y})$ is distribution free over $\mathcal{F}$ and $F \in \mathcal{F}$, then new $\mathbf{Y}$ vectors can be generated by first simulating a sample of size n from the error distribution F . Call this generated vector $\mathbf{r}^*$. The new $\mathbf{Y}$ vector is then generated as $\mathbf{Y}^* = \mathbf{X}\hat{\theta} + \mathbf{r}^*$. This process is repeated b times and $R_n(\theta, \mathbf{X}, \mathbf{Y}^*)$ is computed for each generated $\mathbf{Y}^*$ vector. The distribution of $R_n(\theta, \mathbf{X}, \mathbf{Y}^*)$ can then be approximated with the empirical distribution of the b simulated values of $R_n(\theta, \mathbf{X}, \mathbf{Y}^*)$.

If F is unknown and there is no distribution-free result, the residual vector $\mathbf{r}$ is computed as $\mathbf{r} = \mathbf{Y} - \mathbf{X}\hat{\theta}$ (see **Residuals**). The vector $\mathbf{r}^*$ is then generated by resampling from $\mathbf{r}$, or equivalently sampling from the empirical distribution of the elements of the residual vector. These methods have been studied theoretically and empirically by [16, 22–26].

It has been noted that the **residuals** are not independent of one another while the components of the error vector are. To account for this fact, Efron [27] suggests transforming the residuals before resampling. Another approach is to use external bootstrap resampling. The external bootstrap resamples from a continuous distribution with certain characteristics. See [26] and [28]. When $\mathbf{X}$ is stochastic, the paired resampling algorithm has also been suggested. See Chapter 7 of [12].

The sampling distributions of functions of parameter estimates from time series models can also be approximated using similar methods to those used in linear models (see **Time Series Analysis**). Methods based on resampling residuals from fitted time series models have been proposed. See Chapter 9 of [12] and [29–31] for further details and alternative methods.

References

- [1] Student (1908). The probable error of a mean, *Biometrika* **6**, 1–25.
- [2] Gentle, J.E. (1998). *Random Number Generation and Monte Carlo Methods*, Springer, New York.
- [3] Seffling, R.J. (1980). *Approximation Theorems of Mathematical Statistics*, John Wiley & Sons, New York, pp 56–65.
- [4] Bratley, P., Fox, B.L. & Schrage, L.E. (1987). *A Guide to Simulation*, 2nd Edition, Springer-Verlag, New York, pp 44–123.

-
- [5] Efron, B. (1979). Bootstrap methods: another look at the jackknife, *The Annals of Statistics* **7**, 1–26.
 - [6] Huang, J.S. (1991). Efficiency computation of the performance of bootstrap and jackknife estimators of the variance of L -statistics, *Journal of Statistical Computation and Simulation* **38**, 45–56.
 - [7] Ernst, M.D. & Hutson, A.D. (2000). The exact bootstrap mean and variance of an L -estimator, *Journal of the Royal Statistical Society. Series B* **62**, 89–94.
 - [8] Hinkley, D.V. & Wang, S. (1988). Discussion of “Saddle point methods and statistical inference” by N. Reid, *Statistical Science* **3**, 232–233.
 - [9] Davison, A.C. & Hinkley, D.V. (1988). Saddle point approximations in resampling methods, *Biometrika* **75**, 417–431.
 - [10] Daniels, H.E. & Young, G.A. (1991). Saddle point approximation for studentized mean, with applications to the bootstrap, *Biometrika* **78**, 169–179.
 - [11] Wang, S. (1992). General saddle point approximations in the bootstrap, *Statistics and Probability Letters* **13**, 61–66.
 - [12] Shao, J. & Tu, D. (1995). *The Jackknife and Bootstrap*, Springer, New York.
 - [13] Davison, A.C., Hinkley, D.V. & Schechtman, E. (1986). Efficient bootstrap simulations, *Biometrika* **73**, 555–566.
 - [14] Efron, B. (1990). More efficient bootstrap computations, *Journal of the American Statistical Association* **82**, 171–200.
 - [15] Hall, P. (1989). Antithetic resampling for the bootstrap, *Biometrika* **76**, 713–724.
 - [16] Hall, P. (1992). *The Bootstrap and Edgeworth Expansion*, Springer, New York.
 - [17] Hall, P. & Martin, M.A. (1991). One the error incurred using the bootstrap variance estimate when constructing confidence intervals for quantiles, *Journal of Multivariate Analysis* **38**, 70–81.
 - [18] Wand, M.P. & Jones, M.C. (1995). *Kernel Smoothing*, Chapman & Hall, London.
 - [19] Prakasa Rao, B.L.S. (1983). *Nonparametric Functional Estimation*, Academic Press, New York.
 - [20] Simonoff, J.S. (1996). *Smoothing Methods in Statistics*, Springer, New York.
 - [21] Eicker, F. (1963). Asymptotic normality and consistency of least squares estimators for families of linear regressions, *The Annals of Mathematical Statistics* **34**, 447–456.
 - [22] Adkins, L.C. & Hill, R.C. (1990). An improved confidence ellipsoid for linear regression models, *Journal of Statistical Computation and Simulation* **36**, 9–18.
 - [23] Freedman, D.A. (1981). Bootstrapping regression models, *The Annals of Statistics* **9**, 1218–1228.
 - [24] Freedman, D.A. & Peters, S.C. (1984). Bootstrapping a regression equation: some empirical results, *Journal of the American Statistical Association* **79**, 97–106.
 - [25] Hall, P. (1989). Unusual properties of bootstrap confidence intervals in regression problems, *Probability Theory and Related Fields* **81**, 247–273.
 - [26] Wu, C.F.J. (1986). Jackknife, bootstrap, and other resampling methods in regression analysis (with discussions), *The Annals of Statistics* **14**, 1261–1350.
 - [27] Efron, B. (1982). *The Jackknife, the Bootstrap, and Other Resampling Plans*, SIAM, Philadelphia.
 - [28] Liu, R.Y. (1988). Bootstrap procedures under some non-i.i.d. models, *The Annals of Statistics* **16**, 1697–1708.
 - [29] Efron, B. & Tibshirani, R.J. (1986). Bootstrap methods for standard errors, confidence intervals, and other measures of statistical accuracy, *Statistical Science* **1**, 54–77.
 - [30] Holbert, D. & Son, M.S. (1986). Bootstrapping a time series model: some empirical results, *Communications in Statistics, Series A* **15**, 3669–3691.
 - [31] Bose, A. (1988). Edgeworth correction by bootstrap in autoregressions, *The Annals of Statistics* **16**, 1709–1722.

Related Articles

Bootstrap and Jackknife, Overview; Monte Carlo Methods; Monte Carlo Methods, Univariate and Multivariate; Nonparametric Tests; Resampling for Lifetime Distributions.

ALAN M. POLANSKY

Internal and External Quality Measures

Use of Surveys to Improve Quality

There are two distinct general approaches to the use of statistical methods in quality improvement. These are the *production approach* and the *marketing approach*.

The *production approach* focuses on quality improvement from the management (internal) viewpoint. Quality standards in this framework can be expressed either in nonmonetary or monetary terms and applied in observation as well as inferential studies. The observations studies build on the application of the basic statistical **quality control** toolbox, which includes check sheets, **Pareto** charts, histograms, scatterplots, **control charts**, flowcharts, and other simple tools [1]. The inferential studies build on the application of the design of experiments (**DOE**) with the goal of preventing flaws in products and services due to a poor design [2].

The *marketing approach* focuses on the customer (external) assessment of quality using tools of survey research. Emphasis is on the analysis of quality attributes and analysis of overall quality. Like in the production approach, quality standards in the framework of the marketing approach can be expressed in nonmonetary or monetary terms. When defining quality standards in nonmonetary terms, there is little room to avoid the highly subjective *customer expectations* formed through the following attributes see [3–5]:

- *Search attributes* which are determined and evaluated *before* the product or a service is delivered (e.g., price or design).
- *Experience attributes* which are experienced either *during* or *shortly after* the delivery (e.g., taste or duration of well being).
- *Belief* (or *credence*) *attributes* which may be impossible to evaluate even a *long time after* the delivery (few customers possess technical skills sufficient to evaluate whether the car maintenance procedures were necessary and performed in a proper manner).

Because the customer expectations are unavoidable, companies usually try to *create* such expectations that can and will be met with the actual product or service performance. One way to effectively create expectations is the use of branding (brands function as a quality warranty, signaling reliability and consistency of a product or service).

Apart from customer expectations, *customer intentions* can also be used as subjective nonmonetary quality standards. This is problematic when they reflect the ideal rather than actual buying behavior of customers who act as quality evaluators.

Further differences of the production and marketing approach with respect to the use of statistical methods in quality improvement are summarized in Table 1.

One of the most frequently used standard survey measurement instruments in quality improvement is SERVQUAL (discussed in **SERVQUAL Surveys**).

Integrating Internal and External Quality Measures

The idea to link internal (facilitated by the production approach) and external (facilitated by the marketing approach) quality measures is not a new one. Management problems are often associated with coordinating, evaluating, and using both sets of quality measures to improve decision-making processes [6].

At the general level, the need to integrate the production/operations and marketing approaches, and to address interactions between them with the goal of improving quality has been widely recognized [7, 8]. Kordupleski *et al.* [9] write

“... there must be a link between customer needs and processes that can be directly managed. We should attempt to construct, for each customer need, an internal metric associated as closely as possible with that customer need. If we are successful, there should be a strong statistical relationship between the internal metric and the quality score for the corresponding customer need” (p. 85).

In the service sector, companies have traditionally managed services by manipulating operations attributes and observing market outcomes [10]. The problem is that in their initial efforts to increase customer satisfaction, many companies tend to focus on tracking customer survey ratings over time or

2 Internal and External Quality Measures

Table 1 Comparison of production and marketing approach to quality improvement

Criteria of comparison	Production approach	Marketing approach
Theoretical background	Classical texts on statistical quality control	Texts on service management and marketing
Product bundle	Emphasis on tangible components	Emphasis on intangible components
Predominant focus	Internal (back office – design and production – technical quality)	External (front office – delivery – perceived quality)
Standards	Objective (precisely determinable)	Subjective (customer expectations, customer intentions, and so on)
Statistical toolbox	Traditional SQC toolbox	Survey toolbox
Mode of application	Ex-ante, real time, and ex-post	Rarely in real time, usually ex-post
Goal of application	Design quality into service processes to prevent negative service experiences	Use customer feedback to discover areas in need of quality improvement
Quality of tangible components	Measured directly with an array of different tools	Measured indirectly through expert and customer observations
Quality of intangible components	No tools available	Measured indirectly through expert and customer observations
Resulting measures	Direct process measures	Indirect customer perception measures
Time frame of measurement	Measures available relatively quickly	Measures available with considerable delay
Cost of measurement	Measures obtained at a low cost	Measures obtained at a high cost (due to survey design, implementation, and evaluation)
Frequency of application	Very high (e.g., hourly, daily ...)	Very low (e.g., annually or biannually)
Required level of employee quantitative literacy	Includes some of the simplest methods applicable with a minimum of educational effort	Includes some of the most complex analytical methods – proper training necessary to avoid the danger of simplification through pull-down menus in statistical software packages

benchmarking them against those of competition. In other words, customer survey ratings become a goal in their own right, when they should rather be integrated with internal quality measures. Low customer survey ratings should be understood as clear signals of the necessity to redesign service processes [11] and *vice versa* as it should be possible for companies to predict how changes in design affect customer satisfaction, behavior, and revenues.

Systems for developing and delivering services to the customer have the same inherent possibility for human weakness, error, and failure as the systems used for product development and delivery in the manufacturing sector [12]. Consequently, opportunities for quality improvement in the service

and manufacturing sectors are based on similar premises:

- less-than-satisfactory quality internal to an organization results in higher operating costs;
- less-than-satisfactory quality as delivered to the customer causes a negative impact on revenues.

In other words, manufacturing companies could also greatly benefit from integration of internal and external quality measures.

Simple as such integration might seem in theory it is not easy to implement in practice. A number of problems are encountered exclusively at the level of external quality measures. *Direct external quality measures* such as customer survey ratings

are available with considerable delay. Fortunately, there also exist *indirect external quality measures*, which are available on a continuous basis given that certain prerequisites are fulfilled. *Purchase volume* and *purchase frequency* are two such measures which tend to be closely linked to customer survey ratings.

To explore the nature of relationships between internal and external quality measures, data collected in different types of company/customer contacts (e.g., at the point of sale, through personal or telephone complaints, thank you letters, customer surveys, and so on) should be available simultaneously. Preferably, they should also be linked together by means of a unique customer identifier.

There are not that many cases where companies could get hold of these data simultaneously. Although technical requirements to uniquely identify customers, and to collect, store, and analyze their purchase data, have already been fulfilled, some kind of a formal relationship has to exist between a company and its customers.

Creation of formal relationships to facilitate customer **data collection** seems the easiest in the framework of a loyalty program, although other possibilities also exist. An additional prerequisite for attempts at integrating internal and external quality measures is the company's ability to roam the data warehouse harboring internal and external data, and face the challenge of building adequate integrative models.

References

- [1] Ishikawa, K. (1985). What is total quality control? *The Japanese Way*, Prentice Hall, Englewood Cliffs.
- [2] Ograjenšek, I. (2002). Applying statistical tools to improve quality in the service sector, in *Developments in Social Science Methodology*, A. Ferligoj & A. Mrvar, eds, Faculty of Social Sciences, Ljubljana, pp. 239–251.
- [3] Zeithaml, V.A. & Bitner, M.J. (1996). *Services Marketing*, McGraw-Hill, New York.
- [4] Berry, L.L. & Parasuraman, A. (1991). Marketing services, *Competing Through Quality*, The Free Press, New York.
- [5] Parasuraman, A., Zeithaml, V.A. & Berry, L.L. (1985). A conceptual model of service quality and its implications for future research, *Journal of Marketing* 49, 41–50.
- [6] Collier, D.A. (1991). Evaluating marketing and operations service quality information. A preliminary report: service quality, in *Multidisciplinary and Multinational Perspectives*, S.W. Brown, E. Gummesson, B. Edvardsson & B. Gustavsson, eds, Lexington Books, Lexington, pp. 143–154.
- [7] Easton, G.S. (1995). A Baldrige examiner's assessment of US total quality management, in *The Death and Life of the American Quality Movement*, R.E. Cole, ed, Oxford University Press, New York, pp. 11–41.
- [8] Cook, D.P., Chon-Huat, G. & Chung, H.C. (1999). Service typologies: a state of the art survey, *Production and Operations Management* 3, 318–338.
- [9] Kordupleski, R., Rust, R. & Zahorik, A. (1995). Marketing and total quality management, in *The Death and Life of the American Quality Movement*, R.E. Cole, ed, Oxford University Press, New York, pp. 77–92.
- [10] Bolton, R.N. & Drew, J.H. (1994). Linking customer satisfaction to service operations and outcomes. Service quality, in *New Directions in Theory and Practice*, R.T. Rust & R.L. Oliver, eds, Sage Publications, Thousand Oaks, pp. 173–200.
- [11] Ograjenšek, I. (2002). *Business statistics and service excellence: applicability of statistical methods to continuous quality improvement of service processes*, Doctoral dissertation, University of Ljubljana, Faculty of Economics, Ljubljana.
- [12] Hagan, J. (1994). Management of quality, in *Strategies to Improve Quality and the Bottom Line*, ASQC Quality Press, Milwaukee.

IRENA OGRAJENŠEK

Sampling Designs in Surveys

Introduction

Surveys are very popular. They help to detect structures in a population and their results inform decision making. In market research, a typical question is whether a new product will find a market. In empirical social sciences surveys investigate attitudes, beliefs, and behavior patterns of groups of peoples. Opinion polls are surveys of opinion often in connection with election studies. Surveys carried out by national statistical offices yield important statistics that are then officially published, e.g. the number of unemployed. Most surveys have one thing in common: only a small number of people are asked a series of questions, and their answers are then extrapolated to a larger group. There are very different ways to select respondents for a survey. If statistical inferences or generalizations from the selected persons are planned, sampling designs based on **probability theory** are adequate techniques. For a small selection of recommended books on sampling see [1–6].

Populations

Before starting a survey, an investigator should exactly define which *population* he or she wants to study and what the objective of the survey is, i.e., what information on *population characteristics* should be collected, e.g., the percentage of patients satisfied with a hospital service, the total expenses of a class of customers for a type of product, or the value of articles in an inventory. This is often not an easy task since the definition of this *target population* needs to specify regional, temporal, and substantive aspects. Usually a population of N distinct units (patients, customers, articles) can be identified with the set $U = \{1, \dots, N\}$. Each unit i bears a characteristic value y_i , e.g., a satisfaction indicator with $y_i = 1$ for “satisfaction”, $y_i = 0$ for “dissatisfaction”, the expenses y_i of customer i , the value y_i of article i . Most often, as in the above examples, the interesting population characteristic is the population total $Y = \sum_{i=1}^N y_i$, or, equivalently, the population mean $\bar{Y} = Y/N$. If it is possible to compile a list of all elements of the target population then this is an excellent

starting point for implementing random sampling. Employee attitude surveys are good examples of surveys that can use easily available lists. Sometimes, such lists do not exist and the researcher has to look for an adequate *frame* to obtain access to the elements of interest. In a telephone survey, telephone books are the basis for sample selection and form the so-called sampling frame. The telephone numbers are the sampling units of the sampling frame. Obviously, two *errors of coverage* can appear. First, *overcoverage*, i.e. there are elements in the sample, which do not belong to the target population, such as business telephone numbers in a survey of private households only. Second, people without telephone are not included and therefore have a zero chance of being selected. This is called *undercoverage*. When planning a survey, it is important to carefully consider the amount of overcoverage and undercoverage and its influence on the data and the interpretation of results.

Sampling

A *sample* s of size n from a population $U = \{1, \dots, N\}$ is a sequence $s = (i_1, \dots, i_n)$ of elements $i_1, \dots, i_n$ selected from U . The order $i_1, \dots, i_n$ of the sample components reflects the order of appearance, if the sample results from a draw-by-draw mechanism. Under *sampling with replacement*, repetitions can occur among the components $i_1, \dots, i_n$, whereas under *sampling without replacement* the components are distinct from each other. Associated with the sample elements $i_1, \dots, i_n$ are the *data* d defined by the sequence of pairs $d = \{(i_1, y_{i_1}), \dots, (i_n, y_{i_n})\}$, shortly $d = \{(i, y_i) : i \in s\}$ where $i \in s$ indicates the enumeration of the components of the sample s .

We distinguish between *random samples* and *non-random samples*. A simple example for a nonrandom sample is a television show that holds a vote on the best song on their show – obviously, only people watching the show can actually form a part of the sample. Other examples are quota samples or convenience samples. *Random samples*, as considered in the sequel, result from a nondeterministic selection mechanism, which is characterized by the selection probabilities $p(s)$ for the possible sample s . The probabilities $p(s)$ determine the *sampling design*. Under *simple random sampling (SRS)* all samples have the same probability of being selected. Hence, in case of SRS without replacement $p(s) = 1/\binom{N}{n}n!$ for all s

2 Sampling Designs in Surveys

and in case of SRS with replacement, $p(s) = 1/N^n$ for all s . Given a design $p(s)$, $\pi_i, i = 1, \dots, N$, denotes the *inclusion probability* of unit i in the sample. In the SRS case $\pi_i = n/N$ for $i = 1, \dots, N$. The second-order inclusion probability of selecting unit i and j in the sample is denoted by π_{ij} and $\pi_{ii} = \pi_i$. For $i \neq j$ in SRS without replacement, $\pi_{ij} = n(n-1)/N(N-1)$ in contrast to SRS with replacement, where $\pi_{ij} = \pi_i \pi_j = n^2/N^2$.

National statistical offices often use *systematic sampling* in their surveys. In its simplest form, systematic sampling can be described as follows, assuming for simplicity that $N = n \cdot H$, where H is the sampling interval: randomly draw a number K ($1 \leq K \leq H$) and select the elements $K, K+H, \dots, K+(n-1)H$. In Figure 1 this selection process is illustrated for $n = 6$.

If the elements of the population are suitably ordered we can attain an efficiency gain compared to SRS. A suitable ordering is given, e.g., if the units of the populations are ordered according to the sizes of the y values. On the contrary, if the ordering is bad, systematic sampling can actually generate worse precision than SRS. Such a case arises if the y values within the samples are more homogeneous, i.e. more similar, than in the population.

In many cases the population is partitioned into H disjunct strata $U_h, h = 1, \dots, H$, of size N_h , i.e. $U = U_1 \cup \dots \cup U_H$. Examples are the states of a country or the business divisions of a company. *Stratified selection* is defined by independent selection of n_h elements of each stratum $U_h, h = 1, \dots, H$. The advantage of stratified sampling compared to systematic random sampling can be enormous if the strata are more homogeneous – with respect to essential characteristics – than the population as a whole. The important question is how to allocate the **sample size** n to the strata.

A simple possibility is *proportional allocation* where

$$n_h = \frac{N_h}{N} \cdot n, \quad h = 1, \dots, H \quad (1)$$

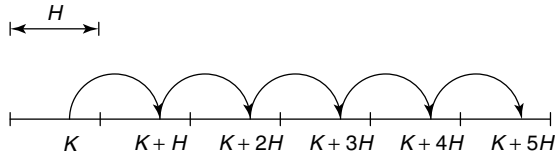


Figure 1 Systematic sampling

i.e., the relative weight n_h/n of a stratum in the sample equals the relative weight N_h/N of the stratum in the population.

Optimal allocation is another allocation procedure, which, however, additionally requires information on variances $S_{yy}(h) = \sum_{i=1}^{N_h} (y_{hi} - \bar{Y}_h)^2 / (N_h - 1)$ of the unit characteristics $y_{h1}, \dots, y_{hN_h}$ in stratum $h, h = 1, \dots, H$. $\bar{Y}_h = \sum_{i=1}^{N_h} y_{hi} / N_h$ denotes the population mean of stratum $h, h = 1, \dots, H$. An example of an optimal allocation is the so-called Neyman allocation where

$$n_h = n \cdot \frac{N_h \sqrt{S_{yy}(h)}}{\sum_{g=1}^H N_g \sqrt{S_{yy}(g)}}, \quad h = 1, \dots, H \quad (2)$$

Since the computation of proportional or optimal allocation does not lead to integers, rounding rules must be taken into consideration.

Cluster sampling selects not single persons but clusters of persons, e.g. classes in a school. This kind of selection procedure is often chosen because of cost considerations. Usually, cluster sampling should be avoided if the clusters are homogeneous. Systematic sampling is a simple example of cluster sampling where the different clusters are defined by the possible samples and only one cluster has to be selected. SRS, stratified sampling, and cluster sampling are special cases of single-stage designs.

Stratified sampling and cluster sampling can also be understood as extreme examples of *multistage sampling*. In multistage sampling the population is divided up into partitions, which in this case are called *primary sampling units* (PSUs). Each PSU contains *secondary sampling units* (SSUs), and so on. The PSUs may be households and the SSUs the persons living therein. Two-stage sampling means that PSUs are selected by a random procedure, and SSUs are then randomly selected from these PSUs. This is a very common procedure in general social surveys.

In *two-phase sampling*, the first step consists of drawing a sample of the population and collecting easily available (inexpensive) information for these sample elements. In the next step, depending on this information a subsample of the sample is selected. For the elements of this subsample relevant but more expensive information can be collected. Inventory samples are examples for multiphase sampling.

In practice, sampling designs are generally very complex and are often a combination of stratified and multistage sampling. Usually, the probabilities of selection of the PSUs are not the same for all PSUs but depend on the size of the SSUs.

Estimation

After selecting the sample, the data are collected for example by interviewing, measuring, or observing. Once the data $d = \{(i, y_i) : i \in s\}$ are collected, the next question is how to estimate unknown quantities of the population such as the percentage of patients satisfied with a hospital service, the total expenses of a class of customers for a type of product, or the value of articles in an inventory. It is highly advisable to use an estimator that is adequate particularly with regard to the underlying sample design. The customary requirements for an adequate estimator are (nearly) unbiasedness and small mean square error or variance. If SRS is used, the sample mean $\bar{y}_s = \sum_{i \in s} y_i / n$ is a good estimator since it is unbiased for the population mean $\bar{Y}$. But if some auxiliary information $x_i, i = 1, \dots, N$, is available it should be also incorporated into the estimator. Denoting the population and sample mean of the x values by $\bar{X}$ and $\bar{x}_s$, respectively, the *ratio estimator* $(\bar{y}_s / \bar{x}_s) \cdot \bar{X}$ has smaller variance than the sample mean provided that the x and y values are nearly proportional for all elements of the population. If selection of the elements is done with unequal probabilities then the *Horvitz–Thompson estimator* $\hat{Y}_{HT} = \sum_{i \in s} y_i / \pi_i$ is design unbiased for the population total if all π_i are positive. The variance of the Horvitz–Thompson estimator is given by

$$V[\hat{Y}_{HT}] = \sum_{i,j=1}^N \frac{y_i y_j}{\pi_i \pi_j} (\pi_{ij} - \pi_i \pi_j) \quad (3)$$

The variance is a measure of quality of an estimator. Since it is unknown it has to be estimated from the sample. If the sample design is complex it can be very difficult to determine π_{ij} . As an example, the computation of the π_{ij} 's for cluster sampling with varying cluster sizes can be cumbersome if the clusters are ordered randomly and the selection of the clusters is done by an extension of systematic sampling. Then *resampling procedures* such as

jackknife or **bootstrap** are applied. If we are not interested in the whole population but in domains and the number of selected elements in a special domain is very small then the precision of a direct estimator such as the sample mean can be very poor. In such cases, techniques developed in *small area estimation* may help by borrowing strength from other sources. For example, prevalence of a particular disease may be adequately estimated from a study at a national level. If it is known that the prevalence of this disease is highly correlated with a set of predictor variables known at the state level then small area estimation can be used to obtain estimates on state level from the national data with greater precision than one would observe in state level, direct survey estimates. For an extensive treatment of small area estimation see [7].

Sample-Size Estimation

Before starting a survey the researchers must determine how many elements have to be selected. Besides the costs, it is necessary to clarify what magnitude of error can be tolerated. Consider the estimation of a population mean $\bar{Y}$ by an estimator $\hat{\bar{Y}}$. Usually, two objectives of controlling the absolute error by guaranteeing $P(|\hat{\bar{Y}} - \bar{Y}| \leq e) \geq 1 - \alpha$ and of controlling the relative error by guaranteeing $P(|(\hat{\bar{Y}} - \bar{Y}) / \bar{Y}| \leq e) \geq 1 - \alpha$ are distinguished. e is called the *margin of error* and $1 - \alpha$ the *confidence level*, say 0.95. To solve both problems at least approximately, **normal distribution** of the sampling distribution of $\hat{\bar{Y}}$ is assumed as well as the existence of a consistent variance estimator $\hat{V}(\hat{\bar{Y}})$, which allows building of **confidence intervals** on grounds of this approximation. In the case of SRS without replacement $P\left(|\hat{\bar{Y}} - \bar{Y}| \leq z_{1-\alpha/2} \sqrt{V(\hat{\bar{Y}})}\right) \geq 1 - \alpha$

with $V(\hat{\bar{Y}}) = 1/(N-1) \sum_{i=1}^N (y_i - \bar{Y})^2 (1 - \frac{n}{N})/n$, where $z_{1-\alpha/2}$ is the constant exceeded with probability $\alpha/2$ by the standard normal random variable. Since the variance of the estimator is unknown it must be estimated. If the objective is controlling the relative error then the *coefficient of variation* $CV_y = \sqrt{V(\hat{\bar{Y}})} / \bar{Y}$ has to be considered. The sample size necessary to achieve the relative error can be

derived as

$$n = \frac{\left(\frac{z_{1-\alpha/2} \cdot C V_y}{e}\right)^2}{1 + \frac{1}{N} \left(\frac{z_{1-\alpha/2} \cdot C V_y}{e}\right)^2} \quad (4)$$

If the a population mean $\bar{Y}$ is a proportion P , e.g., the proportion of unemployed, the sample proportion p has variance $V(p) = P(1 - P)/n \left(1 - \frac{n-1}{N-1}\right)$. If we do not have any idea about the magnitude of P we use the upper bound $V(p) \leq 1/4n \left(1 - \frac{n-1}{N-1}\right) < 1/4n$ where the first bound is adopted for $P = 0.5$. Controlling the absolute error we derive as necessary sample size

$$n = \frac{z_{1-\alpha/2}^2}{4e^2} \quad (5)$$

For example, if the margin of error is 3% we have to select 1067 elements to achieve 0.95 confidence level.

Confidence intervals and sample-size computations for complex sample designs can be found in [7].

Nonresponse

Sometimes, answers for some units of the selected sample are either totally or partly missing. This is referred to as *nonresponse*. Sources for *unit nonresponse* can be not-at-homes, refusals, unable to answer or not found. If a respondent answers some but not all the items, these missing items are referred to as *item nonresponse*. Nonresponse is a very serious problem and the treatment of which requires modeling assumptions. Therefore, it may be easier to deal with nonresponse in the framework of model-based inference. Usually, weighting procedures are applied for treating unit nonresponse and **imputation** methods for item nonresponse. **Calibration** estimators are

used very often nowadays to get a grip on the non-response problem, see [8].

Software

In the past, only specialized statistical software packages offered procedures for sampling. In the meantime, a lot of statistical software packages such as SAS, SPSS, Stata, or R enable the user not only to draw samples but also to analyze data collected by complex sampling procedures. The estimation of the *design effect*, defined formally by Kish in [4] as “the ratio of the actual variance of a sample to the variance of a simple random sample of the same number of elements”, is thereby very helpful. In this way the researcher is protected against error induced by the classical unrealistic assumption that data are independent identically distributed observations yielding confidence intervals that are too small and thus provoking a wrong interpretation of the output.

References

- [1] Chambers, R.L. & Skinner, C.J. (2003). *Analysis of Survey Data*, John Wiley & Sons, New York.
- [2] Cochran, W.G. (1977). *Sampling Techniques*, 3rd Edition, John Wiley & Sons, New York.
- [3] Kalton, G. (1983). *Introduction to Survey Sampling*, SAGE Publications, Newbury Park.
- [4] Kish, L. (1965). *Survey Sampling*, John Wiley & Sons, New York.
- [5] Lohr, S.L. (1999). *Sampling: Design and Analysis*, Duxbury Press.
- [6] Särndal, C.E., Swensson, B. & Wretman, J. (1992). *Model Assisted Survey Sampling*, Springer, New York.
- [7] Rao, J.N.K. (2003). *Small Area Estimation*, John Wiley & Sons, New York.
- [8] Lundström, S. & Särndal, C.E. (2002). *Estimation in the Presence of Nonresponse and Frame Imperfections*, Statistics Sweden.

SIEGFRIED GABLER

Adaptive and Spatial Sampling Designs

Introduction

Simple random sampling and other **basic sampling designs** belong to the class of *noninformative* designs. A sampling design is noninformative if the probability for sample selection does not depend on the characteristic of interest, say y . In contrast, a sampling design is *informative* if the selection of an element depends on the y values observed during the survey. A typical example for an informative sampling design is *sequential sampling* where the selection of a unit depends on the y values of the selected units preceding. In *adaptive sampling*, popularized by Thompson [1], units in the neighborhood of units in the sampling process with certain y values are sampled, too. Examples can be found in surveys assessing the abundance of rare species of animals or trees, in epidemiological surveys, in behavioral surveys [2], or in surveys on hazardous waste sites. A very new application of adaptive sampling is given in the field of web surveys [3].

In *spatial sampling* the characteristic of interest often is defined for a continuous region of space and time. Spatial sampling is an area of survey sampling applicable for sampling in two (or more) dimensions, e.g. sampling of fields.

Adaptive Sampling

In classical sampling we usually have noninformative designs such as simple random sampling. Here the selection probability is free of the characteristic of interest. Adaptive sampling designs are informative since an observed y value of a selected unit satisfies a special condition and the sampling process will then be finished or additional units will be added to the sample from its neighborhood. There are two kinds of designs, *sequential stopping designs* and *adaptive cluster designs*. Sequential stopping designs are not considered here.

The Sampling Process of an Adaptive Sample

In adaptive cluster sampling, in a first step an *initial sample* is selected. The initial sample is based

on a conventional sampling design such as simple random sampling, systematic sampling, cluster sampling, stratified sampling, or unequal probability sampling. Starting with the units of the initial sample, other units of the population in the neighborhood of the units in the initial sample are added. The final sample is called *adaptive sample*. One reason for such a procedure is that the units with special characteristics of interest are clustered.

The following simple example illustrates the whole sampling process.

Example 1 A certain species of plants accumulates in a study area. From the grid of 10×10 squares covering the study area, two squares are randomly selected (Figure 1a). For each selected square with at least one exemplar of the species, its neighbor squares (i.e., one common side) are added to the initial sample (Figure 1b). This process of sampling continues until no further exemplars are encountered.

Estimation

In adaptive cluster designs, conventional estimators used in noninformative designs are no longer unbiased and must be replaced by adequate estimators. These are presented in the section titled “Special Cases for the Initial Sample”.

Efficiency

Under certain conditions, adaptive strategies can be more *efficient* than simple random sampling of equivalent **sample size**. In [2] the following seven conditions are mentioned under which efficiency of adaptive cluster sampling possibly increases compared to conventional random sampling: (a) the within-network variance is a high proportion of the total population variance (i.e., the population is clustered or aggregated with high variability within aggregations); (b) the population is rare; (c) the expected final sample size with adaptive cluster sampling is not much larger than the initial sample size; (d) the costs of observing units in clusters or networks are less than the costs of observing the same number selected at random; (e) the costs of observing units not satisfying the condition are less than the costs of observing units satisfying the condition; (f) the condition for extra sampling may be based on an auxiliary variable easy to measure; and (g) an efficient estimator or Rao–Blackwell

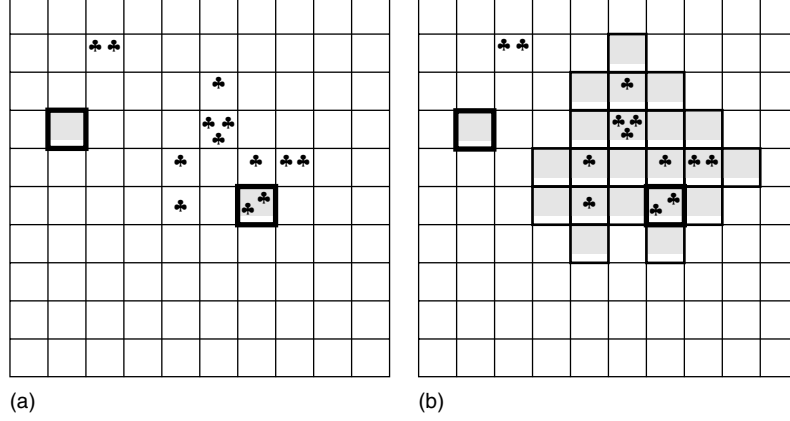


Figure 1 (a) Initial sample quadrants. (b) Selected quadrants after augmentation

improved estimator is used with adaptive cluster sampling.

A Formal Description of Adaptive Cluster Sampling

Consider a population $U = \{1, \dots, N\}$ of size N . The population characteristic of interest is the *population total* $Y = \sum_{i=1}^N y_i$, or, equivalently, the *population mean* $\bar{Y} = Y/N$. A unit where a certain prespecified condition c that depends on its y value is not satisfied is called *edge unit*. In Example 1, there are 92 edge units. A *network* $C(i)$ consists of all units accessible from unit i by augmentation and by subsequently dropping all edge units. $C(i)$ depends on the selection method of the initial sample. An edge unit is also called a *singleton network*. Thus, in Example 1 there are 92 singleton networks, one network of size 2 and 7 networks of size 11. Defining

$$\bar{Y}_{C(i)} = \frac{1}{|C(i)|} \sum_{j \in C(i)} y_j \quad (1)$$

where $|C(i)|$ denotes the number of units in network, it is easy to see that

$$\sum_{i=1}^N \bar{Y}_{C(i)} = Y \quad (2)$$

Thus, estimating $\sum_{i=1}^N \bar{Y}_{C(i)}$ is equivalent to estimating the population total Y . If s is a sample, i.e., a subset of U , and $e(s, y_i | i \in s)$ is an unbiased

estimator of Y , then $e(s, \bar{Y}_{C(i)} | i \in s)$ is unbiased for Y , too. $C(s) = \cup_{i \in s} C(i)$ is an extension of s and is called an *adaptive sample*.

Closely related to adaptive sampling is *network sampling* where rare and clustered populations have to be sampled, e.g. in drug abuse research, and contact patterns are used for sampling. An overview of such link-tracing procedures is given in [4].

Special Cases for the Initial Sample

As already pointed out, the selection of the initial sample is an important component of the adaptive cluster sampling process. It determines the final adaptive sample as well as the estimator. In the section titled “Initial Simple Random Sample without Replacement”, the case will be considered when the initial sample is drawn by simple random sample without replacement. The section titled “Initial Unequal Probability Sampling” considers more general sample designs for the initial sample.

Initial Simple Random Sample without Replacement. Let π_i denote the probability that the initial sample intersects the network $C(i)$ for unit i . Then

$$\pi_i = 1 - \frac{\binom{N - |C(i)|}{n}}{\binom{N}{n}} \quad (3)$$

and the *Horvitz–Thompson estimator* $\hat{Y}_{HT}$ of Y is

$$\hat{Y}_{HT} = \sum_{i=1}^N \frac{y_i}{\pi_i} L_i, \quad (4)$$

$$L_i = \begin{cases} 1 & \text{if the initial sample} \\ & \text{intersects } C(i) \\ 0 & \text{otherwise} \end{cases}$$

Since π_i is constant for each unit i in the k th network, say $\pi_i = \tilde{\pi}_k, k = 1, \dots, K$, where K is the number of distinct networks in the population, we have

$$\hat{Y}_{HT} = \sum_{k=1}^K \frac{\tilde{y}_k}{\tilde{\pi}_k} \tilde{L}_k, \quad (5)$$

$$\tilde{L}_k = \begin{cases} 1 & \text{if the initial sample intersects} \\ & \text{kth network} \\ 0 & \text{otherwise} \end{cases}$$

where $\tilde{y}_k$ is the sum of the y values for the k th network.

If $\tilde{\pi}_{k\ell}$ denotes the probability that networks k and ℓ are both intersected, from classical theory of the Horvitz–Thompson estimator we get

$$E[\hat{Y}_{HT}] = Y \quad (6)$$

$$\text{Var}[\hat{Y}_{HT}] = \sum_{k=1}^K \sum_{\ell=1}^K (\tilde{\pi}_{k\ell} - \tilde{\pi}_k \tilde{\pi}_\ell) \frac{\tilde{y}_k}{\tilde{\pi}_k} \frac{\tilde{y}_\ell}{\tilde{\pi}_\ell}. \quad (7)$$

An unbiased variance estimator is given by

$$\widehat{\text{Var}}[\hat{Y}_{HT}] = \sum_{k=1}^K \sum_{\ell=1}^K \left(\frac{\tilde{\pi}_{k\ell} - \tilde{\pi}_k \tilde{\pi}_\ell}{\tilde{\pi}_{k\ell}} \right) \frac{\tilde{y}_k}{\tilde{\pi}_k} \frac{\tilde{y}_\ell}{\tilde{\pi}_\ell} \tilde{L}_k \tilde{L}_\ell \quad (8)$$

In Example 1 we get as an estimator for the total of the species of plants in the grid

$$\hat{Y}_{HT} = \sum_{k=1}^K \frac{\tilde{y}_k}{\tilde{\pi}_k} \tilde{L}_k = \frac{13}{0.13576} = 95.75893 \quad (9)$$

and for the variance estimation

$$\begin{aligned} \widehat{\text{Var}}[\hat{Y}_{HT}] &= \sum_{k=1}^K \sum_{\ell=1}^K \left(\frac{\tilde{\pi}_{k\ell} - \tilde{\pi}_k \tilde{\pi}_\ell}{\tilde{\pi}_{k\ell}} \right) \frac{\tilde{y}_k}{\tilde{\pi}_k} \frac{\tilde{y}_\ell}{\tilde{\pi}_\ell} \tilde{L}_k \tilde{L}_\ell \\ &= \frac{13^2}{0.13576} \left(\frac{1}{0.13576} - 1 \right) \\ &= 7924.90633 \end{aligned} \quad (10)$$

Initial Unequal Probability Sampling. Let $p(s)$ be the selection probability for the initial sample s , and let π_i, π_{ij} be the sample inclusion probabilities for population members i , respectively for pairs $i \neq j$. Following [5], we can use

$$e(s, \bar{Y}_{C(k)} | k \in s) = \sum_{i=1}^N \bar{Y}_{C(i)} b_{si} L_i \text{ with } E_p[b_{si} L_i] = 1 \text{ for all } i \quad (11)$$

as an estimator for Y . With

$$d_{ij} = E_p[b_{si} L_i - 1][b_{sj} L_j - 1], \alpha_i = \sum_{j=1}^N d_{ij} w_j$$

where $w_i (\neq 0)$ are constants, (12)

$$\begin{aligned} \text{Var}_p[e(s, \bar{Y}_{C(k)} | k \in s)] &= -\frac{1}{2} \sum_{i=1}^N \sum_{j=1}^N d_{ij} w_i w_j \\ &\quad \times \left(\frac{\bar{Y}_{C(i)}}{w_i} - \frac{\bar{Y}_{C(j)}}{w_j} \right)^2 \\ &\quad + \sum_{i=1}^N \frac{\bar{Y}_{C(i)}^2}{w_i} \alpha_i \end{aligned} \quad (13)$$

and an unbiased variance estimator is given by

$$\begin{aligned} &-\frac{1}{2} \sum_{i=1}^N \sum_{j=1}^N d_{ij} w_i w_j \left(\frac{\bar{Y}_{C(i)}}{w_i} - \frac{\bar{Y}_{C(j)}}{w_j} \right)^2 \frac{L_i L_j}{\pi_{ij}} \\ &+ \sum_{i=1}^N \frac{\bar{Y}_{C(i)}^2}{w_i} \alpha_i \frac{L_i}{\pi_i} \end{aligned} \quad (14)$$

provided all π_{ij} are positive.

Spatial Sampling

The increased incidence of thyroid carcinoma among children affected by the Chernobyl fallout, or more generally, the spatial distribution of contamination are examples where the spatial aspect plays an important role. If the objective is to estimate the contaminant concentration exceeding a fixed level, then the question of how to select a sample is somewhat different from the classical *finite-population approach* but related to the *superpopulation approach*. A good overview is given in [6].

The field of spatial sampling is split into two approaches, a model-based one and a design-based one. A closer look at both approaches is given in [7]. The main problem of spatial sampling is how to deal with the presence of spatial **correlation**, which is quite typical in problems concerning contamination. Adaptive cluster sampling as described in the section titled “Adaptive Sampling” is a design-based approach and is commonly used in a fixed-population approach. In a model-based approach the classical sampling framework is extended and the characteristic of interest is defined for a continuous region of space and time that can be described by a spatial stochastic process which has been elaborated in detail in [8].

The index set of the stochastic process corresponds to the unit labels $\{1, \dots, N\}$ of a fixed *target population*. The *sampling frame* is the list of *sampling units* that identifies every unit within the target population. The sampling units of the sampling frame in finite-population sampling now correspond to the collection of points or functions. More precisely, a spatial process is a real (or vector valued) process $\{Z(s) : s \in U\}$. The index set U is a region in space or in space and time.

The random variable $Z(s)$ represents the value of the characteristic of interest at s , a location or point. The total population is the integral of the process over a study region D . There exist various models that are widely used for Z , e.g.,

$$E[Z(s)] = \sum_{j=1}^p \beta_j F_j(s) \quad (15)$$

where the F_j are given functions and β_j are unknown parameters. In applications the underlying models take into consideration a spatial structure in so far that measurements at near sites tend to have more homogeneous, i.e., more similar values than more distant sites. Thus, the **covariance** function

$$C(s, t) = \text{Cov}[Z(s), Z(t)] \quad (16)$$

plays an important role. Usually $C(s, t)$ is assumed to be *stationary*, i.e.,

$$C(s, t) = c(s - t) \quad (17)$$

The function $C(s, t)$ is called *covariogram*.

If observations $Z(s_1), \dots, Z(s_n)$ at sites $s_1, \dots, s_n$ in the study region D are available it may be of interest to predict $Z(t)$ at a site in the study region D or to predict the integral of a linear or nonlinear function of the process over a subregion of D , e.g., the population total.

There exist several criteria to obtain optimal sites, e.g., the integrated **mean squared error** (IMSE) which minimizes

$$IMSE(s_1, \dots, s_n) = \int_D E \left[(\hat{Z}(t) - Z(t))^2 \right] dt \quad (18)$$

where $\hat{Z}(t)$ is a linear unbiased predictor of $Z(t)$. Examples are given in [9].

References

- [1] Thompson, S.K. & Seber, G.A.F. (1996). *Adaptive Sampling*, John Wiley & Sons, New York.
- [2] Thompson, S.K. (1997). Adaptive sampling in behavioral surveys, in *The Validity of Self-Reported Drug Use: Improving the Accuracy of Survey Estimates*, NIDA Research Monograph 167, L. Harrison & A. Hughes, eds, National Institute of Drug Abuse, Rockville, pp. 296–319.
- [3] Thompson, S.K. (2006). Adaptive web sampling, *Biometrics* **62**, 1224–1234.
- [4] Spreen, M. (1992). Rare populations, hidden populations, and link-tracing designs: what and why? *Bulletin de Methodologie Sociologique* **36**, 34–58.
- [5] Chaudhuri, A. & Stenger, H. (2005). *Survey Sampling: Theory and Methods*, 2nd Edition, Taylor & Francis Group LLC.
- [6] Cox, D., Cox, L. & Ensor, K. (1997). Spatial sampling and the environment: some issues and directions, *Environmental and Ecological Statistics* **4**, 219–233.
- [7] Gruijter, J. & Braak, C. (1990). Model-free estimation from spatial samples: a reappraisal of classical sampling theory, *Mathematical Geology* **4**, 407–415.
- [8] Cressie, N. (1991). *Statistics for Spatial Data*, John Wiley & Sons, New York.
- [9] Thompson, S.K. (1992). *Sampling*, John Wiley & Sons, New York.

SIEGFRIED GABLER

Selection and Validation of Response Scales

Theory and Levels of Measurement

Measurement (be it implicit or explicit) is defined as a set of rules for assigning either numbers or numerals to empirical properties of a concept under consideration (for example, employee satisfaction or customer loyalty).

In this context, *numerals* are numbers which are given a qualitative meaning (for example, 0 – male, 1 – female) according to the procedure specified by the rules. Both numbers and numerals facilitate the application of the tools of statistical analysis for descriptive, explanatory and/or predictive purposes [1].

If the rules for assigning either numbers or numerals to empirical properties of a concept under consideration are poor, measurement will be meaningless (which can be due to the problematic wording of survey questions, respondent **bias**, interviewer bias, etc.). Meaningful measurement is achieved only when it has high empirical correspondence with the concept the researcher plans to measure.

In the process of operationalization, the concept the researcher plans to measure is assigned a single variable, or a number of qualitative (categorical or discreet) and quantitative (numerical) variables. The former can be further classified into nominal and ordinal variables; the latter into interval and ratio variables according to the *level of measurement* or *measurement scale* [2].

The *nominal level* is the lowest level of measurement. At this level numerals or other symbols are used to classify objects or observations. Objects or observations with similar characteristics are assigned the same numeral or symbol. The only comparisons that can be made between objects and observations are those of (in)equality. Only one-to-one scale transformations are admissible for nominal variables.

The second level of measurement is the *ordinal level*. Very often, objects or observations are not only classifiable. They also exhibit some kind of relation, allowing for rank order. In addition to (in)equality, comparisons of “greater than” and “less than” can be made. Admissible scale transformations are strictly increasing.

At the *interval level* of measurement absolute differences can be calculated on top of comparisons of (in)equality, “greater than” and “less than”. Admissible scale transformations are any strictly increasing linear ones. Operations such as subtraction and addition are therefore meaningful. However, owing to the fact that interval scales have no natural or absolute zero, relative comparisons of magnitude (operations such as division and multiplication) are not allowed.

The *ratio level* of measurement is guaranteed with the existence of a natural or absolute zero (one for which a universal agreement on its location is achieved). Admissible scale transformations are any strictly increasing proportional ones. For variables with properties of the ratio measurement scale relative comparisons of magnitude are thus allowed – on top of comparisons of (in)equality, “greater than” and “less than”, as well as mathematical operations such as subtraction and division.

To summarize, the level of measurement determines which mathematical operations are possible on a variable, which, in turn, determine the tools of statistical analysis at researchers’ disposal (see Table 1).

Survey Measurement

General types of questions used in surveys are discussed in **Design and Testing of Questionnaires**. In this framework we will take a look at selected specific measurement tools applicable in the survey setting. Among them are list questions, category questions, ranking questions, and rating or scale questions.

List questions present the respondents with a list of responses from which to choose (see Example 1).

The response categories have to be defined clearly and meaningfully to all respondents. Also, all possible relevant options, which might be considered by any individual respondent, need to be accounted for. The response categories are very heterogeneous and may include options such as “yes/no” or “agree/disagree”; furthermore “don’t know” and “not sure”; “other” (in case the list of responses is not exhaustive); as well as “applies/does not apply” [3].

In contrast to list questions, *category questions* are designed with the goal of making respondents choose a single category. This is very useful when collecting data on behavior or attributes. The maximum number

2 Selection and Validation of Response Scales

Table 1 Scales of measurement^(a)

Scale	Acceptable comparisons	Illustrative examples	Acceptable measures of central tendency ^(b)
Nominal	$y_A = y_B$ $y_A \neq y_B$	Gender (male, female) Marital status (single, married, divorced, widowed)	Mode
Ordinal	$y_A > y_B$ $y_A = y_B$ $y_A < y_B$	Social class (lower, middle, upper) Level of agreement (disagree, neutral, agree)	Mode Median
Interval	$d = y_A - y_B$	Temperature	Mode Median Arithmetic mean
Ratio	$c = \frac{y_A}{y_B}$	Income Sales Number of employees	Mode Median Arithmetic mean Geometric mean

^(a)The measure of central tendency which gives the most useful information is printed in bold.

^(b) In behavioral sciences there exists a controversy over the use of mean for variables with the properties of the ordinal measurement scale. From the viewpoint of the measurement theory the use of mean for ordinal data is not permitted because the arithmetic operations are not meaningful on numbers which are not “true measurements” with predefined units of measurement. However, many behavioral scientists do calculate the mean for ordinal data under justification that although the interval difference between two ordinal ranks is not constant, it is often of the same order of magnitude (*see SERVQUAL Surveys* for further discussion).

of categories to be included is dependent on the type of questionnaire. Questions in self-administered and telephone questionnaires should have no more than five response categories [4]. Questions in structured interviews can have more categories provided that:

- a prompt card is used (see Example 2).
- the interviewer categorizes the responses (see Example 3).

Example 2 A category question: use of a prompt card

Example 1 A list question

Please tick the appropriate box(es) for services provided to you by the IT department in the past month.		
Description of service	Provided	Not provided
Hardware purchase		
Hardware maintenance		
Software licence		
agreement update		
Lost file(s) recovery		
Other (please specify):		

Which of the following search engines did you use during the last week? <i>Show respondent the prompt card with the logos of search engines. Read out the name of the search engines one at a time. Tick the appropriate box.</i>			
Search engine	Use	Did not use	Do not know
Google			
MSN search			
Yahoo			
Other (please specify):			

Example 3 A category question: interviewer categorizes the responses

How often did you go to the movies during the last 3-month period? <i>Listen to the respondent's answer and tick the appropriate box.</i>
Two or more times a week Once a week Less than once a week to fortnightly Less than fortnightly to once a month Less often

Example 4 A ranking question

Please number each of the factors listed below in order of importance to you in your choice of a car rental company. Number the most important 1, the next 2, and so on. If a factor has no importance at all, please leave blank.	
Factor	Importance
Opening hours	
24-h hotline	
Price	
Special offers	
Loyalty programme	
Other (please specify):	

Ranking questions are used when researchers want to find out the relative importance of a single response item (for example, of a product feature or the factor determining a purchase decision) in comparison to other applicable response items. Such questions should have very clear instructions and not more than seven or eight items to be ranked as the respondents are capable of ranking accurately only when they can see or remember all items. In a telephone survey, where respondents can only rely

on their memory, there should be a maximum of only three or four items to rank [5].

Example 5 A rating or scale question

For the following statement please tick the box that matches your view most closely. <i>I fear about my job security.</i>
Disagree completely Disagree Neither disagree nor agree Agree Agree completely

When analyzing respondent opinions, researchers usually rely on *rating or scale questions*. The most commonly used is the Likert-style rating scale in which the respondent is asked how strongly he or she (dis)agrees with a statement or a series of statements. The level of (dis)agreement is usually measured on a scale that has from 4 to up to 10 points.

A typical Likert-style scale applied to measure service quality or employee satisfaction would include five points:

1. completely disagree
2. disagree
3. neither disagree nor agree
4. agree
5. completely agree.

A lot of controversy is involved with the issue of whether the number of points should be even (without the neutral middle point) or odd. However, there is a widespread agreement that the number of negative and positive response options for a single question should be in balance.

Two most common variations of the Likert-style rating scale are the numeric rating scale (see Example 6) and the semantic differential (see Example 7). The numeric rating scale is used if the researcher wants the numbers to reflect the feeling of the respondent. The semantic differential is applicable in consumer research projects focused on determining underlying consumer attitudes. Here, the respondent

4 Selection and Validation of Response Scales

is asked to rate a single object or idea on a series of bipolar rating scales described by a pair of opposite adjectives.

Example 6 The numeric rating scale

For the following statement please circle the number which matches your view most closely. <i>This packaged holiday was . . .</i>											
1	2	3	4	5	6	7	8	9	10		
Poor value for money										Good value for money	

Example 7 A semantic differential

On each of the lines below place an x to show how you feel about the quality of lectures in the course on business ethics.		
Good handouts	----	Poor handouts
Boring	----	Inspiring
Appropriate use of multimedia	----	Inappropriate use of multimedia

Reliability and Validity of Measurement

Whenever we attempt to measure something, we want our measures to be reliable and valid, capturing what they are supposed to capture. The observed measurement score might reflect the true score as well as the influence of other factors. These might be *stable personal characteristics* (one person might be more generally talkative than another), *transient personal factors* (such as moods), as well as *situational factors* (time pressure, mode of questionnaire administration, data entry errors, etc.).

Reliability is defined as the measurement instrument's consistency. It tells us whether the same measures yield the same results on different occasions.

Validity is concerned with the following questions:

- whether the research findings are really what they appear to be about (whether we measure what we intended to – the so-called internal validity);
- whether the research findings support the intended interpretation of the variables (the so-called construct validity);
- whether the research findings can be generalized at the level of statistical population from which the sample was drawn (the so-called statistical conclusion validity);
- whether the research findings can be generalized beyond the level of statistical population from which the sample was drawn (the so-called external validity).

Reliability does not imply validity: a reliable measure might be consistent in measuring something yet that something is not necessarily what the measure is supposed to be measuring.

In the survey setting, threats to reliability include participant error, participant bias, interviewer error, as well as interviewer bias. Some of the important threats to validity are [6]:

- history (related traumatic past experiences of the respondents may have misleading effects on research findings);
- testing effects (respondents may realize how to use the measurement instrument to their advantage);
- mortality (respondents may drop out of the study and thus jeopardize generalizability of research findings);
- maturation (change of attitudes, opinions, and knowledge of respondents during the study duration may affect research findings);
- ambiguity about the causal direction between the researched concepts (the famous egg/hen problem makes for a difficult interpretation of research findings).

In order to estimate reliability, researchers can apply either single-administration or multiple-administration methods. Two single-administration methods (the split-half method and the internal consistency method) and two multiple-administration

methods (the test-retest method and the alternate forms method) are briefly described in the following paragraphs.

The *split-half method* estimates reliability by treating each of two or more parts of a measurement instrument as a separate scale. It is based on the Spearman–Brown prediction (or prophecy) formula:

$$r_{xx'} = \frac{2r_{oe}}{1 + r_{oe}} \quad (1)$$

where $r_{xx'}$ is the reliability of the original test and r_{oe} is the reliability coefficient obtained by correlating the scores of the odd statements with the scores of the even statements.

The *internal consistency method* correlates each scale item with the total score, giving the researcher a tool for identifying the items to drop (those with the lowest correlations) and the items to keep (those with the highest correlations). The method is based on the Cronbach's α which measures the extent to which the scale items are connected. The α can be calculated in several ways; the formula for the standardized α is

$$\alpha = \frac{N \cdot \bar{r}}{(1 + (N - 1) \cdot \bar{r})} \quad (2)$$

where N is the number of scale items and $\bar{r}$ is the average of all (Pearson) **correlation** coefficients between them. The value of α will generally increase with increased correlations between the scale items. In behavioral sciences an acceptable level of α is 0.70 or higher, demonstrating an acceptable level of reliability.

The essence of the *test-retest method* is administering the measurement instrument to the same group of persons at two different points in time, and then calculating the correlation between the two sets of observations (scores). The reliability coefficient is thus based on the Pearson product-moment correlation coefficients between two subsequent administrations of the same measure.

In the framework of the *alternate forms method* the researcher develops two alternate versions of a measurement instrument, which are supposed to be equivalent in their ability to measure the desired concept. Both versions are administered to the same group of people, and correlation coefficients are calculated between two sets of observations (scores). The reliability coefficient is thus based on the Pearson product-moment correlation coefficients between two

different forms of measure administered at a single point in time.

Triangulation

In the process of measurement, data can be obtained either in formal or informal settings and can represent verbal (oral and written) as well as nonverbal acts and/or responses [7]. This combination (formal *versus* informal setting; verbal *versus* nonverbal acts/responses) accounts for the following major methods of **data collection**: library research, observation studies, survey research, and experimental research.

Should researchers wish to study nonverbal acts in an informal setting, they might decide on using participant observation as a method of data collection. Should the focus be on verbal responses in a structured setting, a standardized questionnaire might be the data collection tool of choice. Between these two opposite poles there exists a number of other type of setting/type of response combinations.

Each data collection method has its advantages and disadvantages. It is an established fact that research findings are affected by the nature of the selected data collection method. As pointed out by Fiske [8]:

"Each method [of data collection] is ... one discriminable way of knowing" (p. 62).

In other words, by exclusively using one method of data collection researchers run the risk to discover only one piece of the whole puzzle.

In order to minimize the possible effect of the chosen data collection method on the validity and/or scope of the research findings, researchers can use two or more methods of data collection to test hypotheses and measure variables. Discovering different pieces of the puzzle and putting them together by applying different methods of data collection is thereby the essence of *triangulation*. Are the findings of different data collection methods consistent with one another, this greatly increases the validity of research.

References

- [1] Ghauri, P. & Grønhaug, K. (2002). *Research Methods in Business Studies*, FT Prentice Hall, Harlow.
- [2] Stevens, S.S. (1946). On the theory of scales of measurement, *Science* **103**, 677–680.

6 Selection and Validation of Response Scales

- [3] Saunders, M., Lewis, P. & Thornhill, A. (2003). *Research Methods for Business Students*, Prentice Hall, Harlow.
- [4] Fink, A. (1995). *How to Ask Survey Questions*, Sage Publications, Thousand Oaks.
- [5] Kervin, J.B. (1999). *Methods for Business Research*, Addison-Wesley, Reading.
- [6] Robson, C. (2002). *Real World Research*, Blackwell Science, Oxford.
- [7] Frankfort-Nachmias, C. & Nachmias, D. (2000). *Research Methods in the Social Sciences*, Worth Publishers and St. Martin's Press, New York.
- [8] Fiske, D.W. (1986). Specificity of method and knowledge in social sciences, in *Metatheory in Social Science; Pluralisms and Subjectivities*, D.W. Fiske & R.A. Shweder, eds, University of Chicago Press, Chicago, pp. 61–82.

IRENA OGRAJENŠEK

Design and Testing of Questionnaires

Questionnaire and Questionnaire Surveys

Questionnaire is the most frequently used **data collection** tool of survey research. It is constructed with the goals of measuring attitudes, opinions, perceptions and knowledge, and comparing or explaining practices or behavior. In other words, it helps researchers finding out what a selected group of respondents know, think, or feel – something which cannot be done by means of an observation study, or an experiment.

Any form of a questionnaire survey is an interaction between a researcher and a respondent [1]. In a self-administered questionnaire survey with no interviewer-respondent interaction, the researcher speaks to the respondent directly through a questionnaire. In an interviewer-administered questionnaire survey it is the interviewer who reads the words of the researcher to the respondent (in the process he or she might introduce the so-called interviewer **bias** into the research). In both cases it is paramount that every respondent is asked the questions in exactly the same way as the others. Additionally, the same question should mean the same thing to every respondent.

The most frequently used types of questionnaire surveys are presented in Table 1.

Table 2 summarizes the most important characteristics of personal, telephone, mail and on-line questionnaire surveys.

Two factors that determine the questionnaire design requirements to the largest extent are the form

of the questionnaire (paper *versus* electronic), and the mode of questionnaire administration (interviewer-*versus* self-administered questionnaire). We look at them more closely in the next section.

Questionnaire Design

Advance information on the subject matter of the questionnaire survey is usually obtained by means of a *focus group* discussion. A focus group helps researchers to gather data related to the feelings and opinions of a group of people with similar experiences or involved in a similar situation in a neutral setting. Listening to other group members' views encourages focus group participants to voice their own opinion, which would be less accessible without a group interaction [5]. This interaction facilitates coverage of all issues related to the chosen topic, which sets the boundaries of the subsequent questionnaire survey. Although the focus group data are mostly of qualitative nature, they do provide the researcher with a guide to the most pressing matters to be addressed, and the most relevant questions to be asked, in the questionnaire.

The *general questionnaire design principles* pertain to individual questionnaire items regardless of the questionnaire form and mode of administration. According to these principles, each questionnaire item should:

- be explicitly related to goals and objectives of the research;
- be understood by all respondents in the same way;
- ask for information which can be either immediately recalled by respondents or retrieved from their records;

Table 1 Types of questionnaire surveys^(a)

Interviewer-administered questionnaire		Self-administered questionnaire	
On paper	On screen	On paper	On screen
PAPI ^(b) in person	CAPI ^(c)	Diary	CASI ^(d)
PAPI ^(b) by telephone	CATI ^(e)	Mail survey	On-line survey

^(a) Adapted after Biemer and Lyberg [[2], p. 189].

^(b) Paper and pencil interviewing.

^(c) Computer-assisted personal interviewing.

^(d) Computer-assisted self-interviewing.

^(e) Computer-assisted telephone interviewing.

2 Design and Testing of Questionnaires

Table 2 Characteristics of the four most common types of questionnaire surveys^(a)

Characteristic	CAPI ^(b)	CATI ^(c)	Mail survey	On-line survey
Cost	High	Medium	Low	Low
Timeliness	Medium	High	Low	High
Duration of individual data collection process	Long	Short (up to 30 min)	Medium (up to 8 pages)	Medium
Response rate	High	Medium	Low	Low
Sample size	Depends on the number of interviewers	Depends on the number of interviewers	Large	Large
Geographic dispersion of sample units	Depends on available budget	Broad	Broad	Broad
Access	Total population	Telephone owners	Total population	Individuals with Internet access
Control of respondent selection	High	Medium	Low	Medium
Required level of respondent literacy	Low	Low	Functional literacy	Functional and computer literacy
Types of questions	Open or closed, can be complex	Short and simple, open or closed; if closed with a small number of possible answers	Interesting but simple, predominantly closed with a large number of possible answers in a logical order	Interesting but simple, predominantly closed with a large number of possible answers
Diversity of questions	High	Low	Medium	Medium
Use of reference materials	Possible	Only use of audio references possible	Only use of visual references possible	Use of multimedia possible
Possibility of collecting sensitive data	Low	Medium	High	High
Danger of socially desirable answers	High	Medium	Low	Low
Predominant type of bias	Interviewer	Interviewer	Other persons	None
Data entry	If CAPI, simultaneous manual entry; if PAPI, after the data collection process (manual entry or scanning)	Simultaneous manual entry	After the data collection process (manual entry or scanning)	Simultaneous manual entry

^(a) Adapted after Saunders *et al.* [3, p. 284]; Malhotra and Birks [4, p. 234]; Biemer and Lyberg [2, p. 211].

^(b) Computer-assisted personal interviewing.

^(c) Computer-assisted telephone interviewing.

- be constructed in such a way as to prevent introduction of a respondent and interviewer bias into research;
- be formatted in a way that minimizes data entry errors.

If applied consistently, these general principles should guarantee that in practice, every questionnaire functions as planned by the researcher.

Additionally, *specific questionnaire design principles* should be followed for different questionnaire forms and modes of administration.

A *self-administered questionnaire* should provide the respondent with clear explanations and instructions for every questionnaire item or group of items. Also, contact information (a toll-free phone number or an e-mail address) should be provided in case of any respondent queries. This information can be either a part of an accompanying letter (for example, in case of a mail survey) or part of the respondent address at the beginning of an on-line questionnaire.

An *interviewer-administered questionnaire* needs to be designed in such a way that it caters to the needs of both the interviewer and the respondent. To insure that all interviewers ask the same question in the same manner, standards for interviewer behavior, question presentation and probing should be established in advance during the interviewer training sessions which should take place regardless of the length and complexity of the questionnaire.

When designing a *paper questionnaire*, special attention should be paid to easiness and convenience of questionnaire completion on one, as well as easiness and convenience of data entry on the other hand. Clearly linking questions and responses, keeping designated response spaces in the right-hand column area, using different typesets for instructions and questions, and so on facilitates both.

When designing a *questionnaire in the electronic form*, easiness and convenience of questionnaire completion are the sole most important concerns as the recording of responses and data entry take place simultaneously. Speed and reliability of modern internet connections give questionnaire designers a lot of room for development of attractive interfaces and – in case of a self-administered questionnaire – use of multimedia.

Sufficient time and resources should be allocated to questionnaire translation and testing should a survey be conducted in different languages.

Types of Questions

Survey questions take one of two primary forms. *Open questions* require respondents who are both capable and willing to use their own words when answering. This means that answers to open questions must be catalogued, coded and interpreted which makes the use of open questions less applicable in large-scale surveys but very important whenever respondents' own words are essential. The latter is very often the case in the initial phase of questionnaire design (where the number and heterogeneity of possible response options to a single question are assessed). Consider the question in Example 1.

Example 1: An open question and possible answers

Question: What kind of car are you driving?

Answer 1: A red one Answer 2: A really fast

one Answer 3: Old Answer 4: I don't own a car

Answer 5: A Renault

Closed questions have a preselected number of response options which should be mutually exclusive. They must be known in advance which makes closed questions much more difficult to write than open questions.

Closed questions tend to be more efficient and reliable when trying to obtain information from large groups of respondents. They seem to be preferred by respondents who are either unwilling or unable to express themselves while being surveyed [6].

While closed questions mean more work in the questionnaire design phase, they produce standardized data which are easily statistically analyzed.

When designing survey questions, researchers should be aware of respondent characteristics. A *set of general rules* which should be followed in the question design process includes a number of items [5–8]:

- *Explain the purpose of the questionnaire to all respondents – and then ask purposeful logically related questions.* For instance, in a survey about a theme park experience respondents will not expect questions about the next elections.

4 Design and Testing of Questionnaires

- *Ask precise and unambiguous questions.* Questions can be labeled as precise and unambiguous if a number of respondents agree on the words used without prompting.
- *Use time periods that are related to the importance of the question.* Periods of a year or more can be used for major life events (birth of a child, divorce, death of a parent). Periods of a month or less should be used for less important issues.
- *Do not use slang, colloquial expression, jargon, specialist language, and abbreviations.* Whenever possible (unless you are dealing with a very homogeneous group whose members share a special language), stick to conventional wording and rely on standard grammar, punctuation, and spelling.
- *Avoid using biasing words and phrases.* Words such as “nigger” or “abortion”, and phrases such as “freedom fighters”, might trigger emotional responses, or even prejudices which have little to do with the issues addressed by the survey. Biasing words and phrases may differ in different cultures, which is why all questions must be reviewed and pilot tested before they are used.
- *Avoid asking negative questions.* Negative questions such as “Do you agree with the statement that the European Union shouldn’t have the authority to intervene in the member states’ internal affairs?” tend to be misunderstood and/or misinterpreted by respondents because they require an exercise in logical thinking.
- *Avoid asking offensive or insensitive questions.* They might cause embarrassment or pose a threat to the respondent. If included in a mail or on-line questionnaire, these questions negatively influence the rate of response. If at all, such questions should be asked in a face-to-face interview – and even in that case they should be related to past indiscretions rather than the respondent’s present behavior.
- *Avoid the use of leading questions.* It is against the ethical code of research to steer respondents toward the preferred response option.
- *Ask only one question at a time.* In other words, avoid double-barreled questions such as “Do you think the customers and employees will benefit from the newly introduced company initiative?”
- *Do not be afraid to ask value-laden questions.* Sometimes this might be the only way to encourage respondents to give a “true” rather than a

socially acceptable response. However, use this type of questions with caution – you do not want the respondents to get annoyed and either not answer the question or answer it inaccurately.

- *Do not be afraid to ask factual questions.* However, do not degrade them to questions which merely perform the function of a memory test. Also, avoid factual questions which require the respondents to perform calculations.

By adhering to these general rules, you should be able to keep your questionnaire as short as possible, while including all the questions required in achieving your research goals.

Note that it is often best to use questionnaire items that have been tested in the past and found to yield good results. A number of readily available standard measurement instruments are available for different topic areas in a broad number of disciplines (see **SERVQUAL Surveys** for a discussion of the **SERVQUAL**, a rating scale to measure perceived service quality). They facilitate survey comparability in time and across countries, but might limit the scope of the planned research project because of the boundaries associated with question wording and order.

Questionnaire Testing and Piloting

In order to test and evaluate a questionnaire, a wide range of methods can be used. Some (e.g., an expert panel or a peer evaluation) are more applicable in the early stages of questionnaire design, others (e.g., a pilot study) should be used prior to the broad questionnaire distribution. The most important ones are briefly discussed in the following paragraphs.

Both *peer evaluation* and an *expert panel* serve the same purpose: to provide the researcher with preliminary feedback concerning the draft version of a questionnaire. The feedback can be of technical nature and/or pertaining to the issues of content. It serves as a basis for the first few questionnaire revisions.

Although peer and/or expert reviews can be very effective in producing good quality questionnaires, it is also extremely useful to pretest the questionnaire on individuals who are representative of the population to be surveyed. A *split panel test* is an experimental comparison of two or more

versions of the same questionnaire. These are used to evaluate the impact of question order, question wording, and data collection method on survey results.

Cognitive one-to-one interviews are conducted by specially trained interviewers with the following two objectives: to probe the respondents' understanding of individual questions, and to examine the respondents' thinking processes as they provide their answers. Cognitive interviews thus help researchers understand and eliminate potential problems associated with question wording.

A *pilot study* is the most comprehensive and therefore also the most expensive of applicable questionnaire testing and piloting methods. Its scope is broader than that of a split panel test or a cognitive interview because it provides researchers with the feedback on all steps of the data collection process.

Tools of statistical analysis should also be mentioned when it comes to questionnaire design and testing. **Correlation** coefficients might be used to reduce the number of highly related questions. Cronbach's α might be calculated to find out how well a set of items (or variables) measures a single hypothesized unidimensional latent construct. Structural equation modeling can be applied to assess the measurement

error stemming from questionnaire design – and to correct for it [9].

References

- [1] Fowler Jr, F.J. (1988). *Survey Research Methods*, Sage Publications, Newbury Park.
- [2] Biemer, P.P. & Lyberg, L.E. (2003). *Introduction to Survey Quality*, John Wiley & Sons.
- [3] Saunders, M., Lewis, P. & Thornhill, A. (2003). *Research Methods for Business Students*, Prentice Hall, Harlow.
- [4] Malhotra, N.K. & Birks, D.F. (2003). *Marketing Research: An Applied Approach*, Pearson Education Limited, Essex.
- [5] Hussey, J. & Hussey, R. (1997). *Business Research: A Practical Guide for Undergraduate and Graduate Students*, Macmillan Press, London.
- [6] Fink, A. (1995). *How to Ask Survey Questions*, Sage Publications, Thousand Oaks.
- [7] Bregar, L., Ograjenšek, I. & Bavdaž, M. (2005). *Metode raziskovalnega dela za ekonomiste. Izbrane teme*, Faculty of Economics, Ljubljana [In Slovenian].
- [8] Zikmund, W.G. (2003). *Business Research Methods*, Thomson South-Western, Mason.
- [9] Saris, W.E. (1998). *The Effects of Measurement Error in Cross Cultural Research*, ZUMA-Nachrichten Spezial, January 1998.

IRENA OGRAJENŠEK

Data Management in Survey Sampling

Introduction

Traditional **data collection** methods include postal mail surveys, telephone, and face-to-face surveys. Recently, new technologies have led to a rapid emergence of a wide array of new data collection methods. This entry focuses first on new technologies related to data collection. It subsequently turns its attention to technological advances in the field of data management.

New Technologies Used to Collect Data

Interviewer-Administered Surveys

Paper-and-pencil telephone and face-to-face surveys were largely replaced by their computerized counterparts in the 1970s [1]. These well-established computer-assisted survey information collection (CASIC) methods are *computer-assisted telephone interviewing* (CATI) and *computer-administered personal interviewing* (CAPI), respectively.

Recent advances involve miniaturization of input devices in CAPI surveys (using smaller laptop computers such as pen computers and PDAs [2]) and keeping up with changing technology use patterns in society. One such particularly important change is the emergence of mobile-phone-only households. “Dual frame” survey designs have been proposed to keep those households covered in CATI surveys, but these designs introduce additional chances of selection for individuals and households, which have both landline and mobile telephone numbers (*sampling bias*). Statistical correction procedures are necessary to address this problem.

Self-Administered Surveys

Even before a sizable portion of the general population possessed computers, computerized self-administered data collection was implemented in CAPI surveys. Whenever necessary, but usually when sensitive questions were asked, the interviewer could turn over the laptop to the respondent to fill in

that part of the survey questionnaire by himself. This is believed to increase the respondent’s sense of privacy, which would increase the level of honesty of the reports [3]. This type of data collection has been termed *computer-assisted self interview* (CASI). Partly to deal with respondents with limited reading skills, *audio computer-assisted self interview* (ACASI) was developed. A recorded voice reading out the question is played to the respondent (usually through earphones). If no aural cues are provided, one can speak of *video computer-assisted self interview* (VCASI), and reserve *CASI* as a broader term encompassing both ACASI and VCASI. A telephone version of audio computer-assisted self interview (*T-ACASI*) also exists [4]. A similar survey type is *interactive voice response* (IVR). Selected respondents are asked to call a toll-free phone number or are transferred to the survey by the interviewer who has recruited them. Respondents listen to prerecorded questions, voice-read answers, response options, and instructions for which numbers they can use for their response.

Computerized self-administered questionnaires (CSAQ) are sent to respondents in electronic form. Respondents fill them out on their computers and send the responses back to the survey organization. An early form was the *disk by mail* (DBM) system, in which the questionnaire was stored on a floppy disk and sent by postal mail. In contrast, *e-mail surveys* send the questionnaire (either as attachment or in-line) by e-mail. The World Wide Web allows surveyors to put *web surveys* (“forms”) on the Internet (*computer-assisted web interviewing*, CAWI). The web survey address is communicated through e-mail, other on-line communication (such as posting the web survey address in newsgroups, banners on web sites, etc.), or off-line communication (e.g., web survey addresses printed in off-line media). With the arrival of *digital*, *interactive*, and *web TV*, similar surveys could be implemented without requiring respondents to own a computer.

Since mobile phones and PDAs can also connect to the Internet using protocols such as WAP (wireless application protocol) or GPRS (general packet radio service), web surveys can also be completed using these devices. Nonetheless, these devices represent some limitations such as a small screen, low resolution, or slow Internet connections. Recently, *SMS* (*short message service*) surveys have also started to emerge. These surveys target mobile phone

users and typically ask only a small number of questions since SMS messages are limited to a relatively small number of characters.

New technologies may speed up data collection and reduce costs, but some problems do arise. New data collection methods such as web or mobile phone surveys in many countries face the problem of nonexistent *sampling frames*, which adequately cover the general population (coverage error) as well as the problem of low *response rates* (nonresponse error). Convenience samples based on volunteer panels are sometimes used. Efforts have been undertaken to develop strategies to overcome generalization problems associated with convenience samples, for instance, by using randomly drawn *calibration samples* and application of propensity score adjustment (PSA) procedures [5].

Concern also exists about the *quality of data* collected with emerging technologies. Web surveys, for instance, generate higher levels of unit nonresponse, item nonresponse, do-not-know responses, and acquiescence than CATI surveys [6]. With self-administered questionnaires, usability is important to ensure sufficiently high **data quality** levels. New technologies have facilitated data collection needed to evaluate usability such as eye tracking [7] and keystroke registration [8]. New methods, particularly CSAQs, typically also reach relatively distinct population segments (younger, more highly educated Caucasians with higher income levels who are comfortable using new technologies), decreasing the representativity of these samples. Methodological procedures (e.g., offering incentives) and statistical correction procedures (e.g., PSA) are needed to alleviate these problems.

One proposed strategy to combine advantages of new methods with advantages of more traditional methods is to combine different survey modes. *Mixed-mode survey* designs deploy different survey modes either simultaneously or sequentially to follow-up nonrespondents from a previous mode. Although mixed-mode surveys are rapidly gaining popularity, a serious drawback is that they may produce mode effects [9].

New Technologies Used in Data Management

Fieldwork organizations rely on automated systems to manage the interviewing system and the collected

data. At the risk of oversimplifying, interviewing system management refers primarily to *case management*. This includes sample management, on-line call scheduling (in CATI systems), automated call history record keeping, maintenance and reporting of call outcome data, transfer cases to other interviewers, preparation of datasets, and so on [10]. With the increasing popularity of mixed-mode surveys, case management systems must also be capable of transferring sample cases from one survey mode to another.

Recent advances eliminate the need of interviewers to be physically present in the research organization to connect to the system. New systems allow telephone interviewers to operate from their own homes or CAPI interviewers to conduct surveys in the homes of the respondents, while still being connected to the central system, e.g., through the Internet.

Recently, much attention is being paid to *process data*. These data can be used to monitor the fieldwork process. If e.g., data on contact attempts and their outcomes are continually registered and fed into the computerized data management system, data collection protocols can be adapted during the fieldwork. Typically, the aim is to reduce error on an array of relevant variables for which the distribution in the study population is known (e.g., age, type of residence). Research on such “responsive survey designs” is still in the exploration phase [11].

Other aspects of data management are *data coding* and *data editing* (see “**survey data processing**”). These steps can be at least partially automated [12]. Recent advances in data editing include definition and use of a so-called abstract data model of a survey and software that automatically derives from this model the functional form of edits and statistical criteria to clean data [13]. New systems allow both data coding and data editing to be performed during the interview. Such integration is associated with a certain risk in that original data may be lost, eliminating the possibility to evaluate and if necessary correct the automated procedures.

References

- [1] Couper, M.P. & Nicholls, II, W.L. (1998). The history and development of computer assisted survey information collection methods, in *Computer Assisted Survey Information Collection*, M.P. Couper, R.P.

- Baker, J. Bethlehem, C.Z.F. Clark, J. Martin, W.L. Nicholls, II & J.M. O'Reilly, eds, John Wiley & Sons, New York, pp. 1–21.
- [2] Bosley, J., Conrad, F.G. & Uglow, D. (1998). Pen CASIC: design and usability, in *Computer Assisted Survey Information Collection*, M.P. Couper, R.P. Baker, J. Bethlehem, C.Z.F. Clark, J. Martin, W.L. Nicholls, II & J.M. O'Reilly, eds, John Wiley & Sons, New York, pp. 521–541.
- [3] Hewitt, M. (2002). Attitudes toward interview mode and comparability of reporting sexual behavior by personal interview and audio computer-assisted self-interviewing, *Sociological Methods & Research* **31**, 3–26.
- [4] Turner, C.F., Forsyth, B.H., O'Reilly, J.M., Cooley, P.C., Smith, T.K., Rogers, S.M. & Miller, H.G. (1998). Automated self-interviewing and the survey measurement of sensitive behaviors, in *Computer Assisted Survey Information Collection*, M.P. Couper, R.P. Baker, J. Bethlehem, C.Z.F. Clark, J. Martin, W.L. Nicholls, II & J.M. O'Reilly, eds, John Wiley & Sons, New York, pp. 455–473.
- [5] Lee, S. (2006). Propensity score adjustment as a weighting scheme for volunteer panel web surveys, *Journal of Official Statistics* **22**, 329–349.
- [6] Fricker, S., Galesic, M., Tourangeau, R. & Yan, T. (2005). An experimental comparison of web and telephone surveys, *Public Opinion Quarterly* **69**, 370–392.
- [7] Redline, C.D. & Lankford, C.P. (2001). Eye-movement analysis: a new tool for evaluating the design of visually administered instruments (paper and web), in Paper presented at the *American Association of Public Opinion Research (AAPOR) Annual Conference*, Montreal, May 2001.
- [8] Caspar, R.A. & Couper, M.P. (1997). Using keystroke files to assess respondent difficulties with and audio-CASI instrument, in *Proceedings of the Survey Research Methods Section, American Statistical Association*. Available at <http://www.amstat.org/sections/SRMS/Proceedings/>.
- [9] Dillman, D.A. (2000). *Mail and Internet Surveys: The Tailored Design Method*, John Wiley & Sons, New York.
- [10] Connet, W.E. (1998). Automated management of survey data: an overview, in *Computer Assisted Survey Information Collection*, M.P. Couper, R.P. Baker, J. Bethlehem, C.Z.F. Clark, J. Martin, W.L. Nicholls, II & J.M. O'Reilly, eds, John Wiley & Sons, New York, pp. 245–262.
- [11] Groves, R.M. & Heeringa, S.G. (2006). Responsive design for household surveys: tools for actively controlling survey errors and costs, *Journal of the Royal Statistical Society. Series A (Statistics in Society)* **169**, 439–457.
- [12] Speizer, H. & Buckley, P. (1998). Automatic coding of survey data, in *Computer Assisted Survey Information Collection*, M.P. Couper, R.P. Baker, J. Bethlehem, C.Z.F. Clark, J. Martin, W.L. Nicholls, II & J.M. O'Reilly, eds, John Wiley & Sons, New York, pp. 223–243.
- [13] Petrakos, G., Conversano, C., Farmakis, G., Mola, F., Siciliano, R. & Stavropoulos, P. (2004). New ways of specifying data edits, *Journal of the Royal Statistical Society. Series A (Statistics in Society)* **167**, 249–274.

DIRK HEERWEGH

Processing of Survey Data

Introduction

After the data are collected, a number of processing steps must be performed to convert the **survey data** from its raw and unedited state to a verified and corrected state that is ready for analysis and/or dissemination to the users. If the data were collected by *paper and pencil* (PAPI) methods, they will have to be converted into a computer-readable form. Responses to open-ended questions may need to be classified into some number of categories using a coding scheme so that these responses can be tabulated. In addition, a number of operations may be performed on the data to reduce survey errors and missing items.

For example, the data may be “cleaned” by eliminating inconsistencies and addressing unlikely responses (e.g., *outliers*). Survey weights may be computed that account for unequal selection probabilities. These weights may be further refined by a series of “postsurvey adjustments” that are intended to reduce coverage error **bias**, nonresponse bias, and sampling variance. Some survey variables (for example, household income) may have numerous missing values and values may be *imputed* for them. Following these steps, the data contents file should be well documented. Data *masking* and *de-identification* techniques may also be conducted on the file to protect the confidentiality of the respondents. The next section provides a brief overview of these data processing activities.

Overview of the Data Processing Steps

The data processing steps vary depending on mode of **data collection** for the survey and technology available to assist in data processing. The steps involved for processing PAPI questionnaires will be discussed initially. These steps include the following:

- scan editing
- data capture
- data editing
- coding
- file preparation
- data documentation
- analysis.

The steps for *computer-assisted interviewing* (CAI) methods are essentially the same except that the data entry step is not necessary since the data are already keyed into the computer by an interviewer.

The PAPI Questionnaire

The PAPI questionnaire is used to collect information for the variables under study. Some questions corresponding to these variables are closed-ended requiring the interviewer (or respondent) to check a box representing a response alternative. For example, marital status may be coded as “1 = single”, “2 = married”, “3 = divorced”, “4 = widow/widower”, and “5 = never married”. For open-ended questions, a free-format response is written in a blank field on the questionnaire. For example, a response to the question “What is your occupation?” may be “I am a security officer for a bank” or “I am a barber. I cut hair.” The questionnaire may also involve branching to skip around questions that do not apply. These legitimately skipped items need to be distinguished from items requiring a response that were not answered by the respondent. The latter type of skip constitutes item-missing data, which is often addressed by the subsequent processing steps.

Scan Editing

Scan editing is an operation involving several steps. First, as questionnaires are received by the survey organization, they are coded as received into the receipt control system and inspected for obvious problems such as blank pages or missing data for key items that must be completed in order for questionnaires to be usable. In *interviewer-assisted* surveys, the questionnaires that are rejected as being incomplete by this process may be sent back to the interviewers for completion. For mail surveys, the questionnaires might be routed to a telephone *follow-up* process for completion. In cases where there is no follow-up of nonrespondents, the questionnaires may be passed on to the next data processing step (i.e., data entry). Ultimately, questionnaires that are not minimally completed will be coded as nonrespondents due to incomplete data. Also as part of the process, questionnaires may be grouped into small batches called *work units* to facilitate the subsequent processing steps.

Data Capture

In *data entry* or *data capture*, the paper questionnaire responses are converted into computer-readable form (i.e., digitized). Data can be entered manually using keying equipment or automatically by using scanning or optical character recognition devices. To use the latter methods, messy questionnaires may have to be copied onto new, clean forms so that the scanner can read them properly. Keying usually involves some form of **quality control** verification. For example, each questionnaire may be keyed independently by two different keyers. Any discrepancies between the first and second keyings are then rectified by the second keyer. Alternatively, **acceptance sampling** methods (typically, a **single sampling plan**) may be applied to each work unit. Here, only a sample of questionnaires within each work unit is rekeyed. If the number of discrepancies between the two keyings exceeds some threshold value, the entire work unit would be rekeyed. Because of these verification methods, the error rate associated with keying is usually quite low for closed-ended responses – less than 0.5%. However, for verbal responses such as names and addresses, the error rates are substantially higher – 5% or more [1].

Data Editing

Editing is a process for verifying that the digitized responses are plausible and, if not, modifying them appropriately. Editing rules can be developed for a single variable or for several variables in combination. The editing rules may specify acceptable values for a variable (e.g., an acceptable range of values) or acceptable relationships between two or more variables (e.g., an acceptable range for the ratio of two variables). Typically editing identifies entries that are either definitely in error (called *critical edits*) or are highly suspicious of being in error (called *query edits*). Corrective measures must be taken for all critical edits while various rules may be applied to determine which query edits to address in order to reduce the cost of the editing process. This approach is sometimes referred to as *selective editing* [2].

Some surveys specify that respondents should be recontacted if the number of edit failures is large or if key survey items are flagged as in error or

questionable. In this manner, missing, inconsistent, and questionable data can be eliminated by the respondent's input. However, this is not always done either to save costs or due to the impracticality of respondent recontacts. In that case, values may be inserted or changed by means of deducing the correct value based upon other information on the questionnaire or from what is known about the sample unit from prior surveys. Consistency checks, selective editing, deductive editing, and other editing functions can be performed automatically by specially designed computer software.

Coding

Coding is a procedure for classifying open-ended responses into predefined categories that are identified by numeric or alphanumeric code numbers. For example, there may be thousands of different responses to the open-ended question “What is your occupation?” To be able to use this information in subsequent analysis, each response is assigned one of a much smaller number (say 300–400) code numbers, which identify the specific occupation category for the response. To make occupation categories consistent across different surveys and different organizations, a standard occupation classification (SOC) system is used. A typical SOC codebook may contain several hundred occupation titles and/or descriptions with a three-digit code number corresponding to each. In most classification standards, the first digit represents a broad or main category, and the second and third digits represent increasingly detailed categories. Thus, for the response “barber”, a coder consults the SOC codebook and looks up the code number for “barber”. Suppose the code number is 411. Then the “4” might correspond to the main category “Personal Appearance Workers” 41 might correspond to “Barbers and Cosmetologists”, and 411 to “Barber”. In *automated coding*, a computer program assigns these code numbers to the majority of the cases while the cases that are too difficult to be accurately coded by computer are coded manually.

File Preparation

The file preparation step results in an analysis file that serves as the output file and the output file is

used to produce statistics and analyses. This step consists of a number of activities including *weighting*, *imputation*, *postsurvey weighting adjustments*, *statistical disclosure analysis*, and *data suppression*. First, weights have to be computed for the sample units. Often the weights are developed in three steps. First, a base weight is constructed, which is the inverse of the probability of selection for a sample member. Second, the base weight is adjusted to compensate for nonresponse error, and third, further adjustments of the weights might be performed to adjust for frame coverage error depending on availability of external information. These so-called *post-survey adjustments* are intended to achieve additional improvements in the accuracy of the estimate.

Imputation is a process that replaces missing responses with values that may be treated just like other responses in the subsequent analysis. It has been shown that imputation reduces the biases in analysis due to item nonresponse. A common method is hot-deck imputation, which fills in missing values using corresponding values from similar people in the same dataset. Cold-deck imputation, by contrast, selects donors from another dataset.

Finally, if the data file is to be released to the public, a statistical disclosure analysis may be performed to determine the risk that a sample member on the file could be identified by an *intruder* on the basis of a sample member's data. This may involve estimating the probability that an individual in the file can be *reidentified* (i.e., his or her identity can be discovered with certainty) by an intruder. If this probability is unacceptably high, then various techniques are available to suppress or mask some responses to avoid identity disclosure and to protect the confidentiality of the sample member (see [3] for an overview of these methods).

Data Documentation and Analysis

The final processing step is data documentation in which a type of data file users manual is created. This document describes the methods used to collect and process the data and provides detailed information on the variables on the file. For example, each variable on the data file might be linked to one or more questions on the questionnaire. If variables were combined, recoded, or derived, the steps involved in creating these variables is described. The documentation might also include information regarding

response rates for the survey, item nonresponse rates, reliability estimates, or other information on the total **survey error** of key variables.

The new technologies have opened up many possibilities to integrate these data processing steps [4]. Therefore, the sequence of steps might be very different from those described above for some surveys. For example, it is possible to integrate data capture and coding into one step; likewise, data capture and editing can be integrated with coding. It is also possible to integrate editing and coding with data collection through the use of CAI technology. The advantage of this type of integration is that inconsistencies in the data or insufficient information for coding can immediately be resolved with the respondent. This reduces follow-up costs and may also result in better information from the respondent. Many other possibilities for combining the various data processing steps exist. The goal of integration is to increase the efficiency of the operations while improving **data quality**.

Data Quality and Data Processing

As is clear from the previous discussion, the data can be modified extensively during data processing. Hence, data processing has the potential to improve data quality for some variables while increasing the error for others. Unfortunately, knowledge about the errors introduced in data processing is very limited in survey organizations and consequently such errors tend to be neglected. Often operations are sometimes run without any particular quality control efforts and the effects of errors on the overall accuracy as measured by the **mean squared error** (MSE) are often unknown except perhaps for national data series of great importance.

As an example, although editing is intended to improve data quality, it misses many errors and can even introduce new ones. Automation can reduce some of the errors made by manual processing, but might introduce new errors. For instance, in optical recognition data capture operations, the recognition errors are not uniformly distributed across digits and other characters, which can introduce systematic errors (i.e., biases). For these reasons, quality control and quality assurance measures should be a standard part of all data processing operations. For more information on data processing steps and their contributions to survey error, see [5].

References

- [1] Biemer, P.P. & Lyberg, L.E. (2003). *Introduction to Survey Quality*, John Wiley & Sons, Hoboken, p. 224.
- [2] Granquist, L. & Kovar, J. (1997). Editing of survey data: how much is enough? in *Survey Measurement and Process Quality*, L. Lyberg, P. Biemer, M. Collins, E. De Leeuw, C. Dippo, N. Schwartz & D. Trewin, eds, John Wiley & Sons, New York, pp. 415–435.
- [3] Willenborg, L. & de Waal, T. (1996). *Statistical Disclosure Control in Practice*, Springer, New York.
- [4] Couper, M. & Nicholls, B. (1998). The history and development of computer assisted survey information collection methods, in *Computer Assisted Information Collection*, M. Couper, R. Baker, J. Bethlehem, C. Clark, J. Martin, W. Nicholls & J. O'Reilly, eds, John Wiley & Sons, New York, pp. 1–21.
- [5] Biemer, P. & Lyberg, L. (2003). *Introduction to Survey Quality*, John Wiley & Sons, Hoboken, NJ.

PAUL P. BIEMER

Statistical Analysis of Survey Data

Introduction

After **data collection** and entry, survey data need to be prepared for analysis and analyzed. The data preparation stage involves three important steps: data formatting, data editing, and computation of survey weights. The data analysis stage involves the choice of relevant statistical tools according to the measurement scale properties of survey variables. Both stages are discussed in the following sections.

Data Preparation Stage

The data preparation stage starts with the decision on *how to format the data* in view of available computing facilities and analytical software. Analysis of survey data often involves two or more different units of analysis. Consequently, the data file needs to provide for the nested hierarchical structures (e.g., employees within a company, members of a household within a household unit, etc.).

An important part of data formatting procedures is the construction of a code – a set of rules translating survey answers into numerals and numbers (for a definition of numerals *see Selection and Validation of Response Scales*). Codes should be unambiguous. If designed appropriately, they help researchers minimize errors in the data analysis phase.

The following is a set of the commonsense principles to facilitate the **coding** process [1]:

- In the database, response options should be numbered in the same order they appear on the questionnaire.
- Whenever possible, codes should fit numbers in the real world (for example, “a 20-year-old” should be coded as “20”).
- Consistency should be maintained when assigning numerals (for example, “0” should be assigned to all male respondents, “1” to all female respondents in a survey).
- It should be ensured that the codes differentiate between the categories such as “don’t know”, “does not apply”, and “missing”.

Existing coding schemes such as standard classifications (classifying products, activities, occupations, etc.) should be used wherever possible because they save time, they are usually well tested, and they allow comparisons of research results with other surveys [2].

Following the decision on data format the next step is *data editing*; *see Processing of Survey Data*. Errors in survey responses (for example, inconsistent responses or responses that fall outside the range of possibilities) need to be identified and dealt with.

A special challenge is presented by items for which answers have not been obtained. The so-called **item nonresponse problem** occurs when a sampled unit cannot be contacted, does not respond to the request to participate in a survey, or refuses to answer to a particular survey question owing to fatigue, sensitivity, lack of knowledge, or other factors [3].

Item nonresponse results both in a reduced **sample size** and in the possibility that the remaining data set is biased. The problem might be overcome through the process of assigning values for the missing responses by a number of **imputation** methods that have been steadily developed since the 1980s.

If the proportion of missing values is small, a single imputation might be quite reasonable. This is rarely the case in real life, which is why Rubin’s multiple imputation (a **Monte Carlo** technique in which the missing values are replaced by $m > 1$ simulated versions, where m is typically small, ranging from 3 to 10) gained such popularity [4]. Following the process of imputation, each of the simulated complete datasets is analyzed by standard methods. The results are combined to produce estimates and **confidence intervals** that incorporate missing-data uncertainty.

Originally, the use of multiple imputation was limited to large data files from censuses and sample surveys. However, with the advent of new computational methods and software that support the process of multiple imputation, this technique has become increasingly attractive for researchers in the biomedical, behavioral, and social sciences [5].

The final step of data preparation involves the *computation of survey weights*. These are computed for all units of analysis. The construction of weights starts with determining the selection probabilities of sample units. The initial (base) weights for sampled

units are then computed as the inverses of the sample units' selection probabilities and adjusted to compensate both for the eligible nonrespondents and for the nonrespondents with unknown eligibility status. Further adjustments are possible if researchers want to make the adjusted weighted sample distributions for identified key variables conform to known distributions of these variables available from external sources [6–8].

Data Analysis Stage

Survey data analysis has two distinguishing characteristics in comparison to standard statistical analysis [6]:

- First, as already stated, there exists the need to use survey weights to compensate for unequal selection probabilities, nonresponse, and noncoverage. The danger of not using weights lies in the possibility of producing distorted estimates of population values.
- Secondly, there exists the need to compute sampling errors for survey estimates in such a way that a survey's complex sample design is accounted for. While the standard statistical analysis assumes unrestricted sampling, many surveys employ stratified multistage sampling. Given that, in general, sampling errors for estimates from a stratified multistage sample are larger than those from an unrestricted sample of the same size, the application of the standard formulas will overstate the precision of the estimates [9–11]. This means that standard statistical software packages produce invalid **standard error** estimates for survey estimates. It is therefore recommended that researchers resort to the use of special survey analysis software packages such as (in alphabetical order) CENVAR, Epi-Info, IVEware, PC-CARP, SAS, STATA, SUDAAN, and WasVar [11].

Often, the nature of survey data analysis is descriptive, e.g., in the form of means, percentages, and/or totals. The results may be presented in a tabular or graphical form to facilitate understanding and interpretation.

Whereas in narrow technical terms the estimates are usually very simple (the proper use of

weights being the researchers' only concern), important methodological challenges have to be dealt with in advance. The important issues are the definition of a construct like "customer satisfaction" that should be measured, operationalization of the construct, and survey units specification. The goal is to make sure that the construct to be measured is measured in the most suitable way.

For complex constructs (for example, poverty, wealth, or literacy), aggregate measures (indices) are often constructed with the help of multivariate statistical methods (e.g., factor or cluster analysis).

While the production of descriptive estimates remains the predominant form of survey data analysis, examination of relationships between variables (with special emphasis on analysis of **causality**) also merits research attention. Statistical techniques applied to this effect are the analysis of association, contingency, and **correlation**, as well as regression analysis. Based on the type of the model variables, the researcher decides whether to use linear regression (in case of continuous cardinal variables), Poisson regression (for count data), or logistic regression (in case of categorical variables); see [8, 12].

References

- [1] Fowler Jr, F.J. (1988). *Survey Research Methods*, Sage Publications, Newbury Park.
- [2] Saunders, M., Lewis, P. & Thornhill, A. (2003). *Research Methods for Business Students*, Prentice Hall, Harlow.
- [3] Groves, R.M., Dillman, D.A., Eltinge, J.L. & Little, R.J.A. (eds) (2002). *Survey Nonresponse*, John Wiley & Sons, New York.
- [4] Rubin, D.B. (1987). *Multiple Imputation for Nonresponse in Surveys*, John Wiley & Sons, New York.
- [5] Schafer, J.L. (1997). *Analysis of Incomplete Multivariate Data*, Chapman & Hall, New York.
- [6] Kalton, G. (2005). Introduction, *Household Sample Surveys in Developing and Transition Countries*, United Nations, New York, pp. 302–304.
- [7] Muñoz, J. (2005). A guide for data management of household surveys, *Household Sample Surveys in Developing and Transition Countries*, United Nations, New York, pp. 305–334.
- [8] Chromy, J.R. & Abeyasekera, S. (2005). Statistical analysis of survey data, *Household Sample Surveys in Developing and Transition Countries*, United Nations, New York, pp. 389–417.
- [9] Kalton, G., Brick, J.M. & Thanh, L. (2005). Estimating components of design effects for use in sample design, *Household Sample Surveys in Developing and*

- Transition Countries*, United Nations, New York, pp. 95–121.
- [10] Pettersson, H. & Nascimento Silva, P.L. (2005). Analysis of design effects for surveys in developing countries, *Household Sample Surveys in Developing and Transition Countries*, United Nations, New York, pp. 123–143.
- [11] Brogan, D. (2005). Sampling error estimation for survey data, *Household Sample Surveys in Developing and Transition Countries*, United Nations, New York, pp. 447–490.
- [12] Nathan, G. (2005). More advanced approaches to the analysis of survey data, *Household Sample Surveys in Developing and Transition Countries*, United Nations, New York, pp. 419–445.

IRENA OGRAJENŠEK

Survey Error

Introduction

In survey samples, the primary objective usually is to estimate unknown population quantities, such as population means or totals. The estimates invariably do not exactly equal the population values, i.e., there are survey errors. This chapter describes typical sources of survey errors. The context is estimating a population quantity with a randomly selected sample (*see Internal and External Quality Measures; Sampling Designs in Surveys; Adaptive and Spatial Sampling Designs*) of a population.

Survey Errors

Survey errors can be classified in two categories, non-sampling and sampling errors. Nonsampling errors arise due to unplanned mistakes in the **data collection** process, whereas sampling errors arise due to the variation from measuring only a subset of the population. Often nonsampling errors are difficult to detect and quantify. In contrast, sampling errors can be controlled and quantified. Indeed, many published surveys report results only in terms of estimates and sampling errors, without assessments of nonsampling errors (*see Survey Quality and Survey Ethics*). The magnitude of survey errors is affected by the amount of resources available, deadlines for releasing data products, and ethical issues in data collection and dissemination. Although important, these issues are not discussed here.

When describing the effects of survey error, survey methodologists use the terms *bias* and *variance*. To fix ideas, let θ be the population quantity of interest, and let $\hat{\theta}$ be the estimator of θ . The bias of $\hat{\theta}$ is $\text{Bias}(\hat{\theta}) = E(\hat{\theta}) - \theta$, where $E(\hat{\theta})$ is the expected value of $\hat{\theta}$. The variance of $\hat{\theta}$ is $E((\hat{\theta} - E(\hat{\theta}))^2)$. Survey samplers seek to design data collection systems that yield estimates with near-zero bias and small variance. In general, nonsampling errors increase bias, whereas sampling errors increase variance. The combined effect of both is called the *total survey error*, often represented by the mean square error, $MSE(\hat{\theta}) = E((\hat{\theta} - \theta)^2) = \text{Bias}^2(\hat{\theta}) + \text{Variance}(\hat{\theta})$.

We now describe the main sources of survey error. As a framework, we use the categorization

of survey errors presented by Biemer and Lyberg [1]. These include specification error, measurement error, coverage or frame error, mode effect error, nonresponse error, and process error. This is a typical categorization in the research literature; for example, similar errors are described by Weisberg [2].

Specification Error

Specification error arises when the survey does not address the main concepts of interest, instead providing information on different concepts. For example, the organization funding the survey might want to learn about the total income available to families, but the questionnaire asks families about total income from employment. This could render the results useless or even misleading for the concept of interest. Specification error generally results from poor communication between those who define the goals of the survey and those who design the questionnaire. Specification errors can be avoided when both parties carefully review the survey goals and questionnaire.

Measurement Error

In any survey, it is crucial that the collected data actually measure what is intended. This is challenging when surveys require responses from people. First, respondents may interpret words and questions differently than intended by the survey designer. If so, respondents could be answering effectively different questions; this complicates interpretations. Second, respondents may not be able to answer questions because they do not understand them or do not have the necessary information. The former problem arises with technical questions, and the latter with questions about events in the past, which people may remember imprecisely. Third, respondents may be unwilling to answer questions accurately, because they fear their answers will appear socially undesirable.

Determining the extent and impact of measurement error is difficult. The best approach to dealing with measurement error is to tackle the problem at the beginning of the survey process. For example, survey organizers can convene focus groups of potential respondents to evaluate whether people interpret and are able to answer the questions as intended, and alter question wording appropriately based on the feedback from these groups (*see Design and Testing of Questionnaires*). An extensive literature exists for

2 Survey Error

designing questions that generate reliable information. See, for example [1–6].

Coverage Error

The goal of random sampling is to select units that, after suitable weighting, represent the target population of interest. This breaks down when the sampling frame, i.e. the list from which the random sample is selected, does not include all units in the target population. When units missing from the frame have different characteristics than the rest of the target population, survey estimates may be biased. For example, suppose a target population is all people in a city. Sampling phone numbers listed in the telephone book excludes people with unlisted numbers or without telephones. When these people differ from people in the phone book on important characteristics, the sample does not represent the entire population and survey estimates are biased.

The sampling frame might be deficient because of the units it includes rather than excludes. Some units on the frame might not be eligible for the survey. When used for **estimation**, these units might bias results. Some units might be on the frame multiple times, which gives them a higher chance of being selected than other units. Without correction for differential probabilities of inclusion, estimates can be biased.

The key to avoiding coverage error is to use an accurate sampling frame with a one-to-one correspondence with units in the population. Such frames can be expensive to develop or verify. An alternative strategy is to put together multiple frames, without necessarily removing duplicates, in hopes of covering the target population. Using multiple frames requires adjusting the survey weights. One option for checking the frame for undercoverage is to obtain an external file with records that should be in the population and the frame, and estimate the fraction on the file that are missing from the frame. See [1, 2] for more discussion of coverage bias.

Mode Effect Error

Typically **survey data** on individuals are collected through mail or e-mail questionnaires, telephone interviews, or personal interviews. Face-to-face or telephone interviews can influence responses. For example, interviewers might present themselves or

ask questions in ways that lead respondents to be differentially accurate or truthful in their answers. In mail surveys, respondents may have difficulty following the instructions of the questionnaire, for example not adhering to complicated instructions for skipping questions.

Mode effect error is best reduced in the survey design stage. Many survey organizations reduce interviewer effects by training them not to deviate from standardized protocols. Consistency is further improved by computer-aided personal/telephone interviewing (*see Data Management in Survey Sampling*), especially for complicated questionnaires. See [1, 2] for examples of and strategies for dealing with mode effect error.

Nonresponse Error

Ideally, all measurements are recorded for every unit in the sample. Often this is not the case. Some units may not respond to the survey at all, and others may respond only to some items. When the respondents' characteristics differ from those of the nonrespondents, estimates could be biased. For example, people who are satisfied with a product might not respond, whereas those who are unsatisfied might respond. Even without such differences, nonresponse increases variances because estimates are based on fewer observed values than originally intended.

There are many strategies for handling nonresponse. Strategies for reducing its impact at the design stage include, among others, developing clear and to-the-point questionnaires, providing monetary or other incentives for survey completion, providing assurance of confidentiality, and planning for follow-up with nonrespondents (*see Processing of Survey Data*). After data collection, the estimates can be adjusted for nonresponse. Common approaches include modifying the survey weights and **imputation** of missing values. For details on methods for handling nonresponse, see [7].

Process Errors

The final stage of data collection involves turning the raw data into usable files. This process usually involves record checks and edits, recodes, or categorizations of responses, and creation of survey weights (*see Processing of Survey Data*). All these

steps are subject to error, which may be difficult to find when data processing is done with automated procedures. The literature on process errors is not extensive. One recommendation is to establish **quality management** practices (*see Survey Quality and Survey Ethics*) to ensure minimal errors in data processing.

Sampling Errors

Every sample survey is subject to sampling error, since different samples comprise different records. Fortunately, when selection is done randomly, survey analysts can estimate sampling errors with variances. The most common method for doing so is the design-based paradigm, in which estimates of variance are computed using mathematical formulas induced by the probabilities of selection (*see Statistical Analysis of Survey Data*). For example, in a simple random sample with no nonresponse, the estimated variance of the sample mean, $\bar{y}$, is $(1 - n/N)s^2/n$, where n is the **sample size**, N is the population size, and s^2 is the sample variance of the values of y in the data. There are many texts describing sampling error estimates, one classic being [8].

Concluding Remarks

Estimating the total impact of nonsampling and sampling errors is challenging. Unless a gold standard

file is available, in which nonsampling errors are known to be small, there are no standard approaches for estimating the total survey error. Developing such approaches is an area of ongoing statistical research.

References

- [1] Biemer, P. & Lyberg, L. (2003). *Introduction to Survey Quality*, John Wiley & Sons, Hoboken.
- [2] Weisberg, H.F. (2005). *The Total Survey Error Approach. A Guide to the New Science of Survey Research*, The University of Chicago Press, Chicago.
- [3] Groves, R. (1989). *Survey Errors and Survey Costs*, John Wiley & Sons, New York.
- [4] Fowler, F. (1995). *Improving Survey Questions*, Sage, Thousand Oaks.
- [5] Czaja, R. & Blair, J. (1996). *Designing Surveys*, Pine Forge Press, Thousand Oaks.
- [6] Sirken, M., Herrman, D., Schechter, S., Schwarz, N., Tanur, J. & Tourangeau, R. (1999). *Cognition and Survey Research*, John Wiley & Sons, New York.
- [7] Groves, R., Dillman, D., Eltinge, J. & Little, R.J.A. (2001). *Survey Nonresponse*, John Wiley & Sons, New York.
- [8] Cochran, W. (1977). *Sampling Techniques*, John Wiley & Sons, New York.

JEROME P. REITER

Survey Quality and Survey Ethics

Introduction

Various producers of statistics have devoted considerable efforts to improve the quality of their products. Traditionally, surveys are annotated with quality descriptions providing the clients with support to interpret the results. Furthermore, there can be studies on the quality of the entire sample survey, called *quality reports*. The clients often need more profound information on the organization, applied procedures and processes as well as the products. Thus, the approach is changing from a mere focus on product quality to total **quality management** (TQM).

This article describes various perspectives of survey quality, starting from traditional reporting and ending up with recent more comprehensive quality approaches.

Survey Errors and Quality

One of the main issues in ensuring good quality in surveys is to avoid different types of errors which are discussed by **Survey Error**. Survey organizations are very keen in avoiding errors in all survey processes to assure better quality. However, since the errors are inevitable, their magnitude and effects should be evaluated and reported. Estimated standard errors for parameter estimates are probably the oldest quality measures for sample surveys. Standard errors and other estimates based thereupon, like coefficients of variation or **confidence intervals**, have been calculated from statistical samples since the rise of survey methodology in the 1940s and 1950s.

Besides the sampling errors some authors pointed out, quite early, the roots of other types of errors. For example, Mahalanobis [1] gives subtle examples of various error sources in statistical work process. Similarly, the errors in coverage and those created by nonresponse were taken into consideration. **Survey error** concepts, ending up with the most comprehensive one, the total survey error containing both the sampling and nonsampling errors, are analyzed in various textbooks (see e.g., [2]). All sources of error except the basic sampling error can induce systematic

errors which may distort the results. The magnitude of systematic errors should therefore be analyzed and reported using proper statistical methods (see e.g., [3, 4]) even though a profound mean square error (MSE) **estimation** may not be amenable.

The cost aspect is important. Costs affect all stages of survey process: sampling design, planning and preparation of material needed for a survey, **data collection** methods, communication with respondents, data editing and cleaning, as well as analyzing and reporting, see [2, 5]. Survey quality improvements are always subject to cost constraints.

Reporting of Quality

The first official recommendation of presenting survey quality aspects was published by the United Nations [6]. The UN papers laid down much of the basic ideas applied even today. A fundamental requirement is that all characteristics having an impact on **data quality** should be known to the users.

Types of Reports

There are generally three types of reports describing survey quality. The first type are *comprehensive reports* which contain all aspects of a survey starting from the background information, describing all processes, errors and ending in a comparison with other data sources. The US Federal Committee on Statistical Methodology [7] calls them *analytical reports* and the preparation of such a report requires investigation of the whole process. They are often not targeted at laymen, but rather at researchers and professional peers.

The second type are *brief reports* (or quality descriptions) often included as a separate chapter or box in the publication. These are short and meant to give the reader an overall view of the quality, with some key indicators and recommendations to users.

The third type, *quality profiles*, are understood to present a view on how different quality aspects change with time. For example, Eurostat gives short quality profiles for a vast set of structural indicators and indices, and grades their quality by three categories [8]. On the other hand, the National Center for Education Statistics presents a profound quality profile using the total survey error model of three waves of large school surveys in the United States [9].

In general, **time series** are a useful method of giving information on the development of various quality aspects, for example, nonresponse rates.

The quality reports should always indicate how the errors have been adjusted and give advice to the users on possible limitations when using the data.

Quality Indicators

Many proposals have been put forward for a set of indicators to be provided when reporting quality. Preferably, the indicators should cover all relevant aspects of quality. The coverage errors should be evaluated, standard errors and other measures of precision of parameter estimates should always be included. Unit nonresponse rate(s) and, if possible, the effect of unit nonresponse on certain key parameters should be investigated. Proposals for a standardized way of informing about unit nonresponse exist (see e.g., [10, 11]), but so far there is no generally accepted form of reporting, except in certain internationally coordinated surveys. However, the situation is much more unsatisfactory concerning the other indicators of quality. Eurostat has its own standard set of indicators [12] and the US Federal Committee proposes some others [7].

Metadata and Documentation

The increasing dissemination through internet creates a challenge to all survey organizations. The users should be provided with similar quality information, which is embedded directly into the statistical tables and databases. Metadata and documentation systems store and provide information on the various aspects of quality: statistics production and outputs as well as identification of areas for improvement.

There are metadata standards like the special data dissemination standard (SDDS) [13], or statistical data and metadata exchange (SDMX) [14] which partly fill the need. Although the standards are developing fast, there are many issues to be solved before the metadata systems are so complete that a user may explore all quality and process issues as deeply as is needed. Therefore, new inventions are to be expected in standardizing statistical metadata.

Professional Ethics Related to the Quality Issues

Ethical Codes

Professional ethics is a very important constituent with respect to quality of surveys and statistics, in general. It sets the general guidelines for the activity of each statistician or researcher. On this issue, we refer to the code of the International Statistical Institute [15]. It consists of the following topics:

1. obligations to the society
2. obligations to funders and employers
3. obligations to colleagues
4. obligations to subjects

with a number of references. Many agencies in the field of official statistics have created their own codes. Besides the code for professionals, we must mention the guidelines for authorities. Here, the main reference is the United Nations Fundamental Principles for Official Statistics [16].

Marketing and Opinion Research

The marketing research communities have focused on quality and ethical issues for a long time. The first international code dates back to the late 1940s, and the current code was accepted jointly by the International Chamber of Commerce and World Association of Research Professionals [17]. The code focuses especially on the rights of customers to obtain impartial and objective results, addressing also the rights of respondents, professional responsibilities, and the mutual requirements of survey organization and their clients. Similar codes of professional ethics exist in many market research organizations and associations, e.g., [18].

Total Quality Management

Recently, many survey organizations have introduced a holistic approach to quality, often in the form of TQM. It is specifically visible in official statistics, where either national or international requirements have put pressure to accept some form of systematic and standardized practices, both in work processes and output. As a consequence, the organizations have given declarations to promote quality (e.g.,

[19]), accepted quality criteria in the form of policy, guidelines, current best methods, and so on [20], and established a program to realize those policies [21]. Examples of those approaches include the European Statistics Code of Practice [22] or the Data Quality Assessment Framework of the International Monetary Fund (IMF) [23].

TQM is based on dividing the quality concept into three main categories: operational environment, work processes and outputs. Each main category contains many different things to be assessed and measured. Finally, the organization must have a proper way to analyze feedback from customers and other sources, and to learn from the experiences. Introducing a TQM-based quality management is a time-consuming process, and it will often change the whole organization. In addition, it requires good documentation of all phases of processes, self-assessments, as well as internal and external auditing procedures.

In spite of apparently similar main headings, the realization of the policy can vary substantially. As an example, the comparison of the European Statistical Code of Practice and the IMF Data Quality Assessment Framework reveals that they share many common items, but differ in some others [24].

The renowned international standard of the ISO 9000 family [25] is not very common among survey organizations, rather national certification standards have been applied in market research. A recent quality standard ISO 20252 for market, opinion, and social research [25] is intended to overcome that problem. Although it is mainly devoted to the market research community, the systematic approach can also be applied to academic research and official statistics. The main titles of the 20252 standard include *Quality management system requirements*, *Managing the executive elements of research*, *Data collection*, *Data management and processing*, and *Report on research projects*. The standard is still in the introductory phase and the certification procedures are to be developed.

References

- [1] Mahalanobis, P.C. (1946). Recent experiments in statistical sampling in the Indian Statistical Institute, *Journal of the Royal Statistical Society* **109**, 325–337.
- [2] Groves, R.M. (1989). *Survey Errors and Survey Costs*, John Wiley & Sons, New York.
- [3] Lessler, J.T. & Kalsbeek, W.D. (1992). *Nonsampling Errors in Surveys*, John Wiley & Sons, New York.
- [4] Biemer, P., Bollen, K. & Vernunt, J. (2005). Recent research in the evaluation of errors in surveys, in *Proceedings of the 55th Session of the International Statistical Institute*, Sydney 2005, ISI CD-ROM, 2006.
- [5] Biemer, P.P. & Lyberg, L.E. (2003). *Introduction to Survey Quality*, John Wiley & Sons, Hoboken.
- [6] United Nations. (1964). *Recommendations for the Preparation of Sampling Survey Reports (Provisional Issue)*, Statistical Papers, Series C No. 1 Rev. 2. New York.
- [7] United States. (2001). *Measuring and Reporting Sources of Error in Surveys*, Statistical Policy Working Paper, The Federal Committee on Statistical Methodology, New York.
- [8] Eurostat. (2007). *Quality Profiles for Structural Indicators*, <http://www.ec.europa.eu/eurostat/quality> - accessed 13082007.
- [9] US Department of Education, National Center for Education Statistics. (2000). *Quality Profile for SASS Rounds 1–3: 1987–1995*, NCES 2000–308, Washington, DC.
- [10] Lynn, P., Beerten, R., Laiho, J. & Martin, J. (2003). *Towards standardisation of survey outcome categories and response rate calculations*, *Research in Official Statistics*, No. 1/2002, **5**, 61–84.
- [11] The American Association for Public Opinion Research. (2006). *Standard Definitions. Final Dispositions of Case Codes and Outcome Rates for Surveys*, 4th Edition, AAPOR, Lenexa.
- [12] Eurostat. (2005). *Standard Quality Indicators*, <http://www.ec.europa.eu/eurostat/quality> (accessed 13082007).
- [13] International Monetary Fund. *Special Data Dissemination Standard*, <http://dsbb.imf.org/Applications/web/sddshome/> (accessed 13082007).
- [14] SDMX. *Statistical Data and Metadata Exchange*, <http://www.sdmx.org> (accessed 13082007).
- [15] International Statistical Institute. (1986). Declaration on professional ethics, *International Statistical Review*, **54**(2), 227–242.
- [16] United Nations Economic and Social Council. (1994). *Report of the Special Session of the Statistical Commission*, E/1994/29, New York.
- [17] International Chamber of Commerce, ESOMAR. (2005). *ICC/ESOMAR International Code of Marketing and Social Research Practice*, http://www.esomar.org/uploads/pdf/ps_cg_icccode.pdf (accessed 13082007).
- [18] World Association of Public Opinion Research WAPOR. (2005). *Code of Professional Ethics and Practices*, <http://www.unl.edu/wapor/ethics.html> (accessed 13082007).
- [19] Eurostat. (2001). *Quality Declaration of the European Statistical System*, <http://www.ec.europa.eu/eurostat/quality> (accessed 13082007).
- [20] Colledge, M. & March, M. (1997). Quality policies, standards, guidelines, and recommended practices at statistical agencies in *Survey Measurement and Process Quality*, L. Lyberg, P. Biemer, M. Collins, E. de Leeuw, C. Dippo, N. Schwarz & D. Trewin, eds, John Wiley & Sons, New York, pp. 501–522.

4 Survey Quality and Survey Ethics

- [21] Trewin, D. (2002). The importance of a quality culture, *Survey Methodology* **28**(2), 125–133.
- [22] Eurostat. (2005). *European Statistics Code of Practice*, <http://www.ec.europa.eu/eurostat/quality> (accessed 13082007).
- [23] International Monetary Fund. (2003). *Data Quality Assessment Framework*, <http://dsbb.imf.org/Applications/web/dqrs/dqrsdqaf/> (accessed 13082007).
- [24] Laliberté, L. & Defays, D. (2006). Comparison of the IMF's data quality assessment framework (DQAF) and European statistical system quality approaches? An update. Paper presented at the *European Conference on Quality in Survey Statistics Q2006*, Cardiff UK, 24-26 April 2006. http://www.statistics.gov.uk/events/q2006/downloads/W15_Laliberté.doc (accessed 13082007).
- [25] International Standardization Organization. ISO 20252: 2006, Quality management systems—Market, opinion and social research—Vocabulary and service requirements. <http://www.iso.org/iso/en/ISOOnline.frontpage> (accessed 13082007).

KARI DJERF

SERVQUAL Surveys

Explanation of SERVQUAL

Service quality is often analyzed using rating scales to assess individual service attributes. The most widely used and most carefully scrutinized rating scale is SERVQUAL (the name stands for SERVICE QUALity).

The scale is a result of a systematic ongoing study of service quality that begun in 1983 [1–5]. The model defines quality as the difference between customer perceptions and expectations with regard to quality of delivered service. The respondents are asked to answer two sets of questions dealing with the same subject, one set at a general level (e.g., quality of service in financial institutions), and one for a company in question (e.g., quality of service in bank XYZ). The first (general) set of questions are quality expectations (E_i), the second (specific) set are quality perceptions (P_i) as shown in Table 1.

Respondents choose from a modified Likert scale which measures intensity of agreement or disagreement with any given statement. Response options range from “strongly agree” and lesser levels of agreement to “neutral” (neither agree nor disagree) and lesser levels of disagreement to “strongly disagree”. It has to be pointed out that SERVQUAL is based on the assumption of the Likert scale’s cardinal nature. In real life, however, one person’s “complete satisfaction” might be less than another’s “partial satisfaction” (for a more in-depth discussion of Likert scale related critique of SERVQUAL see the section titled “Reexamination and Criticism of SERVQUAL”).

For each of the 22 items (service attributes), a quality judgment can be computed according to the following equation:

$$\text{Perception}(P_i) - \text{Expectation}(E_i) = \text{Quality}(Q_i) \quad (1)$$

Quality is poor if expectations exceed perceptions and vice versa. On the other hand, if low customer expectations are met, quality does exist.

The SERVQUAL score (perceived service quality) is obtained by the following equation:

$$Q = \frac{1}{22} \sum_{i=1}^{22} (P_i - E_i) \quad (2)$$

The $P_i - E_i$ gap scores can be subjected to an iterative sequence of item-to-item **correlation** analyses, followed by a series of factor analyses to examine the dimensionality of the scale using the oblique rotation, which identifies the extent to which the extracted factors are correlated.

The original SERVQUAL scale had 97 items that resulted in ten service quality dimensions. The presently established 22-item scale shown in Table 1 has five quality dimensions defined as follows:

- *tangibles*: physical facilities, equipment, and appearance of personnel;
- *reliability*: ability to perform the promised service dependably and accurately;
- *responsiveness*: willingness to help customers and provide prompt service;
- *assurance*: knowledge and courtesy of employees as well as their ability to convey trust and confidence;
- *empathy*: individual care and attention that company provides its customers.

However, as discussed in the section titled “Reexamination and Criticism of SERVQUAL”, empirically identified five-factor structure cannot be found in all service industries.

Reexamination and Criticism of SERVQUAL

Although SERVQUAL has been widely used in business-to-business and business-to-customer settings, this does not mean the scale has not been subject to reexamination and criticism. The main objections to SERVQUAL pertain to the object of measurement, length of the questionnaire, timing of questionnaire administration, use of the Likert scale, use of ($P_i - E_i$) difference scores, generalization of service quality dimensions, and the static nature of the model. Let us take a closer look at each of them.

2 SERVQUAL Surveys

Table 1 An overview of SERVQUAL dimensions and expectation as well as perception items

Quality dimension	Expectations (E_i)	Perceptions (P_i)
Tangibles	Excellent companies will have modern-looking equipment. The physical facilities at excellent companies will be visually appealing. Employees of excellent companies will be neat in appearance. Materials associated with the service (such as pamphlets or statements) will be visually appealing in an excellent company.	XYZ has modern-looking equipment. XYZ's physical facilities are visually appealing. XYZ's employees are neat in appearance. Materials associated with the service (such as pamphlets or statements) are visually appealing at XYZ.
Reliability	When excellent companies promise to do something by a certain time, they will do so. When customers have a problem, excellent companies will show a sincere interest in solving it. Excellent companies will perform the service right the first time. Excellent companies will provide their services at the time they promise to do so. Excellent companies will insist on error-free records.	When XYZ promises to do something by a certain time, it does so. When you have a problem, XYZ shows a sincere interest in solving it. XYZ performs its service right the first time. XYZ provides its services at the time it promises to do so. XYZ insists on error-free records.
Responsiveness	Employees of excellent companies will tell customers exactly when services will be performed. Employees of excellent companies will give prompt service to customers. Employees of excellent companies will always be willing to help customers. Employees of excellent companies will never be too busy to respond to customer requests.	Employees of XYZ tell you exactly when the service will be performed. Employees of XYZ give you prompt service. Employees of XYZ are always willing to help you. Employees of XYZ are never too busy to respond to your requests.
Assurance	The behavior of employees of excellent companies will instill confidence in customers. Customers of excellent companies will feel safe in their transactions. Employees of excellent companies will be consistently courteous with customers. Employees of excellent companies will have the knowledge to answer customer questions.	The behavior of XYZ's employees instills confidence in you. You feel safe in your transactions with XYZ. Employees of XYZ are consistently courteous with you. Employees of XYZ have the knowledge to answer your questions.
Empathy	Excellent companies will give customers individual attention. Excellent companies will have operating hours convenient to all their customers. Excellent companies will have employees who give customers personal attention. Excellent companies will have the customers' best interests at heart. The employees of excellent companies will understand the specific needs of their customers.	XYZ gives you individual attention. XYZ has operating hours convenient to you. XYZ has employees who give you personal attention. XYZ has your best interests at heart. Employees of XYZ understand your specific needs.

Adapted after A. Parasuraman, L.A. Berry and V.A. Zeithaml, "Refinement and Reassessment of the SERVQUAL Scale", 1991, pp. 446–449.

Object of measurement.

It is not clear whether the scale measures service quality or customer satisfaction [6].

Length of the questionnaire.

The SERVQUAL questionnaire is too long. It could be shortened by elimination of expectation scores (see discussion under *use of $(P_i - E_i)$ difference scores* for more details), elimination of certain items (those without the clear mode) and/or fusion of the interrelated dimensions of *reliability*, *responsiveness* and *assurance* into one dimension called *task-related receptiveness* [7, 8].

Timing of questionnaire administration.

The main issue here is whether to distribute the questionnaire before or after the service experience [9, 10]. In other words, should expectations be solicited before the service experience, or away from the actual point of service delivery and unrelated to an encounter? Some researchers compromise by collecting their data after the service experience at the actual point of service delivery [7]. Consequently, they fear that this might have loaded their results toward performance, while those of other researchers might have been loaded toward expectations.

Use of the Likert scale.

Several issues are problematic here. Firstly, the structure of the Likert scale should be noted. SERVQUAL authors use a seven-point scale, while in many replication studies a five-point scale is adopted to increase response rate and response quality.

Secondly, the already mentioned fact that SERVQUAL is based on the assumption of the Likert scale's cardinal nature should not be overseen. Furthermore, once the respondents have marked the extreme point and want to express an even stronger opinion on the next item, this cannot be reflected in the answer anymore, since the maximum score has already been given [11]. Unfortunately, the researcher has no tool to assess the levels of interrespondent reliability, agreement, and consensus.

Finally, from the viewpoint of statistical analysis, some authors argue that none of these problems matter as long as the answers are normally distributed [6].

However, in practice, the majority of SERVQUAL surveys tend to result in highly skewed customer responses. Hence, an average rating based on the arithmetic mean of the customer responses is likely to be a poor measure of central tendency, and may not be the best indicator of service quality.

Use of $(P_i - E_i)$ difference scores.

Ambiguous definition of expectations in the SERVQUAL model seems to be a major problem. Increasing $P_i - E_i$ scores may not always correspond to increasing levels of perceived quality which impugnes theoretical validity of the SERVQUAL's perceived quality framework [12–14].

Generalization of service quality dimensions.

An empirically identified five-factor structure cannot be found in all service industries. Only the existence of the *tangibles* dimension is confirmed in all replication studies, in which the number of distinct service quality dimensions otherwise varies from one to nine [15]. It seems that the dimensionality of service quality might be determined by the type of service a researcher deals with [16].

The static nature of the model.

There exists a number of long-term service processes (such as education) where both perceptions and expectations (and consequently quality evaluations) change in time. For these service processes, a dynamic model of service quality should be developed [17].

Generalizations of and Alternatives to SERVQUAL

In order to assess its general applicability, the SERVQUAL scale has been replicated in many different service industries both in business-to-business and business-to-customer markets.

Many of the replication studies question the value and purpose of two separate datasets (perceptions and expectations). Some researchers suggest that it might be better not to use difference scores since the factor structure of the answers given to the questions about expectations and perceptions, and the resulting difference scores, are not always identical [18–21]. Additionally, research

shows that performance perception scores alone give a good quality indication. Therefore, a modified SERVQUAL scale using only performance perceptions to measure service quality (called SERVPERF, which stands for SERVICE PERFORMANCE) has been proposed [20].

A three-component measurement model including perceptions, expectations, and perceived importance ratings, or weights of service attributes has also been suggested [10]. A rationale for this model is the following: while customers might expect excellent service on each attribute, all of the attributes involved may not be equally important to them. However, it seems that the inclusion of perceived importance ratings does not substantially improve the explanatory **power** of the model.

An attempt to operationalize service quality as a multiplicative linear compensatory function of performance and importance ratings is also described in the literature [22]. However, when practically applied, this model does not surpass SERVQUAL in performance.

The only approach that, to date, seems to outperform SERVQUAL is SERVPERF, the performance-only measure of service quality.

References

- [1] Parasuraman, A., Zeithaml, V.A. & Berry, L.L. (1985). A conceptual model of service quality and its implications for future research, *Journal of Marketing* **4**, 41–50.
- [2] Parasuraman, A., Zeithaml, V.A. & Berry, L.L. (1988). SERVQUAL: a multiple-item scale for measuring customer perceptions of service quality, *Journal of Retailing* **1**, 12–40.
- [3] Parasuraman, A., Berry, L.L. & Zeithaml, V.A. (1991). Refinement and reassessment of the SERVQUAL scale, *Journal of Retailing* **4**, 420–450.
- [4] Parasuraman, A., Berry, L.L. & Zeithaml, V.A. (1991a). Understanding, measuring and improving service quality: findings from a multiphase research program, in *Service Quality: Multidisciplinary and Multinational Perspectives*, S.W. Brown, E. Gummeson, B. Edvardsson & B. Gustavsson, eds, Lexington Books, Lexington, pp. 253–268.
- [5] Parasuraman, A., Zeithaml, V.A. & Berry, L.L. (1994). Reassessment of expectations as a comparison standard in measuring service quality: implications for further research, *Journal of Marketing* **1**, 111–124.
- [6] Kasper, H., van Helsdingen, P. & de Vries Jr, W. (1999). Services marketing management, *An International Perspective*, John Wiley & Sons, Chichester.
- [7] Johns, N. & Tyas, P. (1996). Use of service quality gap theory to differentiate between food service outlets, *The Service Industries Journal* **3**, 321–346.
- [8] Llosa, S., Chandon, J. & Orsingher, C. (1998). An empirical study of SERVQUAL's dimensionality, *The Service Industries Journal* **2**, 16–44.
- [9] Smith, A.M. (1995). Measuring service quality: Is SERVQUAL now redundant? *Journal of Marketing Management* **1–3**, 257–276.
- [10] Carman, J.M. (1990). Consumer perceptions of service quality: an assessment of the SERVQUAL dimensions, *Journal of Retailing* **2**, 33–55.
- [11] Chidester, J. (1995). Tailoring your survey, *Credit Union Management* **4**, 30–31.
- [12] Teas, R.K. (1993). Consumer expectations and the measurement of perceived service quality, *Journal of Professional Services Marketing* **2**, 33–54.
- [13] Teas, R.K. (1993a). Expectations, performance evaluation, and consumers' perceptions of quality, *Journal of Marketing* **4**, 18–34.
- [14] Teas, R.K. (1994). Expectations as a comparison standard in measuring service quality: an assessment of a reassessment, *Journal of Marketing* **1**, 132–139.
- [15] Buttle, F.A. (1995). What future for SERVQUAL? in *Proceedings of the 24th European Marketing Academy Conference*, M. Bergadaà, ed, ESSEC, Cergy-Pontoise, pp. 211–230.
- [16] Babakus, E. & Boller, G.W. (1992). An empirical assessment of the SERVQUAL scale, *Journal of Business Research* **3**, 253–268.
- [17] Haller, S. (1998). Beurteilung von dienstleistungsqualität, in *Dynamische Betrachtung des Qualitätsurteils im Weiterbildungsbereich*, Deutscher Universitäts-Verlag, Gabler Verlag, Wiesbaden.
- [18] Brown, T.J., Churchill, G.A. & Peter, J.P. (1993). Improving the measurement of service quality, *Journal of Retailing* **1**, 127–139.
- [19] Cronin Jr, J.J. & Taylor, S.A. (1992). Measuring service quality: a reexamination and extension, *Journal of Marketing* **3**, 55–68.
- [20] Cronin Jr, J.J. & Taylor, S.A. (1994). SERVPERF versus SERVQUAL: reconciling performance-based and perceptions-minus-expectations measurement of service quality, *Journal of Marketing* **1**, 125–131.
- [21] Taylor, S.A., Sharland, A., Cronin, J.J. & Bullard, W. (1993). Recreational service quality in the international setting, *International Journal of Service Industry Management* **4**, 68–86.
- [22] Babakus, E. & Inhofe, M. (1993). Measuring perceived service quality as a multi-attribute attitude, *Developments in Marketing Science*, AMS, Coral Gables, pp. 376–380.

IRENA OGRAJENŠEK

Assignable Cause

In process control and improvement, it is effective to discern the variation of process characteristics into variation due to *chance causes* (noise) and due to *assignable causes* (signal). Assignable causes (called *special* causes by W.E. Deming, and *sporadic events* by J.M. Juran) leave traces in the data in

the form of nonrandom patterns, which enables their identification. Chance causes are a collection of multiple influences, none of which dominate the others, whose cumulative effect is often conceptualized as normal uncorrelated variation. Chance causes are called *common causes* by Deming, and *chronic* by Juran. Assignable causes are often regarded as exogenous to the system, whereas chance causes are considered endogenous.

Autocorrelation Function

Introduction

In quality and reliability engineering studies, time-series observations (*see Time Series Analysis*) are often statistically correlated. A typical example is a set of observations taken on a process variable (like temperature) in regular time intervals, where the (fixed) elapsed time between consecutive observations is relatively short. Due to the natural inertia of the process, observations that are close together will probably tend to locally covary (*see Covariance*), even when both the process mean and the process variance do not vary over time (the process is stationary). In this article it is assumed that the time series is stationary (*see Time Series Analysis*). The covariation between pairs of observations in the time series that are equally distant from one another may be captured by their typical autocorrelation (this is the usual Pearson correlation (*see Correlation*); since it is calculated between observations belonging to the same sample of observations, it is denoted as *autocorrelation*). The autocorrelation between the observation at time period t and the one at time period $t + k$ is named *autocorrelation at lag k* , and is denoted ρ_k (due to stationarity, ρ_k is independent of t). The autocorrelation function displays the autocorrelation as a function of k . It plays a central role in time-series modeling. In the section titled “Definitions and Some Formulae” we define the autocorrelation function and the partial autocorrelation function. In the section titled “Estimation”, estimation procedures are discussed.

Definitions and Some Formulae

The Autocorrelation Function

Let z_t denote an observation taken on a stationary process at time t . Let the mean be $E(z_t) = \mu$, and denote the variance by $E[(z_t - \mu)^2] = \sigma_z^2$ (due to the assumed stationarity, both the mean and the variance do not depend on t). The covariance (*see Covariance*) between z_t and z_{t+k} (namely, between any two process values that are k time intervals

apart) is called *autocovariance at lag k* , and it is calculated by

$$\gamma_k = \text{cov}[z_t, z_{t+k}] = E[(z_t - \mu)(z_{t+k} - \mu)] \quad (1)$$

The *autocorrelation at lag k* is the correlation between z_t and z_{t+k} :

$$\rho_k = \frac{E[(z_t - \mu)(z_{t+k} - \mu)]}{\sqrt{E[(z_t - \mu)^2]E[(z_{t+k} - \mu)^2]}} = \frac{\gamma_k}{\gamma_0} = \frac{\gamma_k}{\sigma_z^2} \quad (2)$$

(due to stationarity, the variances of z_t and z_{t+k} are the same, σ_z^2). Obviously, $\rho_0 = 1$.

The expression in equation (2) does not express ρ_k as an explicit function of k . If such a function exists, it will be denoted as *the autocorrelation function* of the process. For example, for a first-order autoregressive model (*see Time Series Analysis*), the autocorrelation function is $\rho_k = \phi_1^k$ ($k \geq 0$), where ϕ_1 is a certain parameter. If no explicit function in k is available, however one can calculate individual autocorrelations, then the *plot* of the autocorrelation coefficients, ρ_k , as a function of the lag, k , is called *the autocorrelation function*, and it is denoted by the set notation $\{\rho_k\}$. Since $\rho_k = \rho_{-k}$, the autocorrelation function is symmetrical around $k = 0$.

The Partial Autocorrelation Function (for AR(p) Models)

One category of time-series models is the group of autoregressive models (*see Time Series Analysis*). For an autoregressive process of order p , AR(p), it may be shown that the autocorrelation function may be given by the difference equation (refer to [1], p. 56)

$$\rho_k = \phi_1 \rho_{k-1} + \phi_2 \rho_{k-2} + \cdots + \phi_p \rho_{k-p}, \quad k > 0 \quad (3)$$

where $\{\phi_i\}$ are p parameters that need to be estimated. Substituting $k = 1, 2, \dots, p$ in equation (3), a set of p equations is obtained. These are called the *Yule-Walker equations*. Assume that the autocorrelations $\{\rho_1, \dots, \rho_p\}$ are known (or estimable from data). Then the p equations, derived from equation (3), may be used to obtain the p unknown parameters of the AR(p) model. If estimates of

2 Autocorrelation Function

the p autocorrelations are available, they may be inserted into the p Yule–Walker equations to obtain the Yule–Walker estimates of the parameters.

A major problem in building a time-series model is identifying the order of the model (analogously with the problem of determining the number of regressor variables in a multiple linear regression model). In autoregressive models, the identification problem is to determine p . Here, the *partial autocorrelation function* is useful. Suppose that we believe that the order of the model is m , and denote by ϕ_{mj} the j th coefficient in the Yule–Walker equations ($j = 1, 2, \dots, m$). Obviously, ϕ_{mm} will be the last coefficient. For the autocorrelation ρ_k we obtain, analogously with equation (3):

$$\begin{aligned} \rho_k &= \phi_{m1}\rho_{k-1} + \phi_{m2}\rho_{k-2} + \dots + \phi_{mm}\rho_{k-m}, \\ k &= 1, 2, \dots, m \end{aligned} \quad (4)$$

These Yule–Walker linear equations may be solved for the m coefficients $\{\phi_{mj}\}$ for successive values of m . The quantity ϕ_{mm} , regarded as a function of m , is called the *partial autocorrelation function*. Written in a matrix form, equation (4) may be written:

$$\begin{bmatrix} 1 & \rho_1 & \rho_2 & \dots & \rho_{m-1} \\ \rho_1 & 1 & \rho_1 & \dots & \rho_{m-2} \\ \cdot & & & & \\ \cdot & & & & \\ \cdot & & & & \\ \rho_{m-1} & \rho_{m-2} & \rho_{m-3} & \dots & 1 \end{bmatrix} \begin{bmatrix} \phi_{m1} \\ \phi_{m2} \\ \cdot \\ \cdot \\ \cdot \\ \phi_{mm} \end{bmatrix} = \begin{bmatrix} \rho_1 \\ \rho_2 \\ \cdot \\ \cdot \\ \cdot \\ \rho_m \end{bmatrix} \quad (5)$$

For example, from equation (5), one obtains by introducing successively $m = 1, 2$ the partial autocorrelations:

$$\phi_{11} = \rho_1; \quad \phi_{22} = \frac{\rho_2 - \rho_1^2}{1 - \rho_1^2} \quad (6)$$

A particular property of an autoregressive process of order p , AR(p), is that:

For $m \leq p$: ϕ_{mm} will be nonzero; For $m > p$: ϕ_{mm} will be zero.

Thus, calculating successively for $m = 1, 2, \dots, p, \dots$, one may find the correct order of an autoregressive model, given a sample of time-series values, by observing the values of the partial autocorrelation function. A sharp decline will be observed in the sample estimates of ϕ_{mm} for $m > p$.

Estimation

The Autocorrelation Function

The most satisfactory sample estimate, $\hat{\rho}_k$, of the autocorrelation, ρ_k , is (refer to [1], p. 31)

$$\hat{\rho}_k = \frac{\hat{\gamma}_k}{\hat{\gamma}_0} \quad (7)$$

where

$$\hat{\gamma}_k = \frac{1}{N} \sum_{t=1}^{N-k} (z_t - \bar{z})(z_{t+k} - \bar{z}), \quad k = 0, 1, 2, \dots, \quad (8)$$

is the sample estimate of the covariance, γ_k , $\bar{z}$ is the sample mean, and N is the **sample size**. Obviously, $\hat{\gamma}_0$ is the sample estimate of the variance.

If the theoretical shape of the autocorrelation function is known (like that for AR(1) introduced earlier), then estimates derived from equation (7) may be used to estimate the parameters of the assumed time-series model. Thus, for an AR(p) model, estimates of the autocorrelations may be inserted into the Yule–Walker equation (3) in order to derive estimates of the parameters $\{\phi_k\}$. For example, for AR(2), we obtain [1], p. 62:

$$\phi_1 = \frac{\rho_1(1 - \rho_2)}{1 - \rho_1^2}; \quad \phi_2 = \frac{\rho_2 - \rho_1^2}{1 - \rho_1^2} \quad (9)$$

The Partial Autocorrelation Function (for AR(p) Models)

Estimation of partial autocorrelations may be conducted by successively fitting autoregressive models (see **Time Series Analysis**) of an increasing order ($m = 1, 2, 3, \dots$), using **least squares**, and then picking up the estimates of the partial

autocorrelations $\{\hat{\phi}_{mm}\}$. These will comprise the last coefficients estimated from each model. Alternatively, inserting the estimated autocorrelations into equation (5) and solving for ϕ_{mm} would yield the required estimate of the partial autocorrelation for the $AR(m)$ model. As explained earlier, these successive partial autocorrelations ($m = 1, 2, \dots, p, \dots$) may help determine the order, p , of the $AR(p)$ model best suited for the time series.

Reference

- [1] Box, G.E.P., Jenkins, G.M. & Reinsel, G.C. (1994). *Time Series Analysis: Forecasting and Control*, 3rd Edition, Prentice Hall, Englewood Cliffs.

HAIM SHORE

Bias of an Estimator

The difference between the expected value (*see **Expectation***) of an estimator and the population parameter it is supposed to estimate is called its

bias. Statistical bias is due to the specific algebraic form of an estimator. Selection bias arises because of inappropriate sampling. Measurement bias arises from a measurement system that is not appropriately calibrated.

Coefficient of Determination (R^2)

The coefficient of determination (usually denoted R^2) is a concept in *analysis of variance* and regression analysis. It is a measure of the proportion of explained *variance* present in the data. Hence, the higher the value of R^2 , the better the model describes the data. When data consists of values $y_1, \dots, y_n$ of a response variable and a model with predictor variables is applied to the data, then R^2 is defined as

$$R^2 = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2} \quad (1)$$

where $\bar{y}$ denotes the average of the observations and $\hat{y}_i$ the prediction of y_i using the fitted model. In case there is no relation between the predictor variable(s) and the response variable, then $\bar{y}$ is the best “model” to explain the data. Hence, the terms $y_i - \bar{y}$ account for deviations in the worst case that there is no relation between the predictor variable(s) and the response variable. The range of values of R^2 depends on the type of model to be fitted; in standard cases like linear least-squares regression models they lie between 0 and 1.

In case of linear regression with only one predictor variable, R^2 equals the square of the *correlation* between the values of the predictor variable and the response variable. Several different formulas of R^2 are known, e.g., the formula

$$R_0^2 = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n y_i^2} \quad (2)$$

is used for linear regression models without intercept term. Using a form of R^2 that does not match the type of model may lead to wrong conclusion, since in general the different forms of R^2 are not equivalent (see [1] for an excellent overview).

Very low values of R^2 (<0.5) indicate a weak relation between the predictor variable(s) and the response variable, while moderately low values of R^2 (<0.8 and >0.5) indicate that the model is not adequate, e.g., terms are missing, or that there is substantial error variation (possibly caused by a large measurement spread). However, high values of R^2 do not necessarily imply that the model is adequate (see [1]). The major reason is that R^2 increases when models are enlarged, even when this is done with irrelevant terms. This effect is only partially taken away by using the adjusted form $R_{\text{adj}}^2 = 1 - \frac{n-1}{n-k-1} R^2$, when there are k predictor variables in a linear regression model (when there is no intercept, then the denominator should be $n - k$). A proper assessment of goodness of fit (see **Lack of Fit**) of a regression model should consist of a balanced judgment taking into account model assumptions, the explained variation by the model as well as the number of parameters in the model. This may be accomplished by analysis of the **residuals**, checking the number of significant effects, and performing a model significance test like the standard F -test in *analysis of variance* (which rigorously accounts for the *degrees of freedom* and hence, also incorporates the number of parameters).

Reference

- [1] Kvålseth, T.O. (1985). Cautionary note about R^2 , *The American Statistician* **39**(4), 279–285.

Related Articles

Analysis of Variance; Correlation; Dependence; Degrees of Freedom; Lack of Fit.

ALESSANDRO DI BUCCHIANICO

Correlation

Correlation is a measure of the strength of the linear relationship between two random variables, which is independent of the measurement scale. Theoretically, the correlation between X and Y is defined as

$$\text{Cor}(X, Y) = \frac{\text{Cov}(X, Y)}{\sqrt{\text{Var}(X)\text{Var}(Y)}} \quad (1)$$

where $\text{Cov}(X, Y)$ denotes the *covariance* (see **Covariance**) of X and Y , and $\text{Var}(X)$ and $\text{Var}(Y)$ are the *variances* (see **Variance**) of X and Y . If X and Y are independent, $\text{Cor}(X, Y) = 0$ (although the converse is not necessarily true). If X and Y are linearly

related (positive or negative), the correlation equals 1 or -1 .

The most often used sample correlation coefficient is the Pearson product-moment correlation. For a set of N pairs of observations $(X_1, Y_1), \dots, (X_N, Y_N)$ it is defined as

$$r = \frac{1}{N-1} \sum_{i=1}^N \left(\frac{X_i - \bar{X}}{S_X} \right) \left(\frac{Y_i - \bar{Y}}{S_Y} \right) \quad (2)$$

with $\bar{X} = 1/N \sum_{i=1}^N X_i$; $\bar{Y} = 1/N \sum_{i=1}^N Y_i$; $S_X^2 = 1/N - 1 \sum_{i=1}^N (X_i - \bar{X})^2$; and $S_Y^2 = 1/N - 1 \sum_{i=1}^N (Y_i - \bar{Y})^2$. We have $-1 \leq r \leq 1$, with equality only if the data are distributed along a perfect line in a scatter plot. See also **Measures of Association**.

Covariance

$$\text{Cov}(X, Y) = E(X - EX)(Y - EY) \quad (1)$$

See also **Correlation**.

The covariance is a measure of the departure from independence of two random variables X and Y :

Covariate

In regression models (and more complex linear models, such as ANCOVA models) one models the relationship between the response (or dependent) variable and one or more explanatory (or

independent) variables. The term *covariate* is used in particular for explanatory variables whose effect is not the main interest in the study (but which are included in the analysis to prevent their influence from affecting the estimation of the parameters of interest). Covariates are also called *concomitant variables*.

Dependence

Dependence and Independence – Definition

Let A and B be two random events. The events A and B are said to be statistically dependent if, and only if (the two definitions below are equivalent):

$$P(A|B) \neq P(A); P(B|A) \neq P(B) \quad (1)$$

This definition implies that two random events are statistically independent if, and only if:

$$P(A \wedge B) = P(A)P(B) \quad (2)$$

The basic definition may be extended to a set of k random events. Analogously, k random variables $\{X_1, \dots, X_k\}$, are mutually statistically independent if, and only if, their joint density function is equal to the product of their individual marginal density functions:

$$\begin{aligned} f_{x_1, x_2, \dots, x_k}(x_1, x_2, \dots, x_k) \\ = f_{x_1}(x_1)f_{x_2}(x_2) \dots f_{x_k}(x_k) \end{aligned} \quad (3)$$

Relationship between Dependence and Correlation

If two random variables are independent they are necessarily uncorrelated (*see Correlation*). But the reverse is not true: two dependent variables can be uncorrelated. The only case when two random variables that are dependent are also necessarily correlated and *vice versa* (namely, two correlated random variables are also necessarily dependent) is when they have a joint **normal distribution**.

Expected Value of a Product of Functions of Independent Random Variables

From equation (3), a very important and useful result is the following. Let $g_i(X_i)$ be a function of a single random variable, X_i . If $\{X_1, \dots, X_k\}$ maintain equation (3), then:

$$E[\prod_{i=1}^k g_i(X_i)] = \prod_{i=1}^k E[g_i(X_i)] \quad (4)$$

For example ($k = 2$, X_1 and X_2 are independent): $E(X_1^m X_2^n) = E(X_1^m)E(X_2^n)$.

HAIM SHORE

Descriptive Statistics

Introduction

Having a collection of observations, one may wish to describe these data so that important information, implied from the sample, becomes clearly visible. One way to do this is to make a plot of the data, such as a histogram, a stem-and-leaf plot, or a box and whisker plot (refer to **Graphical Representation of Data**). When the data convey a relationship between two variables, a scatter plot is a common method. Finally, for a **time series**, the observations are often plotted along the time axis so that patterns may be easily discerned. Often, such plots are not enough, and one wishes to summarize the data by a set of descriptive numbers that convey the most important properties of the data. Descriptive statistics provide this summary information. The sample properties that the descriptive statistics reflect often correspond to properties of the statistical distribution, like central tendency, dispersion, and the asymmetry and “peakedness” of the distribution. Some of these statistics may be used to estimate the distribution’s parameters or moments (see **Moments**).

Descriptive Statistics

Measures of Central Tendency

Mean.

$$\bar{X} = w_1 X_1 + w_2 X_2 + \cdots + w_n X_n \quad (1)$$

where X_i is the i th sample observation, n is the **sample size**, and $\{w_i\}$ are weights so that $\sum_{i=1}^n w_i = 1$.

Median and Sample Quantiles. The median is defined as the 50th percentile (see **Quantiles**) of the sample. For symmetrically distributed data, the sample mean and the sample median will be close (the mean and the median are equal for symmetric distributions). The median is considered a robust indicator for central tendency.

In a more general context, assume that one has n observations, and these are arranged in ascending order: $\{X_{(1)}, X_{(2)}, \dots, X_{(n)}\}$ ($X_{(i)}$ is referred to as the *sample i th order statistic*). The sample p th quantile

(the quantile such that a proportion p of the sample observations are smaller or equal to it) is calculated by the following procedure (refer to [1]):

- Express $(n+1)p$ by $r+g$, where r is an integer and $0 \leq g \leq 1$. For example, for $n=20$, $p=0.75$: $(n+1)p = 21 \times 0.75 = 15.75$ ($r=15$, $g=0.75$).
- The sample p th quantile is:

$$Q_p = (1-g)X_{(r)} + gX_{(r+1)}, r < n \quad (2)$$

For $r=n$, $X_{(n+1)}$ is replaced by $X_{(n)}$. Equation (2) may be used to calculate the median and the interquartile range (“Interquartile Range”). The sample 100 p th *percentile* is equal to the sample p th *quantile*.

Mode. The mode is the value that appears with the highest frequency in the sample.

Measures of Dispersion

Standard Deviation (S) and Coefficient of Variation (CV).

$$S = \left[\left(\frac{1}{n-1} \right) \sum_{i=1}^n (X_i - \bar{X})^2 \right]^{1/2} \quad (3)$$

Note that while S^2 is an unbiased estimator of the population variance, S is not an unbiased estimator of the population standard deviation. An often used measure of *scaled* dispersion, based on the sample standard deviation, is the coefficient of variation:

$$CV = \frac{S}{|\bar{X}|} \quad (4)$$

Being scale less, CV allows comparison of dispersions for samples of observations measured on different scales.

Range.

$$R = X_{\max} - X_{\min} = X_{(n)} - X_{(1)} \quad (5)$$

where $X_{\max}$ and $X_{\min}$ are the largest and the smallest observations in the sample.

Interquartile Range.

$$IQR = Q_{0.75} - Q_{0.25} \quad (6)$$

2 Descriptive Statistics

where $Q_{0.25}$ and $Q_{0.75}$ are the first and the third *quartiles* (see **Quartiles**) (the 0.25 and 0.75 **quantiles**, respectively), calculated by equation (2). The interquartile range contains 50% of the sample data and it is considered a robust descriptor of the sample dispersion (like the median is for central tendency).

Skewness and Kurtosis

Other common descriptive statistics are the sample *skewness* (see **Skewness**) and *kurtosis* (see **Kurtosis**). These represent the shape parameters of a distribution, and they reflect the degree of asymmetry and flatness, respectively, of the distribution.

Reference

- [1] Shore, H. (2005). *Response Modeling Methodology: Empirical Modeling for Engineering and Science*, World Scientific Publishing, Singapore.

Related Articles

Exploratory Data Analysis.

HAIM SHORE

Exponentially Weighted Moving Average (EWMA)

For a time series (*see Time Series Analysis*) $X_1, \dots, X_n$, the **exponentially weighted moving average** at time t is defined recursively for $t = 1, \dots, n$ as

$$Y_t = (1 - \lambda)X_t + \lambda Y_{t-1}, \quad \text{for } 0 < \lambda < 1 \quad (1)$$

(where Y_0 is an initial value) and thus

$$Y_t = (1 - \lambda)X_t + (1 - \lambda)\lambda X_{t-1} + (1 - \lambda)\lambda^2 X_{t-2} + \dots \quad (2)$$

The weights decrease geometrically. The exponentially weighted moving average (EWMA) is a smoother often used in forecasting and process control. *See also Moving Averages.*

Homogeneity of Variances

Introduction

Many important statistical techniques view each observation as a deterministic expression plus random error (noise), and assume homogeneity of variance: the measurement errors have common variance σ^2 . Homogeneity of variance is also known as *homoscedasticity*. Lack of homoscedasticity is known as *heteroscedasticity* (see also **Heteroscedasticity**).

Effects of Heteroscedasticity

For some techniques, heteroscedasticity is fatal. As an example, consider the one-way analysis of variance (see **Analysis of Variance**). Here, observations are drawn from several groups, and the null hypothesis of equality of group means is tested via an F -test (see **Parametric Tests**). The one-way analysis variance model views an observation as the sum of the corresponding group mean and a random error. The F -test compares two estimators of the common variance σ^2 : the factor mean square and the residual mean square (see **Mean Square Error**). The former is always an unbiased estimator of σ^2 , whereas the latter is an unbiased estimator of σ^2 only if the null hypothesis holds. Homogeneity of variance is crucial for the F -test: if we allow the variance of the error to vary from group to group, rejection of the null hypothesis may be due to differences in group variances rather than differences in group means.

Heteroscedasticity may not only affect the validity of statistical tests, it may also have a negative effect on the coverage of **confidence intervals** and the accuracy of estimators.

Detecting Heteroscedasticity

Formal diagnostic tests and diagnostic plots for detecting heteroscedasticity have been proposed for analysis of variance, the linear regression model as well as the nonparametric regression model. The best known are Bartlett's test (see the section on Comparing Variances in **Parametric Tests**) and Levene's test. See also [1–4]. Box and Cox transformations [**Box and Cox Transformation**] may be used to reduce heteroscedasticity, see [5].

References

- [1] Brown, M.B. & Forsythe, A.B. (1974). Robust tests for the equality of variances, *Journal of the American Statistical Association* **69**, 364–367.
- [2] Cook, R.D. & Weisberg, S. (1983). Diagnostics for heteroscedasticity in regression, *Biometrika* **70**, 1–10.
- [3] Dette, H. & Munk, A. (1998). Testing heteroscedasticity in nonparametric regression, *Journal of the Royal Statistical Society. Series B (Statistical Methodology)* **60**, 693–708.
- [4] O'Neill, M.E. & Mathews, K.L. (2002). Levene tests of homogeneity of variance for general block and treatment designs, *Biometrics* **58**, 216–224.
- [5] Ruppert, D. & Aldershof, B. (1989). Transformations to symmetry and homoscedasticity, *Journal of the American Statistical Association* **84**, 437–446.

ALEX J. KONING

Hotelling's T^2

Hotelling's T^2 is the multivariate equivalent of Student's t -statistic, and measures the deviation of a sample average from a hypothesized mean or target vector τ . Working with p -dimensional data (typically assumed to have a multivariate **normal distribution**), one has n p -tuples $\mathbf{X}_1, \dots, \mathbf{X}_n$. With the

p -dimensional vector $\bar{\mathbf{X}}$ denoting the sample average, and $\mathbf{S}^{-1}$ denoting the inverse of the sample **covariance** matrix, Hotelling's T^2 statistic is

$$T^2 = n(\bar{\mathbf{X}} - \tau)' \mathbf{S}^{-1} (\bar{\mathbf{X}} - \tau) \quad (1)$$

In industrial statistics the statistic is used in multivariate control charts (*see* **Multivariate Control Charts Overview**). See also **Multivariate Analysis**.

Kurtosis

The kurtosis of a random variable X is related to the fourth *moment* (see **Moments**) of its distribution:

$$\gamma_2 = \frac{\mu_4}{\sigma^4} - 3 \quad (1)$$

where $\mu_4 = E(X - EX)^4$ and $\sigma^2 = E(X - EX)^2$. Distributions with high kurtosis are typically heavy tailed whereas distributions with low kurtosis are light tailed.

Lack of Fit

In parametric modeling one speaks of lack of fit if there is a systematic discrepancy between the predicted values and the observed values due to the chosen parametric model. Substantial lack of fit occurs, for example, if one approximates a curved relationship by a linear function, or if significant interaction terms are dropped from a model. In the design of experiments it is good practice to

include a few design points that enable the assessment of lack of fit, such as runs in the *center point* (see **Center Points**) of a two-level **factorial** design. If the design and the chosen model allow, the error **sum of squares** in the analysis is typically broken down into a pure error term (due to noise) and a lack-of-fit term (see **Analysis of Variance**).

Leverage

Leverage in Linear Regression Analysis

In regression analysis, leverage is a measure of the influence of a given observation on the regression line due to the location of the values of the independent variables. This may happen when an observation is taken at values of the independent variables far away from other observations. The common *least-squares estimation* method in linear regression analysis forces regression models to be close to observed values taken at remote values of the independent variables (see Figure 1). Hence, small changes in the measured response at remote observations may have a large effect on the coefficients of the regression model. It is important to detect such observations. For linear regression models of the form $Y = X\beta + \varepsilon$, the *least-squares* estimator can be written as $\hat{y} = Hy$, where the hat matrix H is given by $X(X^T X)^{-1} X^T$. The matrix elements h_{ij} of the hat matrix can be interpreted as the amount of influence of observation y_j on fitted value $\hat{y}_i$. As a rule of thumb, observations for which the corresponding diagonal element of the hat matrix exceeds $2p/n$ (where p is the number of terms in the regression model and n is the number of observations) are considered to be influential. This is a practical test, but it only works for observations taken at remote values of the independent variables. There exist other influence measures that quantify the effect of deleting individual observations that apply to all observations. Well-known examples include Cook's D , $DFFITS$, and $DFBETAS$. Let p be the number of terms in the linear regression model and denote the standardized *residuals* by r_i . Then one of the forms for Cook's D is given by $\frac{r_i^2}{p} \left(\frac{h_{ii}}{1-h_{ii}} \right)$. When Cook's D exceeds $4/n$, then the i th observation is considered to have caused a significant shift of the entire parameter vector β due to the i th observation. If one is interested in the shift of a single parameter β_j , then one may use $DFBETA_{j(i)}$ which is defined by $\frac{\hat{\beta}_j - \hat{\beta}_{j(i)}}{s_i \sqrt{c_{ii}}}$, where c_{ii} denotes the i th diagonal element of the matrix $(X^T X)^{-1}$ and $\hat{\beta}_{j(i)}$ is the estimate for β_j based on the data set without the

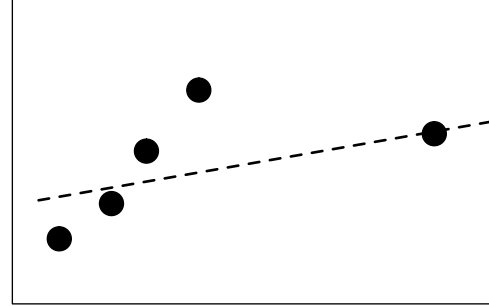


Figure 1 Leverage: the regression line is forced to be close to the remote observation

i th observation. $DFBETA_{j(i)}$ has a cutoff value of $2/\sqrt{n}$. The measure $DFFITS$ measures the difference between the model predictions for the i th observation based on the complete data set and those based on the data set without the i th observation. $DFFITS$ is defined as $\frac{h_{ii}}{1-h_{ii}} \frac{e_i}{s_{(i)} \sqrt{1-h_{ii}}}$, where e_i is the i th *residual* and $s_{(i)}$ denotes the *standard deviation* based on the data set without the i th observation. It has a cutoff value of $2\sqrt{p/n}$. Extensive information on these and other influence measures can be found in [1, 2]. The distance on which these measures are based are similar to those employed in **multivariate statistics**, in particular, the Mahalanobis distance.

References

- [1] Belsley, D.A., Kuh, E. & Welsch, R.E. (1980). *Regression Diagnostics: Identifying Influential Data and Sources of Collinearity*, John Wiley & Sons, New York.
- [2] Cook, R.D. & Weisberg, S. (1982). *Residuals and Influence in Regression*, Chapman & Hall, London.

Related Articles

Least-Squares Estimation; Residuals.

ALESSANDRO DI BUCCHIANICO

Mean Square Error

In *analysis of variance* (see **Analysis of Variance**) the total **sum of squares** is divided into a sum of squares attributed to the factors, and a residual sum of squares (often called the *error sum of squares*). The residual sum of squares divided by the residual **degrees of freedom** is the mean square error (MSE). The square

root of the MSE is often used as an estimator of noise in the measurements. For example, in the simple linear regression analysis $y_i = a + bx_i + \epsilon_i$, $i = 1, \dots, n$, the MSE is given by

$$MSE = \frac{1}{n-2} \sum_{i=1}^n (y_i - \hat{y}_i)^2 \quad (1)$$

with $\hat{y}_i$ the least-squares fit for y_i .

Measures of Location

Measures of location are statistics that yield information on the point of gravity or the center of the data. A more formal description is a function of data that is equivariant under affine transformations, i.e., commutes with the group of transformations $x \rightarrow ax + b$. Common examples include the sample mean $\sum_{i=1}^n x_i/n$, the sample median (the middle value of the data after sorting from small to large in the case of an odd number of observations or the average of the two middle values in the case of an even number of observations), and the sample mode (the value that appears most frequently; this need not be a unique value). When data are normally distributed (*see* **Normal Distribution**), then the sample mean is the optimal estimator for the population mean in the sense that it is the unique unbiased estimator with minimal variance. A disadvantage of the mean is that it is not robust against *outliers*. Well-known alternatives that have a certain resistance against outliers include the median, the shorth (the sample mean of the shortest intervals containing at least half of the observations), *trimmed*

means (sample means when extreme order statistics have been removed), winsorized means (sample means when extreme order statistics have been set to an extreme order statistic), and the Hodges–Lehmann estimator (median of pairwise averages). Many of these sample statistics can be derived from general construction principles or techniques for estimators, like **Maximum Likelihood**, method of **moments**, M-estimators (extension of Maximum Likelihood), L-estimators (linear combinations of order statistics), or R-estimators (rank-based estimators). We refer to [1] for extensive details.

Reference

- [1] Barnett, V. & Lewis, T. (1994). *Outliers in Statistical Data*, John Wiley & Sons, New York.

Related Articles

Maximum Likelihood; Normal Distribution; Outliers; Trimmed Mean.

ALESSANDRO DI BUCCHIANICO

Measures of Scale

Measures of scale are statistics that yield information on the spread of data. A more formal definition is that it is a function of data that changes with a multiplicative factor $|a|$ after applying an affine transformation $x \rightarrow ax + b$. The most well-known example is the standard deviation, which is the square root of the sample variance $\sum_{i=1}^n (x_i - \bar{x})^2 / (n - 1)$ (the sample *variance* is more tractable from a computational point of view, but has the drawback that it does not have the right unit when applied to physical data). A disadvantage of the standard deviation is that it is not robust against *outliers*. This is even more the case for the range (the difference of the largest and the smallest observation), which was a popular alternative to the standard deviation in the precomputer era. The range and the related average moving range (AMR) $\sum_{i=1}^{n-1} |x_{i+1} - x_i| / (n - 1)$ are often used in **quality control**, especially for **control charts** (see **Moving Range and R Charts**). Well-known alternatives that have a certain resistance against *outliers* include the interquartile range (distance between third and first *quartile*), the median

absolute deviation (MAD), which is the median of the absolute deviations from the median), and the length of the shorth (the shortest interval containing at least half of the data, see **Measures of Location**). Many of these sample statistics can be derived from general construction principles or techniques for estimators, like *Maximum Likelihood*, method of **moments**, M-estimators (extension of maximum likelihood), L-estimators (linear combinations of order statistics), or R-estimators (rank-based estimators). We refer to [1] for extensive details.

Reference

- [1] Barnett, V. & Lewis, T. (1994). *Outliers in Statistical Data*, John Wiley & Sons, New York.

Related Articles

Maximum Likelihood; Measures of Location; Outliers; Variance.

ALESSANDRO DI BUCCHIANICO

Monte Carlo Methods

Monte Carlo methods – definition and common uses

Monte Carlo methods refer to computer-intensive methods which are used to mimic *unknown* statistical distributions of *output* random variables, *via* generation of random numbers from *input* random variables with *known* distributions. Both input and output variables (random or deterministic) are tied together by a stochastic model. Application of **Monte Carlo** methods is conducted in the framework of **Monte Carlo simulation** (namely, simulation that is stochastic in some manner). Major areas of application of Monte Carlo methods include:

- Obtaining numerical solutions to certain mathematical problems which are too complicated to solve analytically (such as the calculation of algebraically intractable integrals *via* Monte Carlo integration).
- Finding optimal solutions for problems too complex to solve mathematically (like NP-complete problems) *via* optimum-search techniques (like simulated annealing and **genetic algorithms**).
- System simulation (like physical, biological, and financial systems) in order to study and predict their behavior under various stochastic scenarios.
- Simulating distributions of statistics with unknown distributions for statistical inference (**bootstrap** and jackknife methods).
- Inverse problems, where a large collection of possible models are pseudorandomly generated *via* Monte Carlo methods and the associated likelihood of model's properties are evaluated.

Some techniques employed in Monte Carlo simulation

A major effort in Monte Carlo simulation relates to generating random numbers from the specified distributions. There are several general methods for random number generation, like rejection methods or the simpler inversion method. The latter refers to the case where random numbers need to be generated from a distribution with a known inverse cumulative distribution function (cdf) (*see* **Cumulative Distribution Function (CDF)**). Since any cdf, expressed

as a function of its random variable, is uniformly distributed, introducing a set of numbers, sampled from the uniform distribution on the (0, 1) interval, into the *inverse* cdf would produce a set of random numbers that has the desired distribution. For example, for the inverse cdf of the exponential distribution with parameter λ :

$$x = F^{-1}(P) = -(1/\lambda)\exp(1 - P) \quad (1)$$

where $P = F(X \leq x)$, X is exponentially distributed with parameter λ . Since P is uniformly distributed, uniformly distributed random numbers introduced into equation (1) would produce exponentially distributed numbers. When the distribution does not have an explicit inverse, like the **normal distribution**, special methods have been devised to generate the required distribution (including approximations for the algebraically intractable inverse cdf). Simulation software packages nowadays have random number generators for most commonly used distributions. These methods actually generate *pseudorandom* numbers since deterministic methods are used for the generation of the random numbers.

A second major effort in application of Monte Carlo methods is variance reduction. Variance reduction techniques aim to generate simulation data with minimal variance in order to increase the precision of estimates derived from the simulated data. There are several available variance reducing techniques, like antithetic **resampling**, the method of control variates, indirect **estimation**, and others.

Monte Carlo methods in quality and reliability engineering

Monte Carlo methods are extensively employed in reliability engineering to compute reliability of systems with complex configurations, or systems with components that are not necessarily statistically independent. Another common application is to calculate algebraically intractable integrals that appear in reliability-related calculations. The effects of various policies of preventive maintenance or various warranty policies are also often examined *via* Monte Carlo methods. In quality engineering these methods are used to assess the impact of a product design change, examine the effect of process improvement, and to optimize operating conditions of a running process.

HAIM SHORE

Moving Averages

Introduction

A moving average is a time-dependent weighted average, which is repeatedly recalculated as further observations in a **time series** become available. The name derives from the fact that the average “moves” through the time series. As more observations are collected, they are added to the calculation of the average, and at the same time observations from the most distant past are deleted (or given smaller weights). There are two basic reasons to calculate a **moving average**:

- **Prediction**

A moving average performs *smoothing* of the various components that comprise the time series (like trend, seasonality, or randomness), so that predicting immediate-future values of these components, based on the appropriate moving average, is less affected by fluctuations that had occurred in these components in the past.

- **Adaptive control**

To provide predicted values for a monitored process variable that will reflect adaptation, in a controlled manner, of the forecasting method to changes that have recently occurred in the process. The predicted values, based on the moving average, may then serve as baseline to detect out-of-control states for the process (*see Exponentially Weighted Moving Average (EWMA)*).

Unlike time-series modeling, a moving average does not attempt to model the data. Rather it attempts

to smooth the data so that any *long-term* signal embedded in the observations becomes more visible.

The Standard Moving Average

This average, calculated to predict the observation at time t , and based on the most recent n observations (collected from time $t - n$ until $t - 1$), is the weighted average

$$\bar{X}_t = w_1 X_{t-1} + w_2 X_{t-2} + \cdots + w_n X_{t-n} \quad (1)$$

where the weights sum to 1: $\sum_{i=1}^n w_i = 1$. It is customary to assign smaller weights to more ancient data.

The Exponentially Weighted Moving Average

Alternatively, one could use the **exponentially weighted moving average**:

$$\bar{X}_t = \lambda X_t + (1 - \lambda) \bar{X}_{t-1} = \bar{X}_{t-1} + \lambda(X_t - \bar{X}_{t-1}), \quad 0 \leq \lambda \leq 1 \quad (2)$$

where $\bar{X}_0 = X_0$ (*see for further details Exponentially Weighted Moving Average (EWMA)*). A larger λ would provide better (quicker) adaptation to most recent fluctuations, and thus less smoothing of these fluctuations. It is easily shown that the average in equation (2) is in fact an exponentially-weighted average of all previous observations, where weights become smaller the older an observation is.

HAIM SHORE

Multi-Vari Chart

A Graphical Tool for the Discovery of Sources of Variation

The multi-vari chart is a graphical tool, mainly used in **Exploratory Data Analysis**, which provides a display of the behavior of a quality characteristic in the running process. Variance components and patterns in the data are easily recognized.

Seder [1] introduced the multi-vari chart and compares it with the $\bar{X}$ -bar, R control chart (*see Control Charts, Overview*). It presents an analysis of the variation in a process, differentiating among three main sources: intrapiece (the variation within a piece, batch, lot); interpiece (the additional variation between pieces); and temporal variation. Duncan [2] presents the multi-vari chart as a graphical tool for testing joint homogeneity of means and variance in a nested structure. Bhote [3] discusses the multi-vari chart as a tool useful for generating clues about factors that cause process variation.

To illustrate the multi-vari chart, we consider data from a tea packing process, where the weights of boxes filled with 20 tea bags are determined. During each of 11 days, five consecutive boxes were sampled from each of four packing machines. By Y_{ijk} , $i =$

$1, \dots, 11$; $j = 1, \dots, 4$; $k = 1, \dots, 5$, we denote the weight of the k th box from the j th machine on the i th day. Figure 1 shows the multi-vari chart made from these data. On its y axis, it has the scale of the measurements, grams in this case. On the x axis are the indices from the samples. From each machine sample, the minimum and the maximum weights are indicated by a solid box. For each sample, the two boxes are connected by a vertical line, representing the range of the machine sample. Within each day, the sample means of the four machine samples are connected by a line. Finally, the crosses indicate the day means.

Studying the chart in Figure 1, we see that the sample taken at the seventh day at the second machine, and the eighth day's samples of the first and second machines require a closer examination. Inspection of the corresponding individual weights learned that there are four boxes that have a weight under 76 g or above 84 g. Because there should be 20 tea bags in a box, and the nominal weight of a tea bag is 4 g, it is likely that in these four boxes there is either one tea bag missing, or one tea bag too many.

Besides such errors, the multi-vari chart gives information about variance components. Substantial spread among the crosses (the day means in Figure 1) is indicative for a significant variance component associated with time. Steep rises and drops in the

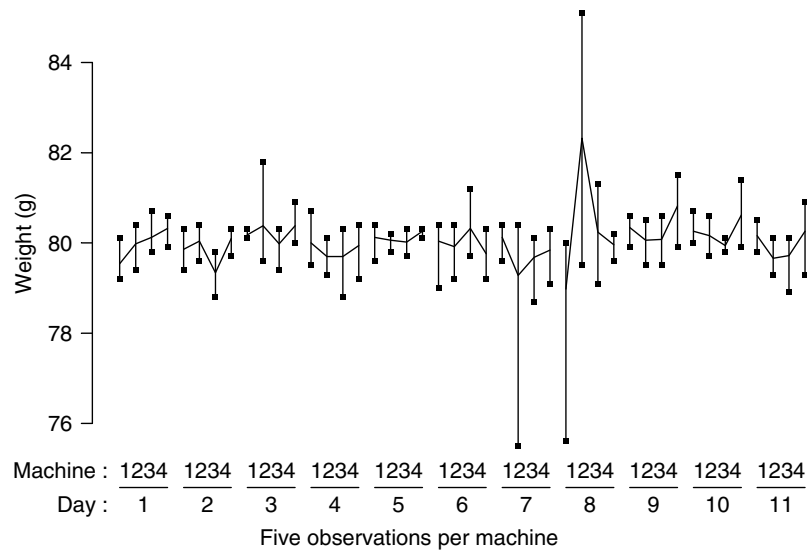


Figure 1 Multi-vari chart for weights of tea boxes

2 Multi-Vari Chart

lines connecting machine means within days shows substantial interpiece or stream-to-stream differences. Tall vertical bars within the machine samples indicate substantial intrapiece spread. Besides variance components, the multi-vari chart also shows trends and systematic patterns between machines and over time.

De Mast *et al.* [4] give a thorough discussion of the multi-vari chart and augment the tool with tests based on control limits. Equipped with these control limits, one could regard the multi-vari chart as an extension of the $\bar{X}$ -bar, R control chart, allowing for two factors instead of one. Another related tool is the analysis of means (ANOM) [5].

References

- [1] Seder, L.A. (1990). "Diagnosis with diagrams", part I and II, *Industrial Quality Control* **6**(1), 11–19 (part I) and 6(2) 7–19 (part II).
- [2] Duncan, A.J. (1986). *Quality Control and Industrial Statistics*, 5th Edition, Irwin Publishing, Homewood.
- [3] Bhote, K.R. (1991). *World Class Quality*, AMACOM, New York.
- [4] De Mast, J., Roes, C.B. & Does, R.J.M.M. (2001). The multi-vari chart: a systematic approach, *Quality Engineering* **13**(3), 437–447.
- [5] Schilling, E.G. (1973). A systematic approach to the analysis of means, *Journal of Quality Technology* **5**(3), 93–108.

JEROEN DE MAST

Nonparametric Tests

Introduction

Nonparametric tests are the counterparts to *parametric tests* (see **Parametric Tests**). Typically, parametric tests assume that the data are drawn from a distribution, which has a finite number of unknown parameters, often the **normal distribution** with unknown mean μ and σ^2 . In contrast, nonparametric tests do not rely on such rather restrictive distributional assumptions, and hence are also known as *distribution-free tests*.

The parameter space is the set of all possible values of the unknown parameters. In mathematical statistics, a statistical method is called *parametric* when the parameter space is finite dimensional, and *nonparametric* when the parameter space is infinite dimensional, such as the space of all possible distribution functions. To complicate matters, there are also semiparametric models in which the parameter vector consists of a finite-dimensional part and an infinite-dimensional part. For instance, if we assume that we are dealing with data from a distribution that is symmetric around some unknown median η , then we are in essence dealing with a semiparametric model, as the parameter vector consists of a finite-dimensional part η and an infinite-dimensional part, a distribution that is symmetric around zero. Below, we shall ignore the distinction between nonparametric and semiparametric methods.

Permutation Tests

Nonparametric tests have been around since the early days of statistical inference. Perhaps the oldest form of nonparametric test is the permutation test, see [1, 2]. If data have been obtained from experiments that involve **randomization**, then we may assess the statistical significance of a permutation test by recalculating the test statistic for every permutation of the data consistent with the randomization. Obviously, permutation tests are computer intensive, which has impeded practical use for a long time. Recently, there is renewed interest in permutation tests.

Rank Tests

The computational demands of a permutation test may be drastically reduced by replacing the observations with their (transformed) ranks before actually performing the test. The resulting test is called a *rank test*.

An illustrative example is the signed rank test, see [3]. This test assumes that n observations are drawn from a distribution that is symmetric around an unknown location θ . The signed rank test may also be applied to paired data.

The null hypothesis to be tested is $H_0 : \theta = \theta_0$. The test assigns ranks to the observations according to their distance to θ_0 . Therefore, the observation closest to θ_0 is assigned rank 1. The test statistic T is obtained by summing the ranks belonging to the observations smaller than θ_0 , and is asymptotically normal.

To assess statistical significance (see **Significance Level**), we recalculate this rank sum for every permutation of the data. For example, if $n = 4$, then under the null hypothesis there are 2^4 equally likely outcomes for the ranks assigned to the observations smaller than θ_0 see Table 1.

Thus, the test statistic T takes values 0, 1, 2, 8, 9, and 10 with probability $1/16$, and values 3, 4, 5, 6, and 7 with probability $1/8$.

It is known that the test statistic T has an approximately normal distribution for $n \geq 20$. Under the null hypothesis, the test statistic T has expected value $n(n+1)/4$ and variance $n(n+1)(2n+1)/24$, and hence approximate critical values of the standardized

Table 1 Derivation of the null distribution of the rank sum test statistic for $n = 4$. The rank sum test statistic is the sum of the ranks assigned to observations smaller than θ_0

Ranks assigned	Rank sum T	Ranks assigned	Rank sum T
–	0	2, 3	5
1	1	2, 4	6
2	2	1, 2, 3	6
3	3	3, 4	7
1, 2	3	1, 2, 4	7
4	4	1, 3, 4	8
1, 3	4	2, 3, 4	9
1, 4	5	1, 2, 3, 4	10

2 Nonparametric Tests

test statistic

$$\frac{T - \frac{n(n+1)}{4}}{\sqrt{\frac{n(n+1)(2n+1)}{24}}} \quad (1)$$

may be determined by means of the standard normal distribution for $n \geq 20$. For $n < 20$, tabulated critical values have to be used.

Another example is the two-sample rank sum test, see [3, 4]. It involves two independent samples. Suppose that the first sample has size n_1 and is drawn from a distribution with population median η_1 and that the second sample has size n_2 and is drawn from a distribution with population median η_2 , the null hypothesis to be tested is $H_0 : \eta_1 = \eta_2$.

The combined sample is ranked from the least to the greatest. The Wilcoxon rank sum test statistic W is the sum of the ranks assigned to the first sample. It is known that the test statistic W has an approximately normal distribution when both sample sizes n_1 and n_2 are larger than 10. Under the null hypothesis, the test statistic W has expected value $n_1(n_1 + n_2 + 1)/2$ and variance $n_1 n_2 (n_1 + n_2 + 1)/12$. Therefore, approximate critical values of the standardized test statistic

$$W - \frac{n_1(n_1 + n_2 + 1)}{2} \quad (2)$$

$$\sqrt{\frac{n_1 n_2 (n_1 + n_2 + 1)}{12}}$$

may be determined by means of the standard normal distribution for $n_1, n_2 \geq 10$. The Mann–Whitney test statistic U is related to W , and is, in fact, given by

$$U = W - \frac{n_1(n_1 + 1)}{2} \quad (3)$$

It follows that the Wilcoxon rank sum test and the Mann–Whitney test are equivalent.

A further example is the Kruskal–Wallis test, which compares the effects of $k > 2$ treatments. To each treatment, there belongs a sample, and these samples are independent. The Kruskal–Wallis test assumes that

$$X_{ij} = \eta + \tau_j + \epsilon_{ij} \quad (4)$$

where X_{ij} is the j th observation in the i th sample, η is an unknown constant, τ_j is the unknown effect

of treatment j , and the error variables, ϵ_{ij} 's, are a random sample from a distribution with median 0. The null hypothesis to be tested is

$$H_0 : \tau_1 = \tau_2 = \dots = \tau_k \quad (5)$$

Rank the combined sample from the least to the greatest, and let $\bar{R}_j$ denote the mean of the ranks assigned to sample j . The Kruskal–Wallis test statistic is defined as

$$H = \frac{12}{N(N+1)} \sum_{j=1}^k n_j (\bar{R}_j - \bar{\bar{R}})^2 \quad (6)$$

where n_j is the size of sample j , $N = n_1 + \dots + n_k$ is the size of the combined sample, and $\bar{\bar{R}} = (N+1)/2$ is the average rank in the combined sample. The null hypothesis is rejected for large values of H . Under the null hypothesis, H approximately follows a χ^2 distribution with $k-1$ **degrees of freedom** for $\min(n_1, \dots, n_k) \geq 5$.

Instead of rank sums $\sum R_i$, one may use sums of transformed ranks $\sum a(R_i)$. For instance, the two-sample Savage test employs the transformed ranks

$$a(R_i) = \sum_{j=N-R_i+1}^N 1/j \quad (7)$$

see [5].

Rank tests may be generalized so as to accommodate censored observations. Generalizations of the two-sample rank sum test and the Savage test are found in [6] and [7], respectively. The latter test is commonly known as the *logrank test* and is related to the Cox regression model.

For an extensive overview of rank tests, refer to [8, 9].

Resampling Methods

The advent of high-speed computing has had a beneficial effect on nonparametric statistics. As computation became fast and cheap, ideas that had been around for some time suddenly became applicable. In particular, this holds true for permutation tests and resampling methods, see [10].

Resampling methods are related to permutation tests. However, instead of permuting the data, subsamples are drawn. The most prominent resampling

method is the **bootstrap**, see [11]. In its simplest form, the bootstrap constructs bootstrap samples by sampling with replacement from the initial sample. Often, the bootstrap is used to construct confidence intervals. However, the bootstrap may also be used to obtain critical values for some test statistic.

Let $X_1, \dots, X_n$ be a sample drawn from P_θ , where θ is a (possibly infinite-dimensional) parameter that belongs to some parameter space. We intend using the test statistic $T_n = T_n(X_1, \dots, X_n)$ to test the null hypothesis $H_0 : \theta \in \Theta_0$ versus $H_1 : \theta \in \Theta \setminus \Theta_0$, where $\Theta_0 \subset \Theta$. Let $\hat{\theta} = \hat{\theta}(X_1, \dots, X_n)$ be a consistent estimator of θ under the null hypothesis, and let $T_n^* = T_n(X_1^*, \dots, X_n^*)$, where $X_1^*, \dots, X_n^*$ is a bootstrap sample drawn from $P_{\hat{\theta}}$. The conditional distribution of T_n^* given $X_1, \dots, X_n$ is the bootstrap estimate for the null distribution of T_n . **Quantiles** of the bootstrap null distribution may act as bootstrap critical values for T_n .

The error in the probability of a type I error is of smaller order for a bootstrap test than for the corresponding asymptotical theory test if and only if T_n is an asymptotic pivotal quantity, that is, the asymptotic null distribution of the test statistic T_n does not depend on θ , see [12]. Therefore, a bootstrap test should preferably be based on an asymptotic pivotal quantity.

For instance, let Θ be the collection of distributions on the real line whose first 10 **moments** are finite, and let $X_1, \dots, X_n$ be a random sample drawn from a distribution belonging to Θ . Let Θ_0 be the set of elements of Θ that have mean μ_0 . Then, the bootstrap samples should be resampled from $X_1^0, \dots, X_n^0$, where $X_i^0 = X_i + \mu_0 - \bar{X}_n$. That is, X_i^0 is obtained by shifting X_i over a distance $\mu_0 - \bar{X}_n$. Note that the mean of the shifted sample $X_1^0, \dots, X_n^0$ is equal to μ_0 , and hence the empirical distribution of $X_1^0, \dots, X_n^0$ belongs to Θ_0 . Moreover,

the bootstrap test may be based on the studentized statistic $T_n = \sqrt{n}(\bar{X}_n - \mu_0)/S_n$, which is an asymptotic pivotal quantity.

References

- [1] Fisher, R.A. (1935). *The Design of Experiments*, Oliver & Boyd, London.
- [2] Pitman, E.J.G. (1937). Significance tests which may be applied to samples from any populations, *Supplement to the Journal of the Royal Statistical Society* **4**, 119–130.
- [3] Wilcoxon, F. (1945). Individual comparisons by ranking methods, *Biometrics Bulletin* **1**, 80–83.
- [4] Mann, H.B. & Whitney, D.R. (1947). On a test of whether one of two random variables is stochastically larger than the other, *Annals of Mathematical Statistics* **18**, 50–60.
- [5] Savage, R.I. (1956). Contributions to the theory of rank order statistics – the two-sample case, *Annals of Mathematical Statistics* **27**, 590–615.
- [6] Gehan, E.A. (1965). A generalized Wilcoxon test for comparing arbitrarily singly-censored samples, *Biometrika* **52**, 203–223.
- [7] Mantel, N. (1966). Evaluation of survival data and two new rank order statistics arising in its consideration, *Cancer Chemotherapy Reports* **50**, 163–170.
- [8] Hollander, M. & Wolfe, D.A. (1999). *Nonparametric Statistical Methods*, John Wiley & Sons, New York.
- [9] Hettmansperger, T.P. & McKean, J.W. (1998). *Robust Nonparametric Statistical Methods*, Edward Arnold, London.
- [10] Efron, B. (1979). Computers and the theory of statistics: thinking the unthinkable, *SIAM Review* **21**(4), 460–480.
- [11] Efron, B. (1979). Bootstrap methods: another look at the jackknife, *Annals of Statistics* **7**, 1–26.
- [12] Beran, R. (1988). Prepivoting test statistics: a bootstrap view of asymptotic refinements, *Journal of the American Statistical Association* **83**, 687–697.

ALEX J. KONING

Normal Distribution

Introduction

This article describes the univariate normal distribution, with brief references also to the bivariate normal distribution and the multivariate normal. For each of these, there are brief historical remarks, and discussions of distributional properties, sampling distributions, and related applications.

The Univariate Normal Distribution

Historical Remarks

The univariate normal distribution is probably the most important distribution in classical statistical theory and methods. This distribution has a “bell-shaped” continuous *probability density function* (see **Probability Density Function (PDF)**). In the history of statistics, it was first discovered by the German mathematician Carl F. Gauss in the early nineteenth century while he was studying certain problems in physics and astronomy. As a result, this distribution is also known as the *Gaussian distribution*.

Distributional Properties

A continuous univariate random variable X is said to follow a normal distribution with parameters μ and σ^2 , if its probability density function $f(x)$ is of the form

$$f(x) = \left[\frac{1}{(2\pi)^{1/2}\sigma} \right] \exp \left[\frac{-(x - \mu)^2}{2\sigma^2} \right] \\ -\infty < x, \mu < \infty, \text{ and } \sigma > 0 \quad (1)$$

in symbols, $X \sim N(\mu, \sigma^2)$. It can be shown by elementary calculus that the mean and the variance of this distribution are μ and σ^2 , respectively.

A random variable Z is said to follow a *standard normal distribution* if $Z \sim N(0, 1)$. The distribution function of Z , $\Phi(z)$, is then given by

$$\Phi(z) = \int_{-\infty}^z \left[\frac{1}{(2\pi)^{1/2}} \right] \exp \left(\frac{-u^2}{2} \right) du \quad (2)$$

Since the function $\exp(-u^2/2)$ is symmetric about 0, $\Phi(z)$ satisfies $\Phi(z) = 1 - \Phi(-z)$ for all z . The $(1 - \alpha)$ th *quantile* (see **Quantiles**) (or $100(1 - \alpha)$ th percentile) of a standard normal distribution, often denoted z_α in most textbooks, is the quantity that satisfies $\Phi(z_\alpha) = 1 - \alpha$, $\alpha \in (0, 1)$. Since the function $\Phi(z)$ cannot be expressed in a closed form, tables for the numerical values of $\Phi(z)$ and z_α are needed. Such tables can be found in most statistics books.

The following theorem, which can be proved by calculus, shows how a normal distribution with mean μ and variance σ^2 is related to the standard normal distribution.

Theorem 1 If $X \sim N(\mu, \sigma^2)$, then the random variable $Z = (X - \mu)/\sigma$ is distributed according to the standard normal distribution.

As a simple application of Theorem 1, it follows that

Fact 1. If $X \sim N(\mu, \sigma^2)$, then,

1. $\Pr[a < X \leq b] = \Phi((b - \mu)/\sigma) - \Phi((a - \mu)/\sigma)$ for all $a < b$;
2. the $(1 - \alpha)$ th quantile of the distribution of X is $\mu + z_\alpha\sigma$.

Thus, tables for the standard normal distribution can be used for any univariate normal distribution.

More detailed distribution properties of the univariate normal distribution can be found in Patel and Read [1] and other related sources.

Sampling Distributions

The χ^2 distribution, Student's t distribution, and the F distribution (see **Probability Density Functions**), generally considered to be the cornerstones of classical statistical analysis, are closely related to the normal distribution. Specifically, let $X_1, X_2, \dots, X_N$ be a random sample of size N from an $N(\mu, \sigma^2)$ distribution; let

$$\bar{X} = \frac{1}{N} \sum_{i=1}^N X_i \quad \text{and} \\ S^2 = \frac{1}{N-1} \sum_i^N (X_i - \bar{X})^2 \quad (3)$$

2 Normal Distribution

denote the sample mean and the sample variance, respectively. Then,

1. $\bar{X}$ has an $N(\mu, \sigma^2/N)$ distribution and $(N-1)S^2/\sigma^2$ has a χ^2 distribution with $N-1$ **degrees of freedom**. Furthermore, $\bar{X}$ and S^2 are independent.
2. $t = N^{1/2}(\bar{X} - \mu)/S$ has a Student's t distribution with $N-1$ degrees of freedom, where $S = \sqrt{S^2}$ is the sample standard deviation.
3. t^2 has an F distribution with degrees of freedom $(1, N-1)$.

Another important sampling distribution result related to the normal distribution is a fundamental theorem in **probability theory**, called the *central limit theorem*.

Theorem 2 Let $X_1, X_2, \dots, X_N$ be a random sample of size N from any population with mean μ and finite variance σ^2 . Then, for every fixed z ,

$$\lim_{N \rightarrow \infty} \Pr\{N^{1/2}(\bar{X} - \mu)/\sigma \leq z\} = \Phi(z) \quad (4)$$

This theorem provides an approximation for the distribution of $\bar{X}$ when N is large. In most applications,

$$\Pr[a < \bar{X} \leq b] \doteq \Phi(N^{1/2}(b - \mu)/\sigma) - \Phi(N^{1/2}(a - \mu)/\sigma) \quad (5)$$

when $N \geq 30$.

Related Applications

In various applications of statistical inference problems – including *estimation* (see **Estimation**) of the parameters and *hypothesis testing* (see **Hypothesis Testing**) – the above sampling distribution results are applied. These include the following:

1. For inference on μ when σ^2 is known, the results for the distribution of $\bar{X}$ and Theorem 2 may be applied.
2. For inference on μ when σ^2 is unknown under the assumption of normality, Student's t distribution may be used.
3. When making statistical inference on σ^2 under the assumption of normality, the χ^2 distribution may be used.
4. The normal distribution may be applied for inference on the difference of two normal means

when their variances are known. Similarly, the Student's t distribution may be applied for the same purpose when the variances are unknown but equal; in the hypotheses-testing problem, this is known as the *two-sample t test* (see **Parametric Tests**).

5. The F distribution may be applied for testing the equality of $k \geq 2$ normal means when the variances are assumed to be equal but unknown. This method is known as the *one-way analysis of variance (ANOVA) method* (see **Analysis of Variance**). When $k = 2$, it reduces to the two-sample t test as a special case.
6. Other results related to Theorem 2 are useful in large-sample inference problems. For example, it is known that under regularity conditions the *maximum-likelihood* (see **Maximum Likelihood**) estimator has an asymptotically normal distribution, and that the asymptotic null distribution of $-2 \ln \lambda$ is χ^2 , where λ is the likelihood-ratio function in a likelihood-ratio test.

The Bivariate Normal Distribution

Historical Remarks

Studies of the bivariate normal distribution seem to begin in the middle of the nineteenth century, and moved forward dramatically when Galton published his work [2] on the applications of **correlation** analysis in genetics. As Karl Pearson noted in his 1920 *Biometrika* paper [3], “In 1885 Galton had completed the theory of bivariate normal correlation” but, because he “was very modest and throughout his life underrated his own mathematical powers, he did not at once write down the equation” of the bivariate normal density function. Consequently, it was Pearson himself who gave a definitive mathematical formulation of the bivariate normal distribution in his 1896 paper [4] on regression and heredity.

Distributional Properties

A two-dimensional random vector (X_1, X_2) is said to have a bivariate normal distribution if their joint density function $f(x_1, x_2)$ is of the form

$$f(x_1, x_2) = [2\pi\sigma_1\sigma_2(1 - \rho^2)^{1/2}]^{-1} \exp\left[-\frac{1}{2}Q_2(x_1, x_2; \boldsymbol{\mu}, \boldsymbol{\Sigma})\right], \quad -\infty < x_1, x_2 < \infty \quad (6)$$

where

$$Q_2(x_1, x_2; \boldsymbol{\mu}, \boldsymbol{\Sigma}) = (x_1 - \mu_1, x_2 - \mu_2) \boldsymbol{\Sigma}^{-1} \times \begin{pmatrix} x_1 - \mu_1 \\ x_2 - \mu_2 \end{pmatrix} \quad (7)$$

defines an ellipse centered at $(\mu_1, \mu_2) \equiv \boldsymbol{\mu}$ (which is the mean vector), the 2×2 matrix

$$\boldsymbol{\Sigma} = \begin{pmatrix} \sigma_1^2 & \rho\sigma_1\sigma_2 \\ \rho\sigma_1\sigma_2 & \sigma_2^2 \end{pmatrix} \quad (8)$$

is the **covariance** matrix, and $\rho \in (-1, 1)$ is the **correlation** (see **Correlation**) coefficient.

The marginal and conditional distributions of a bivariate normal random vector are univariate normal. For details, see Tong [5, Section 2.1].

Sampling Distributions and Related Applications

The sampling distribution results involve the distributions of the sample mean vector $\bar{\mathbf{X}}$, the sample covariance matrix $\mathbf{S}$, the independence property of $\bar{\mathbf{X}}$ and $\mathbf{S}$, and the distribution of the sample correlation coefficient. Those results may be applied for estimation and hypothesis testing purposes. For details, see Anderson [6, Section 2.3] and Tong [5, Sections 2.1 and 2.2].

The Multivariate Normal Distribution

Historical Remarks

The development of the multivariate normal distribution theory, which originated mainly from the studies of regression analysis and multiple and partial correlation analysis, was treated comprehensively for the first time by Edgeworth in his 1892 paper [7]. The development of the sampling distribution theory under the assumption of normality (such as Fisher's work on the distributions of sample correlation coefficients and *Hotelling's* T^2 distribution; see **Hotelling's** T^2) then followed.

Distributional Properties

An n -dimensional random vector $\mathbf{X} = (X_1, \dots, X_n)$ is said to follow a multivariate normal distribution with mean vector $\boldsymbol{\mu} = (\mu_1, \dots, \mu_n)$ and covariance

matrix $\boldsymbol{\Sigma}_{n \times n} = (\sigma_{ij})$, in symbols $\mathbf{X} \sim N_n(\boldsymbol{\mu}, \boldsymbol{\Sigma})$, if its joint probability density function is of the form

$$f(\mathbf{x}) = [(2\pi)^{n/2} |\boldsymbol{\Sigma}|^{1/2}]^{-1} \exp \left[-\frac{1}{2} Q_n(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) \right], \quad \mathbf{x} \in \mathbb{R}^n \quad (9)$$

where

$$Q_n(\mathbf{x}; \boldsymbol{\mu}, \boldsymbol{\Sigma}) = (\mathbf{x} - \boldsymbol{\mu}) \boldsymbol{\Sigma}^{-1} (\mathbf{x} - \boldsymbol{\mu})' \quad (10)$$

The marginal and conditional distributions of a multivariate normal random vector are also normal. For details, see Anderson [6, Sections 2.3, 2.4, and 2.5] and Tong [5, Section 3.3].

Sampling Distributions and Related Applications

The sampling distribution results also involve $\bar{\mathbf{X}}$, $\mathbf{S}$, and their independence property; in particular, the results are related to Hotelling's T^2 distribution and the Wishart distribution (generalizations of the Student's t distribution and χ^2 distribution, respectively). There also exist distributional results on the sample regression equations and various types of sample correlation coefficients. Such results have been found useful for the purposes of prediction and correlation analysis. For details on sampling distributions, see Anderson [6, Chapters 4, 5, and 7], Tong [5, Sections 3.4 and 3.5], and other related sources. For related applications in inference, a classical reference is Anderson [6].

References

- [1] Patel, J.K. & Read, C.B. (1982). *Handbook of the Normal Distribution*, Revised Edition 1996, Marcel Dekker, New York.
- [2] Galton, F. (1888). Co-relations and their measurements, chiefly from anthropometric data, *Proceedings of the Royal Society of London* **45**, 135–145.
- [3] Pearson, K. (1920). Notes on the history of correlation, *Biometrika* **13**, 25–45.
- [4] Pearson, K. (1896). Mathematical contributions to the theory of evolution – III. Regression, heredity and panmixia, *Philosophical Transactions of the Royal Society of London. Series A* **187**, 253–318.
- [5] Tong, Y.L. (1990). *The Multivariate Normal Distribution*, Springer-Verlag, New York.
- [6] Anderson, T.W. (1984). *An Introduction to Multivariate Statistical Analysis*, 2nd Edition, John Wiley & Sons, New York.

4 Normal Distribution

- [7] Edgeworth, F.Y. (1892). Correlated averages, *Philosophical Magazine, Series 5* **34**, 190–204.

Article originally published in Encyclopedia of Biostatistics, 2nd Edition (2005, John Wiley & Sons, Ltd). Minor revisions for this publication by Jeroen de Mast.

Related Articles

Normality Tests

Y.L. TONG

Normality Tests

Introduction

In testing normality, one should distinguish between the simple null hypothesis (parameters μ and σ^2 known) and the composite null hypothesis (parameters μ and σ^2 unknown). The simple null hypothesis is of little practical value.

Under the composite null hypothesis, the unknown parameters μ and σ^2 should be estimated. By using plug-in estimators, we may readily extend tests developed for the simple null hypothesis to the composite null hypothesis. However, plugging in estimators has a complicating effect on the null distribution of the test statistic. This makes it difficult to determine critical values and/or *p values* (see **P Values**).

Distance Methods

Distance methods are based on measuring the distance between a parametric and a nonparametric estimator of some function describing the distribution, such as the cumulative distribution function, the cumulative hazard function, or the moment generating function. In testing normality, the estimated normal cumulative distribution function and the empirical distribution function are typically used.

The asymptotic distribution of the composite null hypothesis versions of the Kolmogorov, the Cramér–von Mises, or the Anderson–Darling tests cannot be captured by a mathematical expression, and hence we have to resort to tabulation. The composite null hypothesis version of the Kolmogorov test is generally known as the *Lilliefors test*; [1] tabulates the null distribution of the test statistic.

Skewness and Kurtosis Tests

Instead of measuring the distance between the empirical distribution function and the estimated cumulative normal probability function, we may focus on certain characteristics of the **normal distribution**, such as the higher **moments**. For instance, it is known that the *skewness* (see **Skewness**) and *kurtosis* (see **Kurtosis**) of any normal distribution are equal to 0 and 3, respectively, and which suggests the use of

the third and fourth sample *moments* (see **Moments**) for testing normality. These tests are referred to as *directional*, as they lack the omnibus property of the distance-based tests. In [2] percentage points for various tests derived from skewness and kurtosis are given.

In econometrics, the Jarque–Bera test, see [3], is commonly used to test normality. This test is also derived from skewness and kurtosis.

Shapiro–Wilk-Type Tests

The normal *probability plot* (see **Probability Plots**) is a popular technique for exploring normality, and is a plot of the observations *versus* their corresponding ideal values, the so-called normal scores. If the data follow a normal distribution, the normal probability plot should be approximately linear. This suggests that we may test normality by quantifying the degree of nonlinearity in the normal probability plot.

The first test along these lines is the Shapiro–Wilk test, see [4]. Unfortunately, this test is computationally demanding. In particular, the Shapiro–Wilk test requires the inversion of a certain $n \times n$ **covariance** matrix V , where n denotes the sample size. In [5] it is argued that for large sample sizes, the covariance matrix V may be replaced by the identity matrix. This yields the Shapiro–Francia test, a simplified version of the Shapiro–Wilk test, which may be used for sample sizes larger than 50. In [6] it is recognized that the Shapiro–Francia test statistic is in fact the squared normal probability plot **correlation**, that is, the squared correlation between the observations and their normal scores; moreover, critical values for the normality probability plot correlation are given; see Figure 1. The Ryan–Joiner test in Minitab, see [7], only differs from the Shapiro–Francia test in the way the normal scores are obtained.

The close relation to the normal probability plot is an attractive aspect of the Shapiro–Wilk-type tests.

Comparison

In [8] results of a comparative study of various normality tests are given. Their findings include the following:

- the Shapiro–Wilk test provides a generally superior omnibus measure of nonnormality;

2 Normality Tests

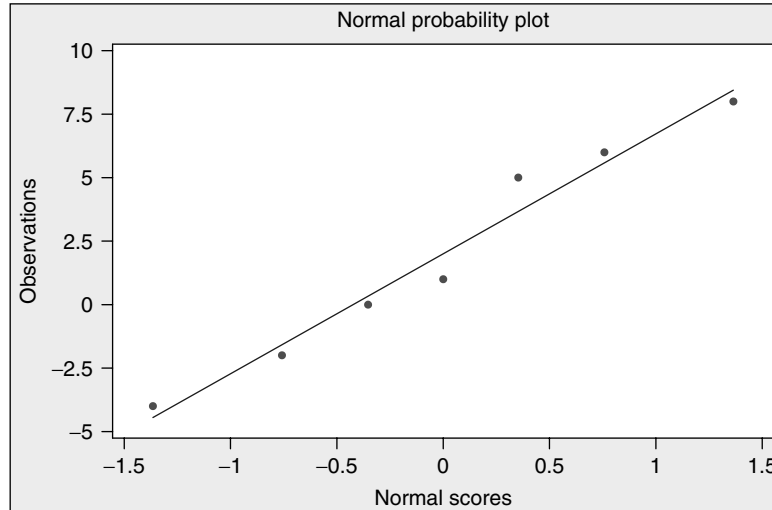


Figure 1 An application of the Shapiro –Francia test. The correlation between the observations 6, 1, –4, 8, –2, 5, 0 and their normal scores is 0.984. The corresponding critical value at the 5% **significance level** is 0.899 [see 6 Table 1]. As $0.984 > 0.899$, the normality hypothesis is not rejected

- the distance-based tests are typically very insensitive;
- a combination of both sample skewness and sample kurtosis usually provides a sensitive judgment, but even their combined performance is usually dominated by the Shapiro–Wilk test.

In [9], the superiority of the Shapiro–Wilk test is confirmed. However, it is pointed out that the findings in [8] are misleading with respect to the distance-based methods: the Cramér–von Mises and Anderson–Darling tests with plug-in estimators perform slightly worse than the Shapiro–Wilk test.

References

- [1] Lilliefors, H.W. (1967). On the Kolmogorov–Smirnov test for normality with mean and variance unknown, *Journal of the American Statistical Association* **62**, 399–402.
- [2] Bowman, K.O. & Shenton, L.R. (1975). Omnibus test contours for departures from normality based on $\sqrt{b_1}$ and b_2 , *Biometrika* **62**, 243–250.
- [3] Jarque, C.M. & Bera, A.K. (1980). Efficient tests for normality, homoscedasticity and serial independence of regression residuals, *Economics Letters* **6**, 255–259.
- [4] Shapiro, S.S. & Wilk, M.B. (1965). An analysis of variance test for normality: complete samples, *Biometrika* **52**, 591–611.
- [5] Shapiro, S.S. & Francia, R.S. (1972). Approximate analysis of variance test for normality, *Journal of the American Statistical Association* **67**, 215–216.
- [6] Filliben, J.J. (1975). The probability plot correlation coefficient test for normality, *Technometrics* **17**, 111–117.
- [7] Ryan, T.A. & Joiner, B.L. (1976). Normal probability plots and tests for normality, Technical report, The Pennsylvania State University, State College. URL <http://www.minitab.com/resources/articles/normprob.aspx>.
- [8] Shapiro, S.S., Wilk, M.B. & Chen, H.J. (1968). A comparative study of various tests for normality, *Journal of the American Statistical Association* **63**, 1343–1372.
- [9] Stephens, M.A. (1974). EDF statistics for goodness of fit and some comparisons, *Journal of the American Statistical Association* **69**, 730–737.

ALEX J. KONING

Observational Studies

An observational study is an empirical investigation that attempts to estimate the effects caused by a factor when it is not possible to perform an experiment. Whereas in an experiment, factor settings

are assigned to runs by the inquirer, in observational studies, the inquirer does not manipulate the factors, but merely records the values that they happen to have. Contrary to experimental data, observational data demonstrate **correlation** rather than causation (*see* **Causality**).

Outliers

Outliers are observations that stand apart from the majority of the observations. Outliers may be caused by reading or typing errors, contamination (data comes from more than one source), faults of the **measurement system**, or incidents in the process. It is important to detect outliers, because statistical analyses may be heavily affected (and even become invalid) when outliers are present. Basically, there are two ways to deal with outliers. One approach is to discard or reject outliers using tests of discordancy (preferably after additional checking of the circumstances under which they were obtained). Another approach is accommodation, i.e., to adjust analyses so that they become robust against possible outliers. The former approach comes down to deletion of the data point or replacement with the missing data value, while the latter approach involves the use of robust procedures such the median or the **trimmed mean** instead of the sample mean.

The box plot is a simple graphical tool to detect outliers. As a rule of thumb, observations that are more than 1.5 or 2 times the interquartile range (IQR, see **Quartiles**) away from the median of the data are marked in box plots as outliers by plotting them as crosses. Formal tests for outliers include the Dixon's tests and the Grubbs tests. Both tests assume normality (see **Normality Tests**). Dixon's test statistic for detecting one outlier on the right is given by $(x_{(n)} - x_{(n-1)}) / (x_{(n)} - x_{(1)})$, where $x_{(i)}$ denotes the i th order statistic, i.e., the i th smallest observation. Similar test statistics exist for one or two outliers on the left or on both sides. Thus $x_{(n)}$ is the largest observation and possibly an outlier. Critical

values of the Dixon's test must be looked up in a statistical table as there is no closed-form formula for them.

The Grubbs tests are based on a transformation of the maximum studentized value of the sample, i.e., the largest among the values $(x_i - \bar{x})/s$, where $\bar{x}$ denotes the sample mean and s the sample standard deviation. The Grubbs test statistic for detecting one outlier on the right is given by $(x_{(n)} - \bar{x})/s$, with critical value $(n-1)t_{n-2,\alpha/n} / \sqrt{n(n-2 + t_{n-2,\alpha/n}^2)}$ where $t_{n-2,\alpha/n}$ is the value such that a Student's t -distribution with $n-2$ **degrees of freedom** has probability α/n of exceeding it.

Caution should be exercised when using such tests, because they may behave poorly for other alternative hypotheses. For example, Dixon's test for one outlier on the right may have poor **power** when there are two outliers on the right. We refer to [1] for a comprehensive overview on this topic and to [2] for a shorter, practical guide.

References

- [1] Barnett, V. & Lewis, T. (1994). *Outliers in Statistical Data*, John Wiley & Sons, New York.
- [2] Iglewicz, B. & Hoaglin, D.C. (1993). *How to Detect and Handle Outliers*, *ASQC Basic References in Quality Control*, ASQ Quality Press, Milwaukee, Vol. 16.

Related Article

Normality Tests.

ALESSANDRO DI BUCCHIANICO

Pareto Chart

Pareto Analysis

Pareto analysis is a problem-solving technique for prioritizing which root causes for a given problem should be handled first. The main idea behind the technique is the empirical observation that 80% of the problems come from 20% of the causes. This empirical observation was named by Juran [1] after the Italian econometrician Pareto (1848–1923), who made this observation during his studies of the distribution of wealth. The 80–20 rule has later been shown to apply in many other contexts. Pareto analysis is one of Ishikawa’s seven basic quality tools.

The Pareto chart is a bar chart where the root causes or categories are shown ordered from high importance or frequency to lower importance or frequency. Pareto analyses may only be performed on data collected in a stable situation, so no trends may be present. Depending on the context, causes may be weighted by frequency, costs, impact. If causes are aggregated into categories, then care must be taken to ascertain that categories allow fair comparisons.

An example adapted from [2] shows defects found in printed circuit boards. The bars indicate the frequency of the errors ordered from left (most frequent) to right (less frequent). As can be seen from the scale on the right-hand side of Figure 1, the three most frequent defects account for slightly over 80% of the defects.

References

- [1] Juran, J.M. (1951). *Juran’s Quality Control Handbook*, 1st Edition, McGraw-Hill.

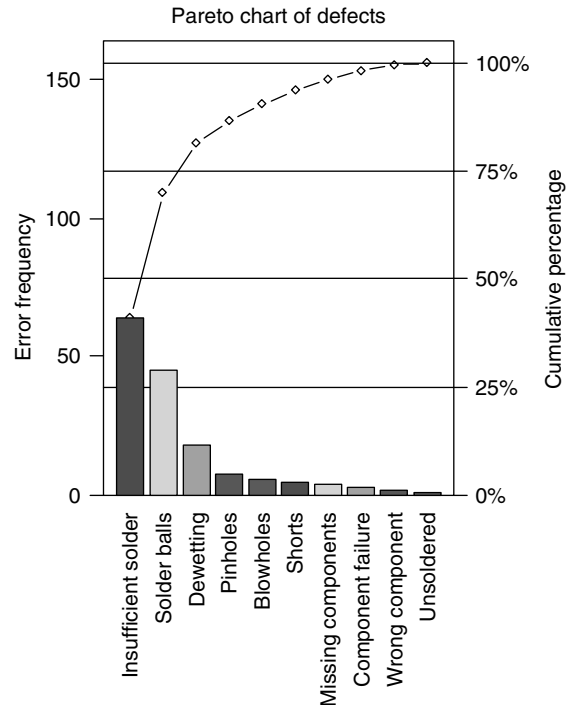


Figure 1 Pareto chart of printed circuit defects

- [2] Montgomery, D.C. & Runger, G.C. (2006). *Applied Statistics and Probability for Engineers*, 4th Edition, John Wiley & Sons.

Related Articles

Cause-and-Effect Diagrams; Failure Modes and Effects Analysis; Graphical Representation of Data.

ALESSANDRO DI BUCCHIANICO

Pooled Variance, Pooled Estimate

When estimators are based on more than one sample, then pooling is the name for obtaining an estimator for a characteristic of the combined samples by combining estimates based on the individual samples. The reason for doing this is that a pooled estimator is more accurate because it is based on more observations. The basic example is the pooled *variance* that appears in **confidence intervals** and tests for the difference of the mean of two samples from *normal distributions* with the same variance. In this case the sample *variance* of the combined sample may be a biased estimator (*see* **Bias of an Estimator**) of the common variance, because it is a measure of deviation from the overall mean (which may be different from the means of the individual samples). The correct way is to use the pooled variance $((n_1 - 1)S_1^2 + (n_2 - 1)S_2^2)/(n_1 + n_2 - 2)$, which is an unbiased estimator of the common *variance* of the two samples. Note that pooling means in this case is not a problem, since the pooled mean $((n_1 - 1)\bar{X}_1 + (n_2 - 1)\bar{X}_2)/(n_1 + n_2)$ equals

the overall mean of the combined sample. However, see [1] for a very readable, practical discussion whether to pool or not. More generally, estimators are pooled in *analysis of variance*.

In case one is not certain whether it is allowed to pool the data, statistical tests may be performed to decide whether to pool or not. Such estimators are known under the name sometimes-pool estimator (see [2, 3] for further details).

References

- [1] Mosteller, F. (1948). On pooling data, *Journal of American Statistical Association* **43**, 231–242.
- [2] Bancroft, T.A. (1944). On biases in estimation due to the use of preliminary tests of significance, *Annals of Mathematical Statistics* **15**, 190–204.
- [3] Han, C.P. & Bancroft, T.A. (1968). On pooling means when variance is unknown, *Journal of American Statistical Association* **63**, 1333–1342.

Related Articles

Bias of an Estimator; Confidence Intervals.

ALESSANDRO DI BUCCHIANICO

Power

The power of a *hypothesis test* (see **Hypothesis Testing**) is the probability that the null hypothesis H_0 is rejected when some other hypothesis H is true. Suppose the hypotheses are about a parameter θ , and suppose H_0 specifies $\theta = \theta_0$, then the power function is defined as $\beta(\theta) = P_\theta(H_0 \text{ rejected})$. The

significance level of the test is $\alpha = \beta(\theta_0)$. For $\theta \neq \theta_0$, $\beta(\theta)$ gives the probability that H_0 is correctly rejected given the effect size $\theta - \theta_0$.

Related Article

Sample-Size Determination.

Probability Density Function (PDF)

Probability density functions (PDFs) are mathematical formulas that provide information about how the values of random variables are distributed. For a discrete random variable X assuming values in a sample space Ω , the density is usually named the *frequency distribution*, defined as $p_X(x) = P(X = x)$, for each x in Ω . If X is a continuous random variable, its density function f_X determines the probability that X is in an interval $[a, b]$,

namely,

$$P(a \leq X \leq b) = \int_a^b f_X(x) \, dx \quad (1)$$

In particular, the **PDF** determines the probability that X is smaller than x , namely, $F_X(x) = P(X \leq x) = \int_{-\infty}^x f_X(x) \, dx$. The function F_X is called the *(cumulative) distribution function of X* (see **Cumulative Distribution Function (CDF)**). The best-known density function is the Gauss curve, which defines the **normal distribution**. (See **Probability Density Functions**.)

Statistical Methods for Counts, Rates, and Proportions

Introduction

Counting, or enumerating, is one of the most basic arithmetic operations. You may count the number

1. of specified “events” occurring in time, space, or volume, like customers arriving at a service counter, flaws on woven fabric, bacteria in a suspension, or events occurring on a discrete item, like stains on a painted door;
2. of discrete items (sample units) possessing one of a specified set of exclusive and exhaustive categories, like pass/fail, or “never”, “rarely”, “frequently”.

In both cases the result is integer valued. In the first case it is common practice to report the number of events or the size-adjusted *rate* of occurrences. In the latter case, the integer-valued *frequency* of occurrence of the category is reported. When the totality of units being investigated is known, the *proportion* of units of the category may also be reported.

Basic Sampling and Distribution Models

Counts of Events. When the events are randomly dispersed in time or space, the distribution of the counts in 1 may be modeled by a *Poisson distribution* with a mean value that is proportional to the size (time span, length, area, or volume) of the observational unit (see **Poisson Processes**). The constant of proportionality is termed the *rate of occurrence*, or *incidence*.

Frequencies of Occurrence. When the classification of the individual items are independent of each other, with a constant probability of belonging to a specific class, then the *multinomial distribution* will model the frequencies effectively. When there are only two classes, the distribution is the *binomial* distribution. The sample *proportion* of observations

belonging to the category in question is an unbiased estimator of the probability.

The binomial and Poisson distributions are one-parameter distributions belonging to the so-called exponential families of distributions. The distributions are skewed with a variance that depends on the mean value. The **skewness** is less pronounced for large values of the mean. When the assumption of independence is violated, this may give rise to *overdispersion*, see below.

Transformations. Traditionally various transformations have been used to achieve approximate normality and **homoscedasticity** for count data such that classical statistical methods could be used. Thus, the *arcsine* was used for observed proportions, and for Poisson counts *square root* is used to achieve homoscedasticity, or $X^{2/3}$ for normality.

However, with more computing power available, more adequate analytical tools have been developed. Instead of transforming the data, contemporary tools respect the distributional shape of the observations and analyze linear models for some transformation of the mean value. The most prominent, the log-linear models for Poisson counts and the logit models for binomial proportions, will be introduced below.

Bayesian Inference

As the likelihood function corresponding to Poisson data is proportional to the **probability density function** of a *gamma distribution*, it follows that the natural conjugate prior distribution for Poisson data is a gamma distribution, and hence, the posterior distribution is again a gamma distribution.

A similar argument shows that the natural conjugate prior distribution for binomial and multinomial data is a *beta* and a multivariate beta distribution, respectively.

Contingency Tables

When events are counted, and classified by two or more categorical variables, the counts are often presented in tables of counts, each variable determining a “dimension” of the table. Similar presentations of frequencies are used when sample units are classified

2 Statistical Methods for Counts, Rates, and Proportions

by two or more categorical variables. Such a table is often referred to as a *contingency table*.

Two-Dimensional Tables

When there are only two dimensions in the table, a hypothesis of no association between the two categories may be formulated as a test for independence, or a test of homogeneity depending on the sampling method being used. In either case, the hypotheses is tested by Pearson's χ^2 test. For a 2×2 table Fisher's exact test is used (*see Parametric Tests*).

Log-Linear Models for Multiway Tables

The so-called log-linear models for multiway tables provide a versatile framework for analyzing association in multidimensional contingency tables. The models are appropriate for counts under unrestricted Poisson sampling and for frequencies under multinomial sampling. The idea is to formulate a *linear model* for the logarithm to the expected counts in individual cells. For a two-dimensional table, the log-linear model is

$$\ln(\lambda_{ij}) = \mu + \alpha_i + \beta_j + \gamma_{ij} \quad (1)$$

with suitable linear restrictions on the parameters.

A hypothesis of no association between the classifying variables is formulated in terms of the *interaction* term γ_{ij} as $H_0 : \gamma_{ij} = 0$.

For multidimensional tables a set of hierarchical (nested) models is formulated, each model spanned by a corresponding set of *sufficient marginals*. The approximate χ^2 distribution of the likelihood-ratio test statistic is utilized to reduce the model by successively removing higher-order interaction terms in analogy with **ANOVA** models for normally distributed observations (*see Analysis of Variance*).

Log-linear interaction models provide the most general terms for relationships among variables. However, the variables investigated by log-linear models are all treated as response variables, and no distinction is made between independent and dependent variables.

Logit Models, Odds

Modeling effects of explanatory variables on the response may be more important than modeling relationships between all variables. A subset of the log-linear models, the so-called logit models for the response variable serves this purpose.

Consider an $r \times c$ **factorial** experiment with a binary response, and let p_{ij} denote the probability of "success" in experiment (i, j) .

The *logit* transformation of a probability, p , is the logarithm to the corresponding *odds*

$$\eta(p) = \ln\left(\frac{p}{1-p}\right) \quad (2)$$

In analogy with equation (1) we now introduce the linear decomposition of the logit η_{ij}

$$\eta_{ij} = \eta(p_{ij}) = \mu + \alpha_i + \beta_j + \gamma_{ij} \quad (3)$$

with suitable linear restrictions on the parameters. The term γ_{ij} is called the *interaction term* in this linear logit model.

When the response is modeled by a multinomial distribution, a generalization of the logit approach may be utilized, see [1].

Graphical Models

The more dimensions a contingency table has, the greater the need for structuring the table in an operational, and interpretable manner. Complex stochastic models can often be constructed by combining simpler models through conditioning. This principle is utilized in the *graphical models*.

Graphical models are log-linear interaction models for contingency tables that can be represented by a simple undirected graph with as many nodes as the table has dimensions, and where the edges between the nodes represent conditional dependencies. It is a particular property for graphical models that they can be given an interpretation in terms of conditional independence in the sense that they allow for a factorization of the joint probability into a product of conditional distributions. The interpretation can be read directly off the graph in the form of a Markov property.

Graphical models with undirected edges are generally called *Markov random fields*, or *Markov networks*. See [2] for an introduction to graphical models. A more advanced text is [3].

Paired Samples. A particular situation arises when the observations are *paired*. This will typically be the case in before and after studies, or when two treatments are applied to each experimental unit. In these cases it is often appropriate to investigate a hypothesis of *marginal symmetry*. When the response is binary, the test reduces to a comparison of the off-diagonal terms with a binomial distribution with $p = 0.5$. The test is often referred to as *McNemar's test*.

Models for paired samples are also used for assessment of agreement between raters. Various metrics have been suggested to quantify agreement, see [4].

Correspondence Analysis

Correspondence analysis is a descriptive/exploratory technique designed to analyze simple two-way and multiway tables containing some measure of correspondence between the rows and columns. The results provide information which is similar in nature to those produced by factor analysis techniques, and they allow one to explore the structure of categorical variables included in the table, see [5] and [6].

Generalized Linear Models

The log-linear models and logit models considered above are special cases of the so-called **generalized linear models**. This class of models extends the *general linear model* for normally distributed variables to models for the mean-value structure for distributions belonging to an exponential family, like Poisson, binomial, or multinomial distributions.

The fundamental idea underlying these models is that instead of formulating linear models for the mean value itself, linear models are formulated for a function of the mean value, the *link function*. Hypotheses are tested by likelihood-ratio tests thereby taking the distributional properties of the observations into consideration. The contribution to the likelihood-ratio test statistic from an observation y_i , and the value, $\hat{\mu}_i$, estimated under the model is called the *deviance*, $d(y_i; \hat{\mu}_i)$. Thus, the deviance is a measure of how well an estimated model fits the observations in terms of log likelihood.

For the binomial distribution with **sample size** n , the deviance is

$$d(y; \hat{p}) = 2n \left\{ y \ln \left(\frac{y}{\hat{p}} \right) + (1 - y) \ln \left(\frac{(1 - y)}{(1 - \hat{p})} \right) \right\} \quad (4)$$

with y denoting the observed proportion. The deviance for Poisson data is

$$d(y; \hat{\lambda}) = 2 \left\{ y \ln \left(\frac{y}{\hat{\lambda}} \right) - (y - \hat{\lambda}) \right\} \quad (5)$$

The *residual deviance* is the sum of the individual terms $d(y_i; \hat{\mu}_i)$. The null distribution of the residual deviance is a χ^2 distribution. When a set of hypotheses is *nested* within each other, the residual deviance may be *partitioned* into contributions from each level just like the partitioning of the total **sum of squares** in an ANOVA context.

The *canonical link* is the function that transforms the mean value into the *canonical parameter* for the exponential family. For the Poisson distribution, the canonical link is the logarithm, and for the binomial distribution it is the logit. Linear models for the canonical link allow simple sufficient statistics for the parameters in the model.

We observe that the log-linear models for Poisson data, and the logit models for binomial data above correspond to modeling the canonical parameter by linear models in the explanatory variables.

Also, it should be noted that other link functions than the canonical link may be used. Thyregod and Spliid [7] give an example of an assembly situation where the logarithmic link provides a better, and more meaningful fit, for binomially distributed data.

An introduction to generalized linear models is given in [8] and [9], see also [10]. For applications in quality-improvement experiments, see [11].

More advanced statistical software has options for analysis of generalized linear models under various link functions.

Regression Models

Generalized linear models with continuous explanatory variables are sometimes referred to as *regression models*.

4 Statistical Methods for Counts, Rates, and Proportions

Poisson Regression

For Poisson data, the canonical link is the *logarithm*, and the log-linear regression model has the form

$$\ln(\lambda_i) = \alpha + \beta x_i \quad (6)$$

corresponding to

$$\lambda_i = \exp(\alpha + \beta x_i) \quad (7)$$

This model is often called *Poisson regression* model. Sometimes, when an *exposure* variable, t_i is present, the log exposure enters as an *offset* variable,

$$\ln(\lambda_i) - \ln(t_i) = \alpha + \beta x_i \quad (8)$$

Logistic Regression

For binomial data, the canonical link is the *logit*, and the linear model for the logit is

$$\ln\left(\frac{p_i}{1 - p_i}\right) = \alpha + \beta x_i \quad (9)$$

corresponding to the expression

$$p_i = \frac{\exp(\alpha + \beta x_i)}{1 + \exp(\alpha + \beta x_i)} \quad (10)$$

for the response probability. The function (10) is the *logistic function*, and hence the model is termed *logistic regression*.

The logistic regression is used to describe *dose-response* relationships when the response is binary. The model may be extended to cover multinomial response distribution, including models for *ordered categorical response*, see [1].

Dedicated statistical software often allows for other choices of link functions for regression models for binary responses. An example is the *complementary log-log* link

$$\ln\{-\ln[1 - p(x)]\} = \alpha + \beta x \quad (11)$$

corresponding to the dose-response function

$$p(x) = 1 - \exp[-\exp(\alpha + \beta x)] \quad (12)$$

The function grows slower away from 0, and approaches 1 faster than the logistic function.

Compound Models, Overdispersion

When there is greater variation in data than predicted by the sampling model, we speak of *overdispersion*.

Often for count data, several events are recorded on the same unit, such that the events are interdependent. This might result in overdispersion. Lee and Nelder [12] discuss the two ways of modeling overdispersion, either by introducing a constant factor, a *dispersion parameter* to the variance, or by using a *hierarchical model* where the mean of the counts is itself a random variable with a specified distribution.

In the hierarchical generalized linear models, the mixing distribution is chosen as the *natural conjugate* to the sampling distribution, see [13]. For Poisson data one would use a gamma distribution as compounding distribution, and the resulting marginal distribution would then be a *negative binomial* distribution. For binomial data, a beta distribution would be used, resulting in the so-called beta-binomial distribution of the counts.

Another common possibility is to use a normal compounding distribution for the canonical parameter (i.e., for the logit (binomial), and for $\ln(\lambda)$ in the Poisson case). This approach is used in the so-called generalized linear mixed models.

References

- [1] Agresti, A. (2002). *Categorical Data Analysis*, 2nd Edition, John Wiley & Sons, New York.
- [2] Borgelt, C. & Kruse, R. (2002). *Graphical Models, Methods for Data Analysis and Mining*, John Wiley & Sons, Chichester.
- [3] Cowell, R.G., Lauritzen, S.L. & Spiegelhalter, D.J. (1999). *Probabilistic Networks and Expert Systems*, Springer-Verlag, New York.
- [4] Fleiss, J., Levin, B. & Paik, M.C. (2003). *Statistical Methods for Rates and Proportions*, 3rd Edition, John Wiley & Sons, New York.
- [5] Greenacre, M.J. (1984). *Theory and Applications of Correspondence Analysis*, Academic Press, London.
- [6] Greenacre, M.J. (1993). *Correspondence Analysis in Practice*, Academic Press, London.
- [7] Thyregod, P. & Spliid, H. (1991). A hypothesis of no interaction in factorial experiments with a binary response, *Scandinavian Journal of Statistics* **18**, 197–209.
- [8] Myers, R.H., Montgomery, D.C. & Vining, G.G. (2001). *Generalized Linear Models: With Applications in Engineering and the Sciences*, John Wiley & Sons, New York.
- [9] Lindsey, J.K. (1997). *Applying Generalized Linear Models*, Springer-Verlag, New York.
- [10] Fahrmeir, L. & Tutz, G. (2001). *Multivariate Statistical Modelling Based on Generalized Linear Models*, 2nd Edition, Springer-Verlag, New York.

- [11] Lee, Y. & Nelder, J.A. (1998). Generalized linear models for the analysis of quality-improvement experiments, *Canadian Journal of Statistics* **26**, 95–105.
- [12] Lee, Y. & Nelder, J.A. (2000). Two ways of modelling overdispersion in non-normal data. Applied statistics, *Journal of the Royal Statistical Society. Series C* **49**, 591–598.
- [13] Lee, Y. & Nelder, J.A. (2001). Hierarchical generalized linear models: a synthesis of generalized linear models, random-effect models and structured dispersions, *Biometrika* **88**, 987–1006.

POUL THYREGOD

***P* Values**

A *P* value is frequently used to report the result of a *hypothesis test* (see **Hypothesis Testing**). A *P* value is defined as the probability, calculated under the null hypothesis, of obtaining a test result as extreme as that observed in the sample. Let the null hypothesis H_0 state that $\theta = \theta_0$, and let T be

a sample statistic, then $P = P_{\theta_0}(T \geq t)$, or $P = P_{\theta_0}(T \leq t)$, or $P = P_{\theta_0}(|T| \geq t)$ (depending on the direction specified by the alternative hypothesis). In the classical approach, one rejects H_0 if the *P* value is smaller than the desired *significance level* (see **Significance Level**) α of the test. Common choices for the significance level are 0.1, 0.05, and 0.01. See also **Power**.

Quartiles

For a continuous variable X , with distribution function F_X , the first, second, and third quartiles are $Q_1 = F^{-1}(0.25)$, $Q_2 = F^{-1}(0.50)$, and $Q_3 = F^{-1}(0.75)$

(that is, the values below which X will be with 25%, 50%, or 75% probability respectively). Note that the second quartile is equal to the median. *See also* **Quantiles**.

Residuals

Residuals could be said to be the difference between the observed values and the fitted values. In the context of linear models $\mathbf{Y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon}$, the ordinary

residuals are defined as $\mathbf{e} = \mathbf{Y} - \mathbf{X}(\mathbf{X}'\mathbf{X})^{-1}\mathbf{X}'\mathbf{Y}$. Other types of residuals are standardized (or studentized) residuals and predicted residuals. Residuals are used for checking model assumptions and detecting **outliers**.

Runs, Runs Tests

A *run* is an uninterrupted sequence of observations all of which satisfy a certain criterion. For example, the sequence 00011010111 has a run of three 0's, then a run of two 1's, and so on. Similarly, we can speak of a run above average, a run below average, a "run up", and a "run down". In the sequence 5, 4, 6, 8, 10, 12, 11 there is a run up of length 4 since there are four increases in a row.

There are tests for randomness based on runs, and in particular in **quality control** runs have long been used as a criterion to detect when a process is out of control. Deming [1] suggested that a run up (or down) of length above six is indicative of a drift in the process. This and other runs tests (or "runs rules") are often used in Shewhart *control charts* (see **Control Charts, Overview**), in particular because the standard Shewhart chart is rather insensitive to small shifts. Nelson [2] lists eight tests for special causes, seven of which are based on runs. These tests are applicable for $\bar{X}$ and X (individuals) charts. They are often referred to as the *Western Electric rules*, as they were given in the company's 1956 *Statistical Quality Control Handbook*.

1. nine points in a row on the same side of the average,
2. six points in a row steadily increasing or decreasing (that is, a run up or a run down of length five),
3. fourteen points in a row alternating up and down,
4. two out of three points in a row more than two times the standard deviation from the average (on the same side),
5. four out of five points in a row more than the standard deviation from the average (on the same side),

6. fifteen points in a row within one time the standard deviation from the average (on either side),
7. eight points in a row more than the standard deviation from the average (on either side).

The run-length distributions of these tests are given (for a range of scenarios) in Champ and Woodall [3, 4], Palm [5], Walker *et al.* [6], and Shmueli and Cohen [7].

References

- [1] Deming, W.E. (1972). Making things right, in *Statistics: A Guide to the Unknown*, J. Tanur, F. Mosteller, W.H. Kruskal, eds, Holden-Day, San Francisco.
- [2] Nelson, L.S. (1984). The Shewhart control chart – tests for special causes, *Journal of Quality Technology* **16**(4), 237–239.
- [3] Champ, C.W. & Woodall, W.H. (1987). Exact results for Shewhart control charts with supplementary runs rules, *Technometrics* **29**(4), 393–399.
- [4] Champ, C.W. & Woodall, W.H. (1997). Signal probabilities of runs rules supplementing a Shewhart control chart, *Communications in Statistics: Simulation and Computation* **26**, 1347–1360.
- [5] Palm, A. (1990). Tables of run-length percentiles for determining the sensitivity of Shewhart control charts for averages with supplementary runs rules, *Journal of Quality Technology* **22**, 289–298.
- [6] Walker, E., Philpot, J.W. & Clement, J. (1991). False signal rates for the Shewhart control chart with supplementary runs tests, *Journal of Quality Technology* **23**(3), 247–252.
- [7] Shmueli, G. & Cohen, A. (2003). Run-length distribution for control charts with runs and scans rules, *Communications in Statistics: Theory and Methods* **32**(2), 475–495.

JEROEN DE MAST

Sample-Size Determination

Introduction

The objective of most inquiries is to collect data that may be turned into useful information for decision making. In a deterministic world, a single observation would most often suffice to provide the required information. For example, to determine the velocity of an object moving at a constant speed, a single observation of the distance traveled over a certain time would provide the required value of speed. Most often, however, the object of inquiry behaves randomly. This randomness might be captured by a statistical distribution (with constant parameters). Most often this distribution is unknown to us. Therefore, the conclusions derived from the data would frequently be associated with uncertainty. This uncertainty needs to be quantified so that a proper sample size be determined, in advance, which guarantees that uncertainty in the conclusions does not exceed prespecified thresholds. For example, a sample size may be required that would ensure a sample estimate of the proportion of rejects in the process that does not deviate more than 10% from the correct value, with confidence of at least 99%.

Occasionally, in addition to an element of randomness, the studied process itself may not be stable. For example, its mean may not be constant. In such cases, a relational model needs to be formulated, which will capture the observed systematic variation. Since the latter is often “clouded” with random variation (unlike the speed example given earlier), the ability to build a reliable model is often dependent on the amount of data collected. Again, sample size becomes a critical element that would determine the quality of the conclusions reached.

Finally, conclusions derived from observations may themselves be associated with costs that one wishes to minimize. For example, in **statistical process control** (SPC), not detecting lack of control, due to too small a sample size, may incur costs that one wishes to avoid. Formulating a model that comprises all costs associated with the sampling

process may be helpful in determining an optimal sample size that minimizes the expected total cost.

In this article, we address four common scenarios where a sample size needs to be determined: **hypothesis testing** (“Sample-Size Determination in Hypothesis Testing”), parameter **estimation** (“Sample-Size Determination in Parameter Estimation”), relational modeling (“Sample-Size Determination in Relational Modeling”), and optimal sampling planning (“Optimized Sample-Size Determination”).

Sample-Size Determination in Hypothesis Testing

In hypothesis testing (*see* **Hypothesis Testing**), there are two mutually exclusive claims about the true “state of nature”. These competing claims are described by the null hypothesis, which embodies the current state of our knowledge, and the alternative hypothesis, which reflects the claim that one wishes to test. Since the randomness in the data collected may lead to the adoption of a wrong hypothesis, the objective of sample-size determination is to guarantee that the probabilities to err do not exceed prespecified maximal values. The latter are denoted by α and β , respectively, for errors of the first and the second types. Quite often, the calculation of these probabilities to determine an appropriate sample size requires knowledge of certain parameters of the assumed data-generating distribution. Therefore, a preliminary small sample of observations is occasionally needed to provide initial estimates of the unknown parameters (most often those associated with the variance). The following numerical example demonstrates this.

Numerical Example

A certain improvement effort is scheduled for a chemical process, aimed to maximize its daily yield. A designed experiment is planned, according to which two samples of daily yields will be collected. One sample would be taken from the presently running process, and another would be taken after the implementation of an improvement effort. It is assumed that the variance of daily yield does not change as a result of the process of improvement. The two samples are to be of equal size, and it is desired

2 Sample-Size Determination

that testing the effectiveness of the improvement effort will guarantee maximum error probabilities of $\alpha = 1\%$, $\beta = 0.5\%$. Identification of an increase in yield of 10% (or more) is desired (smaller increase is regarded as inconsequential).

The required sample size can be determined as follows. Let μ_1 and μ_2 be the means of the process's yield before and after the improvement. It is required that the test has **power** of at least $1 - \beta = 0.995$ to identify an increase in yield of at least 10%. The two hypotheses are as follows:

$$H_0 : \mu_1 = \mu_2; H_1 : \mu_2 - \mu_1 \geq 0.10\mu_1 \quad (1)$$

Assume that the standard deviation of the daily yield, σ , is known, and that the sample average can be considered normally distributed. Under H_0 , the critical value for the sample average, above which H_0 will be rejected, is

$$\bar{X}_{cr} = \mu_1 + \frac{\sigma}{\sqrt{n}} Z_{1-\alpha} \quad (2)$$

where Z_p is the p th quantile of the standard **normal distribution** (100p% of the distribution does not exceed this value). Similarly, under H_1 :

$$\bar{X}_{cr} = 1.1\mu_1 + \frac{\sigma}{\sqrt{n}} Z_\beta \quad (3)$$

From these two equations, the required sample size is given by:

$$n = \left[\frac{\sigma(Z_{1-\alpha} - Z_\beta)}{0.1\mu_1} \right]^2 \quad (4)$$

Note, that equation (4) requires knowledge of σ . If the latter was unknown, a preliminary study would be needed that would provide the required sample estimate.

Sample-Size Determination in Parameter Estimation

In parameter estimation (*see Estimation*), a point estimate is derived from the data, which has some desirable characteristics (like unbiasedness, consistency, or minimum variance). In addition, an associated **confidence interval** (*see Confidence Intervals*) is constructed, which defines, with a prespecified confidence of $1 - \alpha$, an interval around the estimate, where the true value of the parameter is likely to be.

Specifying the desirable precision of the point estimate (the width of the confidence interval) and the required confidence level ($1 - \alpha$), the needed sample size is easily calculated from the distribution of the sample estimate.

Numerical Example

An estimate of the variance of the lifetime of a certain electronic component is required for **reliability analysis**. The following variance estimate, based on a sample of lifetime observations of n items, is used:

$$\hat{\sigma}^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2 \quad (5)$$

where $\bar{X}$ is the sample mean.

We determine a sample size n such that the upper limit (UL) of the confidence interval is no more than 25% larger than the lower limit (LL), while $\alpha = 10\%$. Assume that lifetimes may be considered as approximately normally distributed. It is known that $\hat{\chi}^2 = (n-1) \frac{\hat{\sigma}^2}{\sigma^2}$ is χ^2 distributed with $n-1$ **degrees of freedom** (σ^2 is the unknown true lifetime variance). Therefore a confidence interval may be defined by

$$LL = \frac{(n-1)\hat{\sigma}^2}{\chi_{1-\alpha/2}^2(n-1)} < \sigma^2 < \frac{(n-1)\hat{\sigma}^2}{\chi_{\alpha/2}^2(n-1)} = UL \quad (6)$$

where $\chi_p^2(n-1)$ is the p th quantile of a χ^2 distribution with $n-1$ degrees of freedom. It is required that

$$\frac{UL}{LL} = 1.25 = \frac{\chi_{1-\alpha/2}^2(n-1)}{\chi_{\alpha/2}^2(n-1)} \quad (7)$$

A simple search with some available software yields $n = 436$ (for $\alpha = 0.10$).

Sample-Size Determination in Relational Modeling

Relational modeling refers here to empirical modeling *via* multiple linear regression. The issue of sample size needed to obtain a reliable linear regression model is seldom addressed extensively in the literature (refer, for example, to [1, 2]), probably due

to the diversity of relevant considerations that may be involved. We expound some of these considerations.

A first consideration is whether the observations are collected in a designed experiment or in an observational study (see **Observational Studies**). In the former case, design optimality may affect to a large extent the amount of data required. An **optimal design** reduces the required sample size. For example, if the design aspires to reduce the variance of the required estimate (derived from the model), then design optimality would mean a reduced sample size. In the latter case (observational studies), there are far less opportunities to reduce the required sample size by the sampling procedure and the design of the **data collection** process. This implies that we can only *calculate* the required sample size, dependent on our objective.

In a more general perspective, sample-size determination in empirical modeling (like linear regression) depends on the modeling objective. Two major considerations, related to the quality of the model, are goodness of fit and stability. One may overfit the model, which implies that too many terms are included in the model relative to the available sample size, or one may underfit the model (underuse of the available information). In both cases, the model may not appropriately reflect the information provided by the available sample and either the model needs to be changed, or a larger sample size is required (when overfitting). A good indicator for proper model fitting is the AIC statistic (or the corrected AIC_c), which one wishes to minimize (for example, when several alternative models are examined). Good sources to learn of this consideration are [3, 4].

A second consideration is the model's stability. If a fitted model proves to be of low stability, a larger sample size is required to obtain satisfactory accuracy for the model parameters' estimates. The stability of a model may be learned from the relative widths of the confidence intervals associated with the parameters' estimates (the width of the confidence interval, divided by the parameter's point estimate, should not be larger than 1, and preferably much smaller than that). Another good indicator for stability is prediction error **sum of squares** (PRESS; refer to [5] for details). This statistic essentially measures how well the model can provide *interpolated* predicted values. Thirdly, good *extrapolation* capability testifies that the model is stable and therefore

the sample size, required to estimate its parameters, can be relatively small. A source for model instability may be high intercorrelations between the estimates of parameters (a phenomenon known as *ill-conditioning* or *multicollinearity*; see **Collinearity**). Designed orthogonality between the regressor variables when data are collected (for example, by use of **orthogonal arrays** in the experimental design) may circumvent this problem.

Other considerations and requirements may be valid too. For example, a desirable accuracy for the model's predicted value, for certain values of the regressors, may dictate a certain necessary sample size (the latter appears in the computational formula for the associated confidence interval). Further discussion of this vast area of research may be found in the above sources and the references therein.

Optimized Sample-Size Determination

In previous sections, we addressed sample sizes needed to fulfill certain constraints regarding the precision of conclusions derived from sample data. In this section, we demonstrate that sample size may be determined on the basis of optimization considerations as well. In quality and reliability engineering, the problem of optimizing sample-size determination has been addressed extensively, both with respect to sampling plans (in the context of **acceptance sampling**) and with respect to SPC. Both issues are usually dealt with, to some extent, in any textbook about statistical **quality control** (for example, [6, 7]). The reader is referred to these sources. In this section, we demonstrate how optimized sample size can be derived in the context of hypothesis testing.

Numerical Example

A certain design change in a certain product is suggested. It is believed that as a result of the change, the lifetime of the product will increase. If that increase is 50% or more, an increase of \$200 in revenue per unit is expected. Implementing the design change in the manufacturing facility is evaluated to cost \$10 per unit. It is suggested that the design change will first be tested off-line. Producing a unit in the experimental line will cost \$1000. Currently, annual production is 50 000 units. The optimized

4 Sample-Size Determination

sample size will be determined for a planning horizon of 1 year.

We determine how many units should be produced experimentally (off-line) before a decision is taken as to whether the design change is implemented. The problem can be stated as that of optimal hypothesis testing. Denoting the mean lifetime after the design change by μ and the current mean lifetime by μ_0 , the hypotheses are

$$H_0 : \mu = \mu_0; H_1 : \mu \geq 1.5\mu_0 \quad (8)$$

Let α be the probability of wrongly adopting the design change (in fact, the latter does not extend the product lifetime by the desirable 50% or over). Similarly, denote by β the probability of wrongly rejecting the design change. Let n be the sample size. The annual total cost includes the following:

- Sampling cost: $\$(1000)n$.
- Annual lost profit for wrongly accepting H_0 : $\$(200 - 10)(50\,000)$.
- Annual cost for wrongly accepting H_1 : $\$(10)(50\,000)$.

The expected cost for the planning horizon (1 year):

$$C(\alpha, \beta, n) = (10)(50\,000)\alpha + (190)(50\,000)\beta + (1000)n \quad (9)$$

where the decision variables, $\{\alpha, \beta, n\}$, are interrelated (similarly with equations (2) and (3)), by:

$$\bar{X}_{cr} = \mu_0 + \frac{\sigma}{\sqrt{n}}Z_{1-\alpha} = 1.5\mu_0 + \frac{\sigma}{\sqrt{n}}Z_{\beta} \quad (10)$$

and σ is the standard deviation of the lifetime of the product. Note that α and β , or equivalently $Z_{1-\alpha}$ and Z_{β} , uniquely determine n in equation 10.

To solve this optimization problem, the probabilities $\{\alpha, \beta\}$ may be *explicitly* represented in equation (9) by the corresponding Z values, using the highly accurate approximation for the cumulative distribution function **CDF** of the normal distribution (refer to [8]):

$$\Pr(Z \leq z) = \Phi(z) = \frac{1}{2}[1 + g(-z) - g(z)] \quad (11)$$

where

$$g(z) = \exp\{-\log(2) \exp\{\left(\frac{\alpha}{\lambda/S_1}\right) \times [(1 + S_1 z)^{\lambda/S_1} - 1] + S_2 z\}\} \quad (12)$$

and

$$\lambda = -0.61228883; S_1 = -0.11105481; S_2 = 0.44334159; \alpha = -6.37309208 \quad (13)$$

Once both equations (9) and (10) are expressed explicitly in terms of $\{Z_{1-\alpha}, Z_{\beta}, n\}$, the minimal value of equation (9), subject to constraint (10), may be identified. This solution will provide the optimal $\{Z_{1-\alpha}^*, Z_{\beta}^*, n^*\}$.

For our problem, assuming $\mu_0 = 1000$ h and $\sigma = 70$ h, we introduce for n from equation (10) into equation (9) to obtain from the latter the optimal values of:

$$n^* = 1.233 \approx 2; \alpha^* = (1.6)10^{-4}; \beta^* = (7.2)10^{-6}; x_{cr}^* = 1227 \quad (14)$$

and the total expected cost is \$1383.

References

- [1] Cook, R.D. & Weisberg, S. (1999). *Applied Regression Including Computing and Graphics*, John Wiley & Sons.
- [2] Draper, N.R. & Smith, H. (1998). *Applied Regression Analysis*, 3rd Edition, John Wiley & Sons.
- [3] Burnham, K.P. & Anderson, D.R. (2002). *Model Selection and Multi-Model Inference- A Practical Information-Theoretic Approach*, 2nd Edition, Springer-Verlag, New York.
- [4] Shore, H. (2005). *Response Modeling Methodology: Empirical Modeling for Engineering and Science*, World Scientific Publishing, Singapore.
- [5] Myers, R.H. & Montgomery, D.C. (2002). *Response Surface Methodology*, John Wiley & Sons, New York.
- [6] Montgomery, D.C. (2001). *Introduction to Statistical Quality Control*, 4th Edition, John Wiley & Sons.
- [7] Kenett, R.S. & Zacks, S. (1998). *Modern Industrial Statistics: Design and Control of Quality and Reliability*, Duxbury Press (An International Thomson Publishing Company), Belmont.
- [8] Shore, H. (2005). Accurate RMM-based approximations for the CDF of the normal distribution, *Communications in Statistics – Theory and Methods* **34**, 507–513.

HAIM SHORE

Sigma Metric (Sigma Level)

Why is it Called Six Sigma?

The sigma level has been proposed as a measure of conformance quality [1]. There is a one-to-one correspondence between a process's sigma level and the fraction of nonconforming items it produces. Various definitions are current, but the basic idea is to express conformance quality as a Z value. That is, the characteristic under study is transformed such that it has a standard **normal distribution**; the position of the specification limit on this transformed scale is the Z value. For a normally distributed $N(\mu, \sigma^2)$ characteristic X with upper specification limit USL , this amounts to

$$Z = \frac{(USL - \mu)}{\sigma} \quad (1)$$

There are various conventions for dealing with the case that there is a lower and an upper specification limit. One option is to consider only the nearest specification limit:

$$Z = \min\left(\frac{(USL - \mu)}{\sigma}, \frac{(\mu - LSL)}{\sigma}\right) \quad (2)$$

Alternatively, one could compute the so-called **benchmark** Z value, which is

$$Z = \Phi^{-1}(d) \quad (3)$$

with Φ the standard normal distribution function, and d the total fraction nonconforming: $d = \Phi((LSL - \mu)/\sigma) + 1 - \Phi((USL - \mu)/\sigma)$. If X has a non-normal distribution the Z value is also computed from equation (3), be it that d should be computed differently.

If the sigma level is estimated from a data sample collected in a relatively short period, the estimate is often called the *short-term sigma level*. If one assesses the long-term sigma level from a short-term data sample, one could apply a *1.5 sigma shift*. This controversial rule of thumb says that the long-term sigma level is typically 1.5 smaller than the short-term level. Following this rule of thumb, a *six sigma process* (i.e., $Z_{\text{short term}} = 6.0$) has a long-term level $Z_{\text{long term}} = 4.5$, which corresponds to a fraction nonconforming of $d = 1 - \Phi(4.5) = 3.4$ ppm (parts per million). It is this conformance level that gave the Six Sigma method (see **Six Sigma Method**) its name.

Reference

- [1] Harry, M. & Schroeder, R. (2000). *Six Sigma*, Currency Doubleday, New York, p. 7.

RONALD J.M.M. DOES AND JEROEN DE MAST

Six Sigma Method

Six Sigma in a nutshell

From the late 1980s onwards, many companies started programmes aimed at quality and efficiency improvement under the name **Six Sigma**. Motorola was the first company to organize its improvement activities under this name, but it was, in particular, the programme's adoption in 1995 by General Electric which gave it enormous momentum. It borrows many principles and techniques from twentieth century quality engineering and industrial statistics, but its impact and completeness set it apart from earlier comparable initiatives. It could be seen as a successor of initiatives such as Taguchi's methodology, the Shainin method, and total **quality management** [1]. An overview of relevant literature is provided by Brady and Allen [2].

The approach is often characterized by its emphasis on decision making based on quantitative data, its focus on bottom-line results, and its customer driven approach. From the perspective of business economics, Six Sigma is a programme for organizing projectwise improvement of routine operations in companies. It offers a template organizational structure in which project leaders are named *black belts* or *green belts*, and project owners are named *champions*. To the project leader Six Sigma offers a stepwise strategy and an extensive collection of techniques and methods; De Koning and De Mast [3] provide a reconstruction of the method based on a multitude of sources.

A substantial part of the techniques and methods that Six Sigma prescribes for projects consist of methods borrowed from industrial statistics and quality engineering; Hoerl [4] gives an overview of a typical curriculum of a Six Sigma black belt course. Typical methods used in Six Sigma projects are statistically designed experiments, **gauge R&R** studies, and process capability analysis. In addition to these statistical and quality engineering tools the programme offers techniques borrowed from marketing and **project management**.

Six Sigma's stepwise strategy is often referred to as *DMAIC*, which stands for its main phases Define, Measure, Analyze, Improve, and Control. The objective of these phases can be described as follows [3]:

Define: problem selection and benefit analysis.

Measure: translation of the problem into a measurable form, and measurement of the current situation.

Analyze: identification of influence factors and causes that determine the behavior of critical to quality (CTQ) characteristics.

Improve: design and implementation of adjustments to the process to improve the performance of the CTQ characteristics.

Control: empirical verification of the project's results, and adjustment of the process management and control system in order that improvements are sustainable.

In many accounts these five phases are further subdivided in steps [3]. The programme prescribes that problems are quantified and parameterized in the form of indicators called *CTQ characteristics*. Improvement actions are based on the discovery of causal influence factors, sometimes called *leverage variables* or *the X's*. After 2000 the method acquired the status of standard approach in some nonmanufacturing industries as well, in particular, in healthcare and **finance** [5].

References

- [1] De Mast, J. (2004). A methodological comparison of three strategies for quality improvement, *International Journal of Quality and Reliability Management* **21**, 198–213.
- [2] Brady, J.E. & Allen, T.T. (2006). Six Sigma literature: a review and agenda for future research, *Quality and Reliability Engineering International* **22**, 335–367.
- [3] De Koning, H. & De Mast, J. (2006). A rational reconstruction of Six-Sigma's breakthrough cookbook, *International Journal of Quality and Reliability Management* **23**, 766–787.
- [4] Hoerl, R.W. (2001). Six Sigma black belts: What do they need to know? *Journal of Quality Technology* **33**, 391–406.
- [5] Van den Heuvel, J., Does, R.J.M.M., Bogers, A.J.J.C. & Berg, M. (2006). Implementing Six Sigma in the Netherlands, *Joint Commission Journal on Quality and Patient Safety* **32**, 393–399.

Related Article

Sigma Metric (Sigma Level)

RONALD J.M.M. DOES AND JEROEN DE MAST

Skewness

The skewness of a random variable X is related to the third *moment* (see **Moments**) of its distribution:

$$\gamma_1 = \frac{\mu_3}{\sigma^3} \quad (1)$$

where $\mu_3 = E(X - EX)^3$ and $\sigma^2 = E(X - EX)^2$. Skewness is a measure of the symmetry of density functions, with zero skewness corresponding to a symmetrical density.

Standard Error

Estimators are random variables, and therefore they have a probability distribution. The standard error (SE) is the standard deviation of the distribution of an estimator. The smaller the standard error, the more precise (or, as it is put more commonly, *efficient*)

is the estimator. For estimators (approximately) having a **normal distribution** (such as the sample average), an approximately 95% *confidence interval* (see **Confidence Intervals**) is obtained taking plus and minus two times the standard error. See also **Estimation**.

Sum of Squares

Sums of squares are used in *analysis of variance* and regression analysis to quantify the effect of sources of variation. Such sources may be explanatory variables and factors, but also variation due to **lack of fit** (e.g., sums of squares for curvature when one fits a linear regression model to data that cannot be described by a linear model). The total sum of squares SS_{total} is given by $\sum_{i=1}^n (y_i - \bar{y})^2$. So, the total sum of squares is really the sum of squared deviations from the mean. The reason is that the simplest model for data is the sample mean. All other models should be compared with that model and hence, the sum of squares reflects the variation in the data not explained by the mean. *Analysis of variance* depends on the decomposition $SS_{\text{total}} = SS_{\text{model}} + SS_{\text{error}}$. The terms for SS_{model} depend on the precise form of the model. For example, for a balanced one-way analysis of variance model with n observations and t treatments, $SS_{\text{model}} = \sum_{i=1}^t (y_{i.} - y_{..}/n)^2$.

For unbalanced designs (i.e., different cell sizes or unequal numbers of treatment combinations), there are different decompositions of the sums of squares which unlike in the balanced case, do not yield the same result. Type I sums of squares are based on

a sequential approach, where the order of entry of terms into the overall model is important. When the order of entry of terms is based on term hierarchy, one speaks of type II sums of squares. Type III sums of squares are based on comparing estimates of marginal population means. These sums of squares are the correct choice when one is interested in the importance of *main effects* and *interactions* in unbalanced designs. Closely related to type III sums of squares are type IV sums of squares. These are developed for the case of missing cells. In case there are no missing cells, type III and type IV sums of squares coincide. We refer to [1] for practical guidance in using the correct sum of squares.

Reference

- [1] Yandell, B.S. (1977). *Practical Data Analysis for Designed Experiments*, Chapman & Hall, London.

Related Articles

Analysis of Variance; Interactions; Main Effects.

ALESSANDRO DI BUCCHIANICO

Transfer Function

The term *transfer function* is borrowed from system theory, and refers to the *mathematical equation*

that is used to represent the relation between the inputs and outputs of a system. Transfer functions are used extensively in process optimization and control.

Trimmed Mean

The trimmed mean is a robust estimator for the mean. It is computed by removing a proportion of the largest and smallest observations, and averaging the remaining values. For example, the 10% trimmed

mean of the data 6, 2, 10, 14, 9, 8, 22, 15, 13, 82, and 11 is found removing the 10% smallest and largest observations (2 and 82 in this case). From the remainder we take the average, which is 12.0.

Variance

Variance as a Property of a Distribution

Variance is the most commonly used measure for dispersion. Given a random variable X , with **probability density function** $f(x)$ and mean μ , the variance is defined as

$$\sigma_x^2 = E(X - \mu)^2 = \int_{-\infty}^{\infty} (x - \mu)^2 f(x) dx \quad (1)$$

for a continuous variable, and by

$$\sigma_x^2 = E(X - \mu)^2 = \sum_{j=0}^{\infty} (x - \mu)^2 f(x), \quad x \geq 0 \quad (2)$$

for a nonnegative discrete variable (where $f(x)$ in equation (2) is the probability function). The variance is thus the second central moment (*see Moments*) of the distribution.

A short formula to calculate the variance, based on the first two noncentral **moments**, is

$$\sigma_x^2 = E(X^2) - [E(X)]^2 \quad (3)$$

The variance σ_T^2 of a linear transformation $T = a_0 + a_1 X$ is $\sigma_T^2 = a_1^2 \sigma_x^2$.

Let $\{X_1, X_2, \dots, X_k\}$ be k independent random variables, with standard deviations $\{\sigma_1, \sigma_2, \dots, \sigma_k\}$ respectively. Define the weighted sum:

$$S = a_0 + a_1 X_1 + \dots + a_k X_k \quad (4)$$

The variance of S , σ_S^2 , is

$$\sigma_S^2 = a_1^2 \sigma_1^2 + a_2^2 \sigma_2^2 + \dots + a_k^2 \sigma_k^2 \quad (5)$$

In case the $\{X_1, X_2, \dots, X_k\}$ are correlated, with ρ_{ij} denoting the **correlation** between X_i and X_j , the variance of S is

$$\begin{aligned} \text{Var} \left(a_0 + \sum_{i=1}^k a_i X_i \right) &= \sum_{i=1}^k a_i^2 \text{Var}(X_i) \\ &+ 2 \sum_{i=1}^{k-1} \sum_{j=i+1}^k (a_i a_j \sigma_i \sigma_j \rho_{ij}) \end{aligned} \quad (6)$$

A *random* sum is defined by

$$S = X_1 + X_2 + \dots + X_M \quad (7)$$

where $\{X_1, X_2, \dots, X_M\}$ are independently and identically distributed (i.i.d) random variables, each with mean μ_X and variance σ_X^2 , and M is a random variable with a mean, μ_M , and a variance, σ_M^2 (M is assumed to be independent of any X_i). The mean and variance of S are, respectively,

$$\mu_S = \mu_M \mu_X \quad (8)$$

$$\sigma_S^2 = \mu_M \sigma_X^2 + \sigma_M^2 \mu_X^2 \quad (9)$$

Let $\{X, Y\}$ be a pair of random variables with finite variances. Given the conditional variance of Y , $V(Y|X = x)$, and the conditional mean of Y , $E(Y|X = x)$, the *unconditional* variance of Y is:

$$V(Y) = E_x\{V(Y|X = x)\} + V_x\{E(Y|X = x)\} \quad (10)$$

where E_x is the expected value, calculated with respect to X , and V_x is the variance, calculated with respect to X . Equation (10) expresses the *law of total variance*.

For a catalog of some probability density functions, and the associated formulas for variance calculation, refer to **Probability Density Functions**.

Variance as a Population Parameter

Suppose that we have a population of N units, and on this population we have defined a variable, X . The population variance, σ_p^2 , is calculated by

$$\sigma_p^2 = \frac{1}{N} \sum_{i=1}^N (X_i - \mu_p)^2 \quad (11)$$

where X_i is the X value associated with unit i , and μ_p is the population mean.

The Sample Variance

The sample variance is a statistic commonly used as an unbiased, consistent, and minimum-variance estimator of the population variance. Suppose that a random sample of n observations have been collected, and denote the value of the i th observation by X_i . An

2 Variance

unbiased sample estimator of the population variance, based on a sample with n observations, is

$$\hat{\sigma}_x^2 = \frac{1}{(n-1)} \sum_{i=1}^n (X_i - \hat{\mu})^2 \quad (12)$$

where $\hat{\mu}$ is the sample average. Note, that although this sample estimator of the population variance is unbiased, its square root, often used to estimate the standard deviation, is biased. In **statistical process control** (SPC), the sample range is often employed in a formula that estimates the standard deviation.

Equation (12) is a variance estimator for an infinite population or a population that is large enough to be considered so. When the population is finite, random sampling from the population may be with replacement (each sampled unit is *thrown* back into the population for a possible repeat sampling) or without replacement. For the latter case, equation (12) is no longer appropriate to estimate the variance. It may be shown that the expected value of the statistic in equation (12) is no longer the population

variance but rather (refer to [1, p. 171]):

$$E(\hat{\sigma}_x^2) = \sigma_x^2 \left(1 + \frac{1}{N-1}\right) \quad (13)$$

where N is the *population* size. Similarly, the variance of the sample mean for a random sample of n items, drawn without replacement, is [1, p. 170]

$$V(\bar{X}) = \frac{\sigma_x^2}{n} \left(1 - \frac{n-1}{N-1}\right) \quad (14)$$

Equations (13) and (14) are often useful in sampling without replacement conducted within **acceptance sampling**.

Reference

- [1] Kenett, R.S. & Zacks, S. (1998). *Modern Industrial Statistics: Design and Control of Quality and Reliability*, Duxbury Press (An International Thomson Publishing Company), Belmont.

HAIM SHORE

Measures of Association

Introduction

Measures of association quantify the statistical dependence of two or more categorical variables. Many measures of association have been proposed; in this article, we restrict our attention to the ones that are most commonly used. We broadly divide the measures into those suitable for (a) unordered (nominal) categorical variables, and (b) ordered categorical variables.

Measures of Association for Unordered Categorical Variables

As an example, consider the contingency table shown in Table 1, with sex (s) as the row variable and employment status (e) as the column variable.

Measures Based on the Odds Ratio

The odds that an event will occur are computed by dividing the probability that the event will occur by the probability that the event will not occur. For males, the odds of being employed equal the probability of being employed divided by the probability of not being employed:

$$\begin{aligned} odds_{y|m} &= \frac{p(y|m)}{1 - p(y|m)} = \frac{p(y|m)}{p(n|m)} \\ &= \frac{42/75}{33/75} = \frac{0.56}{0.44} = 1.27 \end{aligned} \quad (1)$$

Similarly, for females:

$$\begin{aligned} odds_{y|f} &= \frac{p(y|f)}{1 - p(y|f)} = \frac{p(y|f)}{p(n|f)} \\ &= \frac{40/165}{125/165} = \frac{0.242}{0.758} = 0.319 \end{aligned} \quad (2)$$

These odds are large when employment is likely, and small when it is unlikely. The odds can be interpreted as the number of times that employment occurs for each time it does not. For example, since the odds of employment for males is 1.27, approximately 1.27 males are employed for every unemployed male, or,

Table 1 Employment status of males and females

		Employment status		Total
		Employed (y)	Not employed (n)	
Sex	Male (m)	42	33	75
	Female (f)	40	125	165
	Total	82	58	40

in round numbers, 5 males are employed for every four males that are unemployed. For females, there is approximately one employed female for every three unemployed females.

The association between sex and employment status can be expressed by the degree to which the two odds differ. This difference can be summarized by the odds ratio (α):

$$\alpha = \frac{odds_{y|f}}{odds_{y|m}} = \frac{p(y|f)/p(n|f)}{p(y|m)/p(n|m)} = \frac{1.27}{0.32} = 3.98 \quad (3)$$

This odds ratio indicates that for every one employed female, there are four employed males. Although not necessary, it is customary to place the larger of the two odds in the numerator.

Odds ratios are not symmetric around one: An odds ratio larger than one by a given amount indicates a smaller effect than an odds ratio smaller than one by the same amount [1]. While the magnitude of an odds ratio less than one is restricted to the range between zero and one, odds ratios greater than one and are not restricted, allowing the ratio to potentially take on any value. If the natural logarithm ($\ln$) of the odds ratio is taken, the odds ratio is symmetric above and below one, with $\ln(1) = 0$. For example, to take the 2×2 example above, the odds ratio of a male being employed, compared to a female, was 3.98. If we reverse that and take the odds ratio of a female being employed compared to a male, it is $0.32/1.27 = 0.25$. These ratios are clearly not symmetric about 1.00. However, $\ln(3.98) = 1.38$ and $\ln(0.25) = -1.38$, and these ratios are symmetric.

Yule's Q Coefficient of Association. Both α and $\ln(\alpha)$ range from $-\infty$ to $+\infty$. to restrict the odds ratio within the interval -1 to $+1$, Yule introduced the coefficient of association, Q , for 2×2 tables. Its

2 Measures of Association

definition is [2]:

$$Q = \frac{n_{11}n_{22} - n_{12}n_{21}}{n_{11}n_{22} + n_{12}n_{21}} = \frac{\alpha - 1}{\alpha + 1} \quad (4)$$

Therefore, if odds ratio (α) expressing the relationship between sex and employment is 3.98,

$$\begin{aligned} Q &= \frac{(42)(125) - (33)(40)}{(42)(125) + (33)(40)} \\ &= \frac{3.98 - 1}{3.98 + 1} = 0.59 \end{aligned} \quad (5)$$

Yule's Q is algebraically equal to the 2×2 Goodman–Kruskal γ coefficient (described below), and thus measures the degree of concordance or discordance between two variables.

Yule's Coefficient of Colligation. As an alternative to Q , Yule proposed the *coefficient of colligation* [2]:

$$\hat{Y} = \frac{\sqrt{n_{11}n_{22}} - \sqrt{n_{12}n_{21}}}{\sqrt{n_{11}n_{22}} + \sqrt{n_{12}n_{21}}} = \frac{\sqrt{a} - 1}{\sqrt{a} + 1} \quad (6)$$

The coefficient of colligation between sex and employment is

$$\begin{aligned} \hat{Y} &= \frac{\sqrt{(42)(125)} - \sqrt{(33)(40)}}{\sqrt{(42)(125)} + \sqrt{(33)(40)}} \\ &= \frac{\sqrt{3.98} - 1}{\sqrt{3.98} + 1} = 0.33 \end{aligned} \quad (7)$$

The coefficient of colligation has a different interpretation than the coefficient of association, and is interpreted as a Pearson product–moment **correlation** coefficient r (see below), but is not algebraically equivalent to one.

Both Q and $\hat{Y}$ are symmetric measures of association, and are invariant under changes in the ordering of rows and columns.

Measures Based on Pearson's χ^2

$$\chi^2 = \sum_i \sum_j \frac{(n_{ij} - \hat{n}_{ij})^2}{n_{ij}} \quad (8)$$

where $\hat{n}_{ij}$ is the expected frequency in cell ij , and can be usefully transformed into several measures of association [3].

The Φ Coefficient. The ϕ coefficient is defined as [4]:

$$\Phi = \sqrt{\frac{\chi^2}{N}} \quad (9)$$

For sex and employment,

$$\Phi = \sqrt{\frac{23.12}{240}} = 0.31 \quad (10)$$

The Φ coefficient can vary between 0 and 1, and is algebraically equivalent to $|r|$. However, the lower and upper limits of Φ in a 2×2 table are dependent on two conditions. In order for Φ to equal -1 or $+1$, (a) $(n_{11} + n_{12}) = (n_{21} + n_{22})$, and (b) $(n_{11} + n_{21}) = (n_{12} + n_{22})$.

The Pearson Product-Moment Correlation Coefficient. The general formula for the *Pearson product–moment correlation coefficient* is

$$r = \frac{\sum_i (X_i - \bar{X})(Y_i - \bar{Y})}{N s_x s_y} \quad (11)$$

where X and Y are two continuous, interval-level variables. The categories of a dichotomous variable can be coded 0 and 1, and used in the formula for r . In a 2×2 contingency table, the calculations reduce to [5]:

$$r = \frac{n_{11}n_{22} - n_{12}n_{21}}{\sqrt{n_{1+}n_{2+}n_{+1}n_{+2}}} \quad (12)$$

For sex and employment, $r = \{(42)(125) - (33)(40) / \sqrt{[(75)(165)(82)(158)]}\} = 0.31$.

A symmetric measure of association is also invariant to row and column order, r varies between -1 to $+1$. From the formula, it is apparent that $r = 1$ if $n_{12} = n_{21} = 0$, and $r = -1$ if $n_{11} = n_{22} = 0$. In a standardized 2×2 table, where each marginal probability = 0.5, $r = \hat{Y}$; otherwise $|r| < |\hat{Y}|$ except

when the variables are independent or completely related.

Measures of Predictive Association

Another class of measures of association for nominal variables is measures of prediction analogous in concept to the multiple correlation coefficient or R^2 in regression analysis [2](see **Coefficient of Determination (R^2)**). When there is an association between two nominal variables X and Y , then knowledge about X allows one to obtain knowledge about Y , more knowledge than would have been available without X . Let Δ_Y be the dispersion of Y and $\Delta_{Y.X}$ be the conditional dispersion of Y given X . A measure of prediction

$$\phi_{Y.X} = 1 - \frac{\Delta_{Y.X}}{\Delta_Y} \quad (13)$$

compares the conditional dispersion of Y given X to the unconditional dispersion of Y , similar to how the multiple correlation coefficient compares the conditional variance of the dependent variable to its unconditional variance [3]. When $\phi_{Y.X} = 0$, X and Y are independently distributed; when $\phi_{Y.X} = 1$, X is a perfect predictor of Y .

We describe four measures that operationalize the idea underlying $\phi_{Y.X} = 0$.

The first measure is the *asymmetric lambda coefficient* of Goodman and Kruskal [6]:

$$\begin{aligned} \lambda(Y = c|X = r) &= \frac{\sum_i \max_j p_{ij} - \max_j p_{+j}}{1 - \max_j p_{+j}} \\ \lambda(X = r|Y = c) &= \frac{\sum_j \max_i p_{ij} - \max_i p_{+i}}{1 - \max_i p_{+i}} \end{aligned} \quad (14)$$

The second is the *symmetric lambda coefficient* (λ) of Goodman and Kruskal [6]:

$$\lambda = \frac{\sum_i \max_j p_{ij} + \sum_j \max_i p_{ij} - \max_j p_{+j} - \max_i p_{+i}}{2 - \max_j p_{+j} - \max_i p_{+i}} \quad (15)$$

The third is the *asymmetric uncertainty coefficient* of Theil [7]:

$$\begin{aligned} U(Y = c|X = r) &= \frac{\left(-\sum_i (p_{i+}) \ln(p_{i+}) \right) + \left(-\sum_j (p_{+j}) \ln(p_{+j}) \right) - \left(-\sum_i \sum_j (p_{ij}) \ln(p_{ij}) \right)}{-\sum_j (p_{+j}) \ln(p_{+j})} \\ U(X = r|Y = c) &= \frac{\left(-\sum_i (p_{i+}) \ln(p_{i+}) \right) + \left(-\sum_j (p_{+j}) \ln(p_{+j}) \right) - \left(-\sum_i \sum_j (p_{ij}) \ln(p_{ij}) \right)}{-\sum_i (p_{i+}) \ln(p_{i+})} \end{aligned} \quad (16)$$

The fourth measure of predictive association is U , Theil's *symmetric uncertainty coefficient* [7]:

$$U = \frac{2 \left[\left(-\sum_i (p_{i+}) \ln(p_{i+}) \right) + \left(-\sum_j (p_{+j}) \ln(p_{+j}) \right) - \left(-\sum_i \sum_j (p_{ij}) \ln(p_{ij}) \right) \right]}{\left(-\sum_i (p_{i+}) \ln(p_{i+}) \right) + \left(-\sum_j (p_{+j}) \ln(p_{+j}) \right)} \quad (17)$$

We use the contingency table shown in Table 2 to illustrate, in turn, the computation of the four measures of predictive association.

The first measure, asymmetric $\lambda(Y|X)$, is interpreted as the relative decrease in the probability of incorrectly predicting the column variable Y between not knowing X and knowing X . As such, lambda can also be interpreted as a proportional reduction in

4 Measures of Association

Table 2 The computation of the four measures of predictive association

	Y_1	Y_2	Y_3	Y_4	Total
X_1	28	32	49	36	145
X_2	44	21	17	12	94
X_3	46	31	53	63	193
Total	118	84	119	111	

variation measure – how much of the variance in one variable is accounted for by the other? In the equation for $\lambda(Y|X)$, $1 - \max p_{ij}$ is the minimum probability of error from the prediction that Y is a function of X , and $1 - \max_j p_{+j}$ is the minimum probability of error from a prediction that Y is a constant over X . For the data we obtain:

$$\begin{aligned}\lambda(Y|X) &= \frac{((49 + 44 + 63)/432) - (119/432)}{1 - (119/432)} \\ &= \frac{0.36 - 0.28}{1 - 0.28} = 0.11\end{aligned}\quad (18)$$

Thus, there is an 11% improvement in predicting the column variable Y given the knowledge of the row variable X . Asymmetric λ has the range $0 \leq \lambda(Y|X) \leq 1$.

In general, $\lambda(Y|X) \neq \lambda(X|Y)$; $\lambda(X|Y)$ is the relative decrease in the probability of incorrectly predicting the row variable X between not knowing Y and knowing Y – hence, the term “asymmetric”. For these data,

$$\begin{aligned}\lambda(X|Y) &= \frac{((46 + 32 + 53 + 63)/432) - (193/432)}{1 - (193/432)} \\ &= \frac{0.45 - 0.44}{1 - 0.44} = 0.02\end{aligned}\quad (19)$$

Symmetric λ is the average of the two asymmetric lambdas and has the range $0 \leq \lambda \leq 1$. For our example,

$$\lambda = \frac{0.36 + 0.45 - 0.28 - 0.44}{2 - 0.28 - 0.44} = \frac{0.09}{1.28} = 0.07 \quad (20)$$

Theil’s asymmetric uncertainty coefficient, $U(Y|X)$, is the proportion of uncertainty in the column variable Y that is explained by the row variable X , or, alternatively, $U(X|Y)$ is the proportion of uncertainty in the row variable X that is explained

by the column variable Y . The asymmetric uncertainty coefficient has the range $0 \leq U(Y|X) \leq 1$. For $U(Y|X)$, we obtain for these data:

$$U(Y|X) = \frac{(1.06) + (1.38) - (2.40)}{1.38} = 0.03 \quad (21)$$

and for $U(X|Y)$,

$$U(X|Y) = \frac{(1.06) + (1.38) - (2.40)}{1.06} = 0.04 \quad (22)$$

The symmetric U is computed as

$$U = \frac{2[(1.06) + (1.38) - (2.40)]}{(1.06) + (1.38)} = 0.03 \quad (23)$$

Both the asymmetric and symmetric uncertainty coefficients have the range $0 \leq U \leq 1$.

The quantities in the numerators of the equations for the asymmetric lambda and uncertainty coefficients are interpreted as measures of variation for nominal responses: In the case of lambda coefficients, the variation measure is called the *Gini concentration*, and the variation measure used in the uncertainty coefficients is called the *entropy* [2]. More specifically, in $\lambda(Y|X)$, the numerator represents the variance of the Y or column variable:

$$V(Y) = \sum_i \max_j p_{ij} - \max_j p_{+j} \quad (24)$$

and in $\lambda(X|Y)$, the numerator represents the variance of the X or row variable:

$$V(X) = \sum_j \max_i p_{ij} - \max_i p_{i+} \quad (25)$$

Analogously, the numerator of $U(Y|X)$ is the variance of Y ,

$$\begin{aligned}V(Y) &= (-\sum_i (p_{i+}) \ln(p_{i+})) \\ &\quad + (-\sum_j (p_{+j}) \ln(p_{+j})) \\ &\quad - (-\sum_i \sum_j (p_{ij}) \ln(p_{ij}))\end{aligned}\quad (26)$$

and the numerator of $U(X|Y)$ is the variance of X ,

$$\begin{aligned}V(X) &= (-\sum_i (p_{i+}) \ln(p_{i+})) \\ &\quad + (-\sum_j (p_{+j}) \ln(p_{+j})) \\ &\quad - (-\sum_i \sum_j (p_{ij}) \ln(p_{ij}))\end{aligned}\quad (27)$$

However, in these measures, there is a positive relationship between the number of categories and the magnitude of the variation, thus introducing ambiguity in evaluating both the variance of the variables and their relationship.

Measures of Association for Ordered Categorical Variables

Measures of Concordance/Discordance

A pair of observations is *concordant* if the observation that ranks higher on X also ranks higher on Y . A pair of observations is *discordant* if the observation that ranks higher on X ranks lower on Y . The number of concordant pairs is

$$C = \sum_{i < i'} \sum_{j < j'} n_{ij} n_{i'j'} \quad (28)$$

where the first summation is over all rows $i < i'$, and the second summation is over all pairs of columns $j < j'$. The number of discordant pairs is

$$D = \sum_{i < i'} \sum_{j > j'} n_{ij} n_{i'j'} \quad (29)$$

where the first summation is over all rows $i < i'$, and the second summation is over all columns $j > j'$.

To illustrate the calculation of C and D , consider the 3×3 contingency table shown in Table 3 cross-classifying income with education.

In this table, the number of concordant pairs is

$$\begin{aligned} C &= 302(331 + 250 + 155 + 185) \\ &+ 105(250 + 185) + 409(155 + 185) \\ &+ 331(185) = 524\,112 \end{aligned} \quad (30)$$

Note that the process of identifying concordant pairs amounts to taking each cell frequency (e.g., 302) in turn, and multiplying that frequency by each of the cell frequencies “southeast” of it (e.g., 331, 250, 155, 185).

The number of discordant pairs is

$$\begin{aligned} D &= 23(409 + 331 + 15 + 155) \\ &+ 105(409 + 15) + 250(15 + 155) \\ &+ 331(15) = 112\,915 \end{aligned} \quad (31)$$

Discordant pairs can readily be identified by taking each cell frequency (e.g., 23) in turn, and multiplying that cell frequencies by each of the cell frequencies “southwest” of it (e.g., 409, 331, 15, 155).

Let $n_{i+} = \sum_j n_{ij}$ and $n_{+j} = \sum_i n_{ij}$. We can express the total number of pairs of observations as

$$\frac{n(n-1)}{2} = C + D + T_X + T_Y - T_{XY} \quad (32)$$

where $T_X = \sum_i n_i(n_{i+} - 1)/2$ is the number of pairs tied on the row variable X , $T_Y = \sum_j n_j(n_{+j} - 1)/2$ is the number of pairs tied on the column variable Y , and $T_{XY} = \sum_i \sum_j n_{ij}(n_{ij} - 1)/2$ is the number of pairs from a common cell (tied on X and Y). In this formula for $n(n-1)/2$, T_{XY} is subtracted because pairs tied on both X and Y have been counted twice, once in T_X and once in T_Y . Therefore,

$$\begin{aligned} T_X &= \frac{(430 \times 429) + (990 \times 989) \\ &\quad + (355 \times 354)}{2} \\ &= 644\,625 \end{aligned} \quad (33)$$

Table 3 Income for three levels of education

		Income				
		Low	Medium	High	Total	Proportion
Education	LT than high school	302	105	23	430	0.243
	High school	409	331	250	990	0.557
	GT than high school	15	155	185	355	0.200
	Total	726	591	458	1775	1.00
	Proportion	0.409	0.333	0.258	1.00	

$$T_Y = \frac{(726 \times 725) + (591 \times 590) + (458 \times 457)}{2} = 542\,173 \quad (34)$$

$$T_{XY} = \frac{(302 \times 301) + (105 \times 104) + \dots + (185 \times 184)}{2} = 249\,400 \quad (35)$$

$$\begin{aligned} \frac{n(n-1)}{2} &= (524\,112) + (112\,915) + (644\,625) \\ &\quad + (542\,173) - (249\,400) \\ &= 1\,574\,425 \end{aligned} \quad (36)$$

Kendall's τ Statistics. Several measures of concordance/discordance are based on the difference $C - D$. Kendall's three *tau* (τ) statistics are among the most well known of these, and can be applied to $r \times c$ contingency tables with $r \geq 2$ and $c \geq 2$. The row and column variables do not have to have the same number of categories. The (τ) statistics have the range $-1 \leq 0 \leq +1$.

The three (τ) statistics are [8]:

$$\tau_a = \frac{2(C - D)}{N(N - 1)} \quad (37)$$

$$\tau_b = \frac{(C - D)}{\sqrt{[(C + D + T_Y - T_{XY})(C + D + T_X - T_{XY})]}} \quad (38)$$

$$\tau_c = \frac{2m(C - D)}{N^2(m - 1)} \quad (39)$$

where m = the number of rows or column, whichever is smaller. For our data, the tau statistics are equal to

$$\tau_a = \frac{2(524\,112 - 112\,915)}{1775(1775 - 1)} = 0.26 \quad (40)$$

$$\begin{aligned} \tau_b &= \frac{(524\,112 - 112\,915)}{\sqrt{[(524\,112 + 112\,915 + 542\,173 - 249\,900) \\ &\quad \sqrt{(524\,112 + 112\,915 + 644\,625 - 249\,900)]}} \\ &= 0.42 \end{aligned} \quad (41)$$

$$\tau_c = \frac{3 \times 2(524\,112 - 112\,915)}{1775^2(3 - 1)} = 0.39 \quad (42)$$

One note of qualification is that $\tau - a$ assumes there are no tied observations; thus, strictly speaking, it is not applicable to contingency tables. It is generally more useful as a measure of rank correlation. For 2×2 tables, τ_b simplifies to the Pearson product-moment correlation obtained by assigning any scores to the rows and columns consistent with their orderings.

Goodman and Kruskal's γ . Goodman and Kruskal suggested yet another measure, *gamma* (γ) based on $C - D$ [6]:

$$\gamma = \frac{C - D}{C + D} \quad (43)$$

For our example data, $\gamma = (524\,112 - 112\,915)/(524\,112 + 112\,915) = 0.65$. γ has the range $-1 \leq 0 \leq +1$, equal to $+1$ when the data are concentrated in the upper-left to lower-right diagonal, and equal to -1 for the converse pattern. Although γ does equal zero when the two variables are independent, two variables can be completely dependent and still have a value of γ less than unity. For 2×2 tables, γ is equal to Yule's Q .

Measures Based on Derived Scores

Some methods for measuring the association between ordinal variables require assigning scores to the levels of the ordinal variables. When a contingency table is involved, scale values are assigned to row and column categories, the data are treated as a grouped frequency distribution, and the Pearson product-moment correlation is computed.

Specifically, let x_i and y_j be the values assigned to the rows and columns. The Pearson product-moment correlation for grouped data is then

$$r = \frac{C P_{XY}}{\sqrt{SS_X SS_Y}} \quad (44)$$

where $C P_{XY}$ is the sum of cross products,

$$C P_{XY} = \sum_{i,j} x_i y_j n_{ij} - \frac{\left(\sum_i x_i n_{i+} \right) \left(\sum_j y_j n_{+j} \right)}{n_{++}} \quad (45)$$

and SS_X and SS_Y are the sums of squares,

$$SS_X = \sum_i x_i^2 n_{i+} - \frac{\left(\sum_i x_i n_{i+}\right)^2}{n_{++}} \quad (46)$$

$$SS_Y = \sum_j y_j^2 n_{+j} - \frac{\left(\sum_j y_j n_{+j}\right)^2}{n_{++}} \quad (47)$$

Therefore, if for our 3×3 contingency table for income and education, we assign the values, “1”, “2”, and “3” to the three levels of each variable, we have

$$CP_{XY} = 6863 - \frac{(3475)(3282)}{1775} = 437.68 \quad (48)$$

$$SS_X = 7585 - \frac{(3475)^2}{1775} = 781.83 \quad (49)$$

$$SS_Y = 7212 - \frac{(3282)^2}{1775} = 1143.54 \quad (50)$$

$$r = \frac{437.68}{\sqrt{(781.83)(1143.54)}} = 0.46 \quad (51)$$

The value of $r = 0.46$ is also known as the *Spearman rank correlation coefficient* r_s , sometimes referred to as Spearman’s rho [9]. The Spearman rank correlation coefficient can also be computed by using:

$$r_s = \frac{6 \sum d^2}{n(n^2 - 1)} \quad (52)$$

where d is the difference in the rank scores between X and Y . Computed by using this formula,

$$r_s = \frac{6(20706)^2}{1775(1775^2 - 1)} = 0.46 \quad (53)$$

Instead of using the rank scores of “1”, “2”, and “3” directly, we can use *ridit scores* [10]. Ridit scores are cumulative probabilities; each ridit score represents the proportion of observations below category j plus half the proportion within j . We illustrate the calculation of ridit scores using the income and education contingency table. Beneath

each marginal total is the proportion of observations in the row or column. The ridit score for the low income category is $0.5(0.409) = 0.205$; for medium income, $0.409 + 0.5(0.333) = 0.576$; and for high income, the ridit score is $0.409 + 0.333 + 0.5(0.258) = 0.871$. The ridit scores for less than high school education is $0.5(0.243) = 0.122$; for high school education, $0.243 + 0.5(0.557) = 0.522$; and for greater than high school education, the ridit score is $0.243 + 0.557 + 0.5(0.200) = 0.900$. Applying the Pearson product-moment correlation formula to the ridit scores, we obtain $r_s = 0.47$.

References

- [1] Rudas, T. (1998). *Odds Ratios in the Analysis of Contingency Tables*, Sage Publications, Thousand Oaks.
- [2] Wickens, T.D. (1986). *Multiway Contingency Tables Analysis for the Social Sciences*, Lawrence Erlbaum, Hillsdale.
- [3] Agresti, A. (1984). *Analysis of Ordinal Categorical Data*, John Wiley & Sons, New York.
- [4] Fleiss, J.L. (1981). *Statistical Methods for Rates and Proportions*, 2nd Edition, John Wiley & Sons, New York.
- [5] Reynolds, H.T. (1984). *Analysis of Nominal Data*, Sage Publications, Newbury Park.
- [6] Goodman, L.A. & Kruskal, W.H. (1972). Measures of association for cross-classification, IV, *Journal of the American Statistical Association* **67**, 415–421.
- [7] Everitt, B.S. (1992). *The Analysis of Contingency Tables*, 2nd Edition, Chapman & Hall, Boca Raton.
- [8] Gibbons, J.D. (1993). *Nonparametric Measures of Association*, Sage Publications, Newbury Park.
- [9] Hildebrand, D.K., Laing, J.D. & Rosenthal, H. (1977). *Analysis of Ordinal Data*, Sage Publications, Newbury Park.
- [10] Bross, I.D.J. (1958). How to use ridit analysis, *Biometrics* **14**, 18–38.

Further Reading

Liebetrau, A.M. (1983). *Measures of Association*, Sage Publications, Newbury Park.

SCOTT L. HERSHBERGER AND DENNIS
G. FISHER

Article originally published in *Encyclopedia of Statistics in Behavioral Science* (2005, John Wiley & Sons, Ltd). Minor revisions for this publication by Jeroen de Mast.

Probability Density Functions

Introduction

Probability density functions (PDFs) (*see* **Probability Density Function (PDF)**) are mathematical formulas that provide information about how the values of random variables are distributed. For a discrete random variable (one taking only particular values within some interval), the formula gives the probability of each value of the variable. For a continuous random variable (one that can take any value within some interval), the formula specifies the probability that the variable falls within a particular range. This is given by the area under the curve defined by the **PDF**. Here a list of the most commonly encountered PDFs and their most important properties is provided. More comprehensive accounts of density functions can be found in [1, 2].

Probability Density Functions for Discrete Random Variables

Bernoulli

A simple PDF for a random variable, X , that can take only two values, for example, 0 or 1. Tossing a single coin provides a simple example. Explicitly the density is defined as

$$P(X = x_1) = 1 - P(X = x_0) = p \quad (1)$$

where x_1 and x_0 are the two values that the random variable can take (often labeled as “success” and “failure”) and P denotes probability. The single parameter of the Bernoulli density function is the probability of a “success”, p . The expected value of X is p and its variance is $p(1 - p)$.

Binomial

A PDF for a random variable, X , that is the number of “successes” in a series of n independent Bernoulli variables. The number of heads in n tosses of a coin provides a simple practical example. The binomial density function is given by

$$P(X = x) = \frac{n!}{x!(n-x)!} p^x (1-p)^{n-x},$$

$$x = 0, 1, 2, \dots, n \quad (2)$$

The probability of the random variable taking a value in some range of values is found by simply summing the PDF over the required range. The expected value of X is np and its variance is $np(1 - p)$. Some examples of binomial density functions are shown in Figure 1. The binomial is important in assigning probabilities to simple chance events such as the probability of getting three or more sixes in 10 throws of a fair die, in the development of simple statistical significance tests such as the sign test (*see* **Nonparametric Tests**) and as the error distribution used in logistic regression. See also the multinomial distribution defined later in this article.

Geometric

The PDF of a random variable, X , that is the number of “failures” in a series of independent Bernoulli variables before the first “success”. An example is provided by the number of tails before the first head, when a coin is tossed a number of times. The density function is given by

$$P(X = x) = p(1 - p)^{x-1}, \quad x = 1, 2, \dots \quad (3)$$

The geometric density function possesses a *lack of memory property*, by which we mean that in a series of Bernoulli variables the probability of the next n trials being “failures” followed immediately by a “success” remains the same geometric PDF regardless of what the previous trials were. The mean of X is $1/p$ and its variance is $(1 - p)/p^2$. Some examples of geometric density functions are shown in Figure 2.

Negative Binomial

The PDF of a random variable, X , that is the number of “failures” in a series of independent Bernoulli variables before the k th “success”. For example, the number of tails before the 10th head in a series of coin tosses. The density function is given by

$$P(X = x) = \frac{(k+x-1)!}{(x-1)!k!} p^k (1-p)^x, \quad x = 0, 1, 2, \dots \quad (4)$$

2 Probability Density Functions

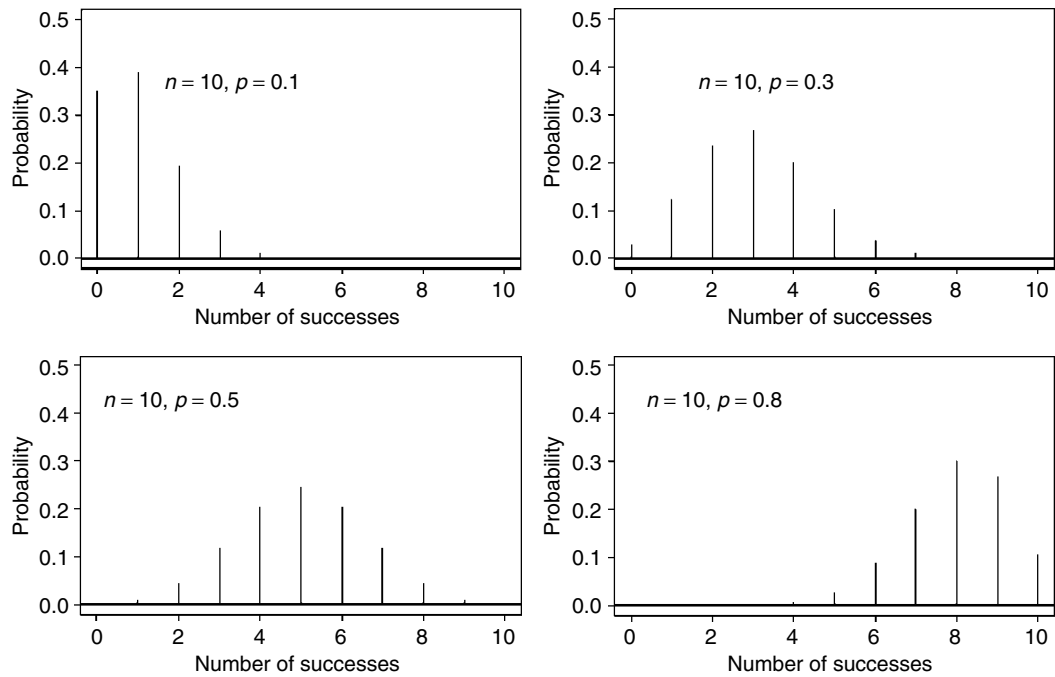


Figure 1 Examples of binomial density functions

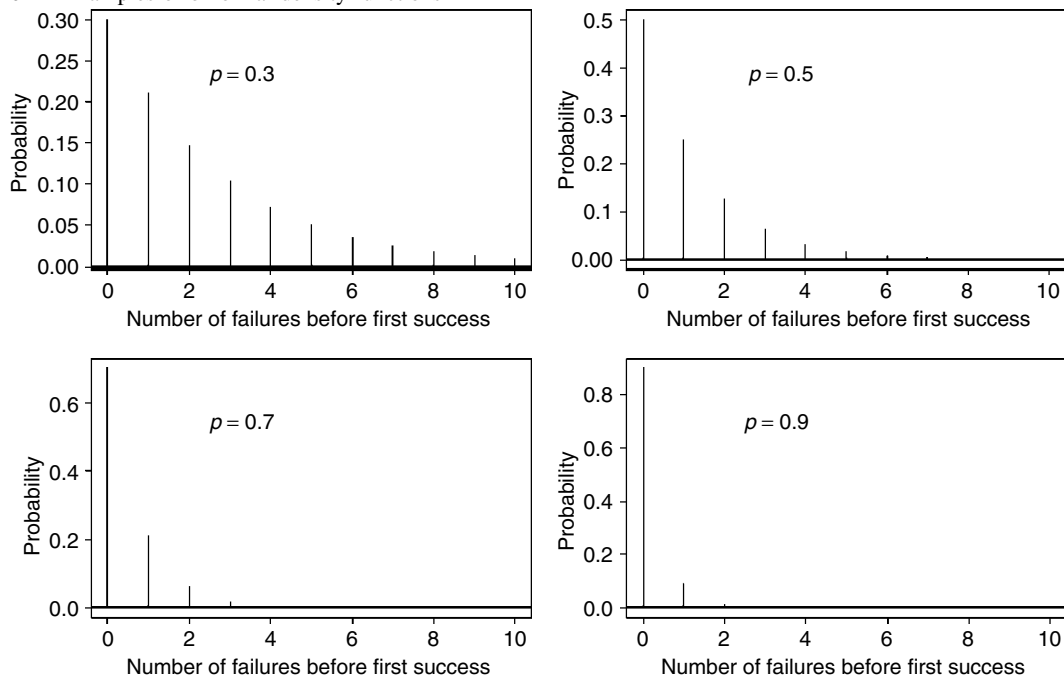


Figure 2 Examples of geometric density functions

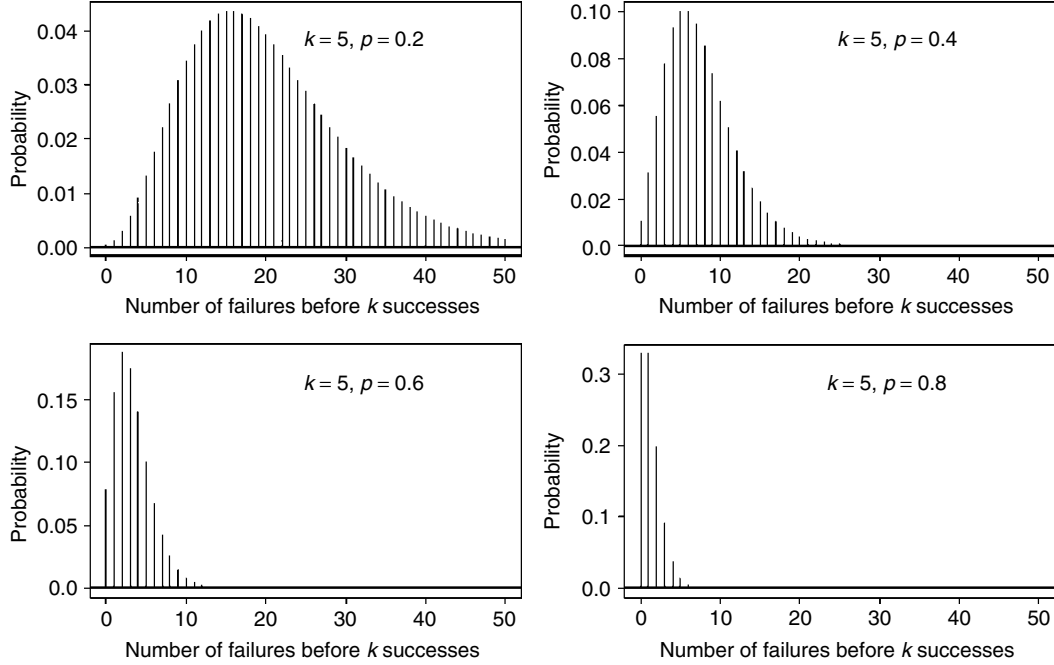


Figure 3 Examples of negative binomial density functions

The mean of X is $k(1-p)/p$ and its variance is $k(1-p)/p^2$. Some examples of the negative binomial distribution are given in Figure 3. The density is important in discussions of overdispersion in the context of **Generalized Linear Models**.

Hypergeometric

A PDF associated with *sampling without replacement* from a population of finite size. If the population consists of r elements of one kind and $N-r$ of another, then the hypergeometric is the PDF of the random variable X defined as the number of elements of the first kind when a random sample of n is drawn. The density function is given by

$$P(X = x) = \frac{r!(N-r)!}{x!(r-x)!(n-x)!(N-r-n+x)!} \frac{N!}{n!(N-n)!} \quad (5)$$

The mean of X is nr/N and its variance is $(nr/N)(1-r/n)[(N-n)/(N-1)]$. The hypergeometric density function is the basis of Fisher's

exact test used in the analysis of sparse contingency tables.

Poisson

A PDF that arises naturally in many instances, particularly as a probability model for the occurrence of rare events, for example, the emission of radioactive particles. In addition, the Poisson density function is the limiting distribution of the binomial when p is small and n is large. The Poisson density function for a random variable X taking integer values from 0 to ∞ is defined as

$$P(X = x) = \frac{e^{-\lambda} \lambda^x}{x!}, \quad x = 0, 1, 2, \dots, \infty \quad (6)$$

The single parameter of the Poisson density function, λ , is both the expected value and variance, that is, the mean and variance of a Poisson random variable are equal. Some examples of Poisson density functions are given in Figure 4.

4 Probability Density Functions

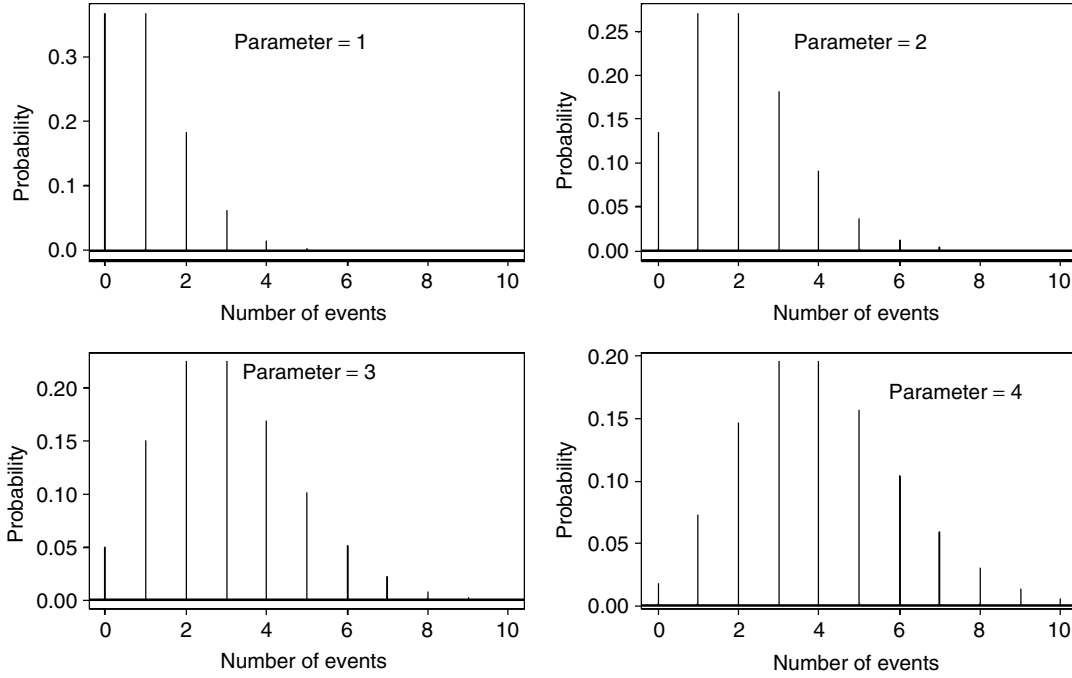


Figure 4 Examples of Poisson density functions

Probability Density Functions for Continuous Random Variables

Normal

The PDF, $f(x)$, of a continuous random variable X defined as follows

$$f(x) = \frac{1}{\sigma\sqrt{2\pi}} \exp\left[-\frac{1}{2} \frac{(x - \mu)^2}{\sigma^2}\right] \quad (7)$$

$-\infty < x < \infty$

where μ and σ^2 are, respectively the mean and variance of X . When the mean is zero and the variance one, the resulting density is labeled as the *standard normal*. The normal density function is bell-shaped as is seen in Figure 5, where a number of normal densities are shown.

In the case of continuous random variables, the probability that the random variable takes a particular value is strictly zero; nonzero probabilities can only be assigned to some range of values of the variable. So, for example, we say that $f(x)dx$ gives the probability of X falling in the very small interval,

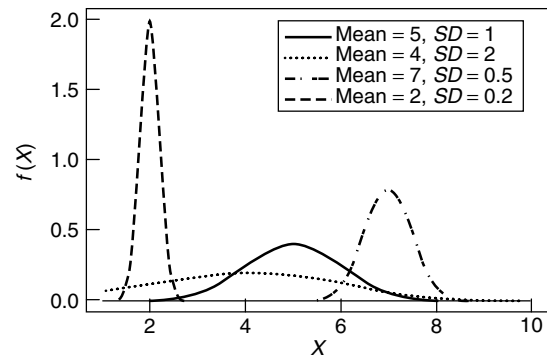


Figure 5 Normal density functions

dx , centered on x , and the probability that X falls in some interval, $[A, B]$ say, is given by integrating $f(x)dx$ from A to B .

The normal density function is ubiquitous in statistics. The vast majority of statistical methods are based on the assumption of a normal density for the observed data or for the error terms in models for the data. In part this can be justified by an appeal to the central limit theorem. The density function first

appeared in the papers of de Moivre at the beginning of the eighteenth century and some decades later was given by Gauss and Laplace in the theory of errors and the least squares method. For this reason, the normal is also often referred to as the *Gaussian* or *Gaussian–Laplace*.

Uniform

The mean of such a random variable is having constant probability over an interval. The density function is given by

$$f(x) = \frac{1}{\beta - \alpha}, \quad \alpha < x < \beta \quad (8)$$

The mean of the density function is $(\alpha + \beta)/2$ and the variance is $(\beta - \alpha)^2/12$. The most commonly encountered version of this density function is one in which the parameters α and β take the values 0 and 1 respectively, and is used in generating quasi-random numbers.

Exponential

The PDF of a continuous random variable, X , taking only positive values. The density function is given by

$$f(x) = \lambda e^{-\lambda x}, \quad x > 0 \quad (9)$$

The single parameter of the exponential density function, λ , determines the shape of the density function, as we see in Figure 6, where a number of different exponential density functions are shown. The mean of an exponential variable is $1/\lambda$ and the variance is $1/\lambda^2$.

The exponential density plays an important role in some aspects of reliability engineering and also gives the distribution of the time intervals between independent consecutive random events such as particles emitted by radioactive material.

Beta

A PDF for a continuous random variable, X , taking values between 0 and 1. The density function is defined as

$$f(x) = \frac{\Gamma(p+q)}{\Gamma(p)\Gamma(q)} x^{p-1} (1-x)^{q-1}, \quad 0 < x < 1$$

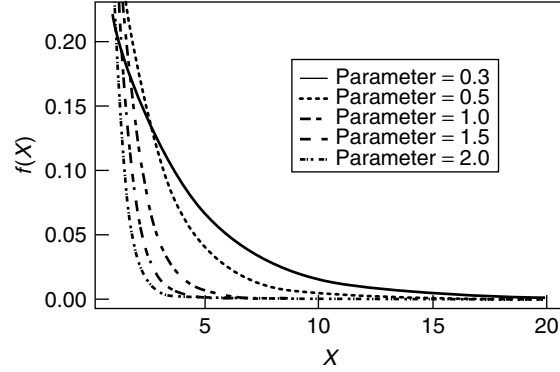


Figure 6 Exponential density functions

where $p > 0$ and $q > 0$ are parameters that define particular aspects of the shape of the beta density and Γ is the gamma function (see [3]). The mean of a beta variable is $p/(p+q)$ and the variance is $pq/(p+q)^2(p+q+1)$.

Gamma

A PDF for a random variable X that can take only positive values. The density function is defined as

$$f(x) = \frac{1}{\Gamma(\alpha)} x^{\alpha-1} e^{-x}, \quad x > 0 \quad (11)$$

where the single parameter, α , is both the mean of X and also its variance.

The gamma density function can often act as a useful model for a variable that cannot reasonably be assumed to have a normal density function because of its positive **Skewness**.

Chi-Squared

The PDF of the **sum of squares** of a number (ν) of independent standard normal variables given by

$$f(x) = \frac{1}{2^{\nu/2} \Gamma(\nu/2)} x^{(\nu-2)/2} e^{-x/2}, \quad x > 0 \quad (12)$$

that is, a gamma density with $\alpha = \nu/2$. The parameter ν is usually known as the **degrees of freedom** of the density. The mean and variance of a χ^2 distribution are ν and 2ν respectively. This density function

6 Probability Density Functions

arises in many areas of statistics, most commonly as the null distribution of the chi-squared goodness of fit statistics, or as the distribution associated with estimators of variance.

Student's T

The PDF of a variable defined as

$$t = \frac{\bar{X} - \mu}{s/\sqrt{n}} \quad (13)$$

where $\bar{X}$ is the arithmetic mean of n observations from a normal density with mean μ and s is the sample standard deviation. The density function is given by

$$f(t) = \frac{\Gamma\left\{\frac{1}{2}(\nu+1)\right\}}{(\nu\pi)^{1/2}\Gamma\left(\frac{1}{2}\nu\right)} \left(1 + \frac{t^2}{\nu}\right)^{-\frac{1}{2}(\nu+1)},$$

$$-\infty < t < \infty \quad (14)$$

where $\nu = n - 1$. The shape of the density function varies with ν , and as ν gets larger the shape of the t -density function approaches that of a standard normal. This density function is the null distribution of Student's t -statistic used for testing hypotheses about population means (see **Parametric Tests**).

Fisher's F

The PDF of the ratio of two independent random variables each having a χ^2 distribution, divided by their respective degrees of freedom. The form of the density is that of a beta density with $p = \nu_1/2$ and $q = \nu_2/2$, where ν_1 and ν_2 are, respectively, the degrees of freedom of the numerator and denominator χ^2 variables. Fisher's F density is used to assess the equality of variances in, for example, the F -test and the analysis of variance (see **Parametric Tests**).

Weibull distribution

Named after the Swedish Professor Waloddi Weibull, the class of Weibull distributions is sometimes thought of as a generalization of the exponential distribution (which it includes as a special case). The PDF is

$$f(x) = \frac{c}{b} \left(\frac{x-a}{b}\right)^{c-1} \exp\left\{-\left[\frac{(x-a)}{b}\right]^c\right\}, x > a \quad (15)$$

The parameters a and b are location and scale parameters (a is sometimes called the *threshold*, and is often assumed to be 0), while c is a shape parameter. For $a = 0$ and $c = 1$ we obtain the density of the exponential distribution. Weibull distributions are used often in lifetime and time-to-failure models.

Lognormal distribution

If for some value of a , $\ln(X - a)$ has a **normal distribution** with mean μ and variance σ^2 , then X has a lognormal distribution with parameters a , μ and σ . Reparameterizing as $b = \exp(\mu)$, a , b , and c are location, scale, and shape parameters. The PDF is

$$f(x) = \frac{1}{\sigma(x-a)\sqrt{2\pi}} \exp\left[-\frac{1}{2\sigma^2} \left\{\log\left(\frac{x-a}{b}\right)\right\}^2\right],$$

$$x > a \quad (16)$$

The mean and variance are $E(X) = a + b \exp(\sigma^2/2)$ and $Var(X) = b^2 \exp(\sigma^2)[\exp(\sigma^2) - 1]$. The lognormal distribution is a possible model for variables that have a skewed distribution.

Multivariate Density Functions

Multivariate density functions play the same role for *vector random variables* as the density functions described earlier play in the univariate situation. Here we shall look at two such density functions, the multinomial and the multivariate normal.

Multinomial

The multinomial PDF is a multivariate generalization of the binomial density function described earlier. The density function is associated with a vector of random variables $X' = [X_1, X_2, \dots, X_k]$ which arises from a sequence of n independent and identical trials each of which can result in one of k possible mutually exclusive and collectively exhaustive

events, with probabilities $p_1, p_2, \dots, p_k$. The density function is defined as follows:

$$P(X_1 = x_1, X_2 = x_2, \dots, X_k = x_k) = \frac{n!}{x_1! x_2! \dots x_k!} p_1^{x_1} p_2^{x_2} \dots p_k^{x_k} \quad (17)$$

where x_i is the number of trials with outcome i . The expected value of X_i is np_i and its variance is $np_i(1 - p_i)$. The **covariance** of X_i and X_j is $-np_i p_j$.

Multivariate Normal

The PDF of a vector of continuous random variables, $\mathbf{X}' = [X_1, X_2, \dots, X_p]$ defined as follows

$$f(\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_p) = (2\pi)^{-p/2} |\Sigma| \times \exp \left[-\frac{1}{2} (\mathbf{x} - \boldsymbol{\mu})' \Sigma^{-1} (\mathbf{x} - \boldsymbol{\mu}) \right] \quad (18)$$

where $\boldsymbol{\mu}$ is the mean vector of the variables and Σ is their covariance matrix.

The simplest version of the multivariate normal density function is that for two random variables and known as the *bivariate normal density function*. The density function formula given above now reduces to

$$f(x_1, x_2) = \frac{1}{2\pi\sigma_1\sigma_2\sqrt{1-\rho^2}} \exp \left\{ -\frac{1}{1-\rho^2} \times \left[\frac{(x_1 - \mu_1)^2}{\sigma_1^2} - 2\rho \frac{(x_1 - \mu_1)}{\sigma_1} \times \frac{(x_2 - \mu_2)}{\sigma_2} + \frac{(x_2 - \mu_2)^2}{\sigma_2^2} \right] \right\} \quad (19)$$

where $\mu_1, \mu_2, \sigma_1, \sigma_2, \rho$ are, respectively, the means, standard deviations, and **correlation** of the two variables. Perspective plots of a number of such density functions are shown in Figure 7.

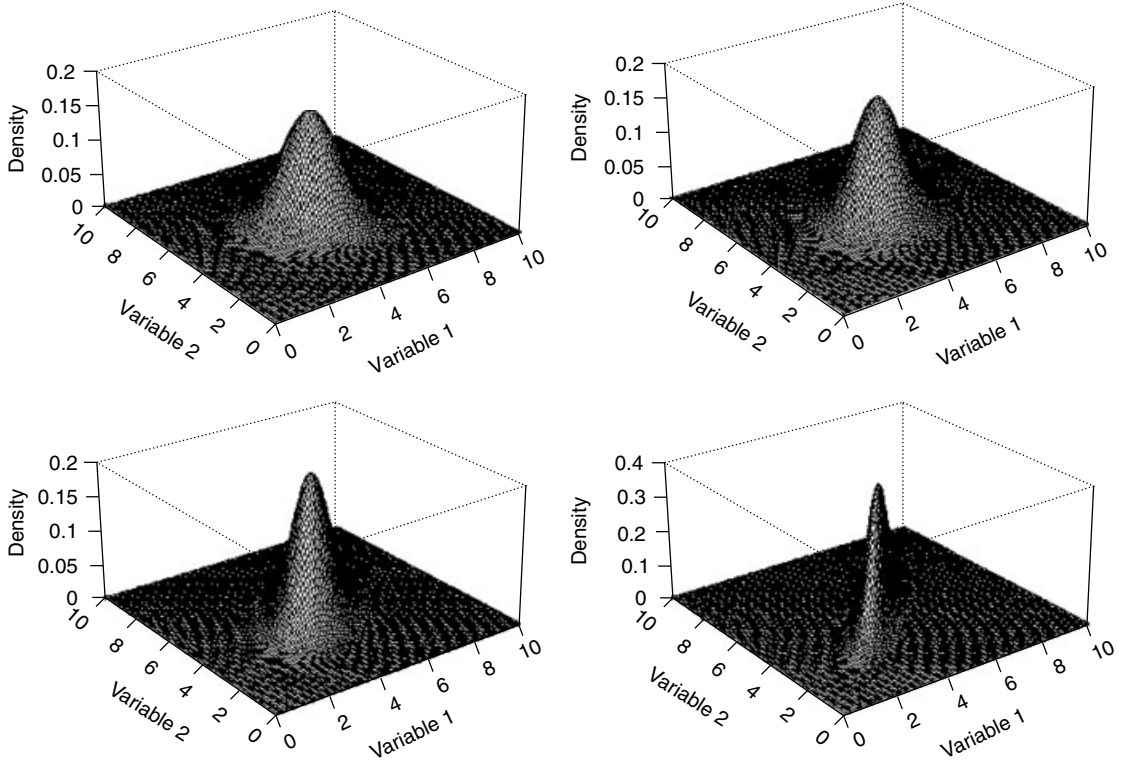


Figure 7 Four bivariate normal density functions with means 5 and standard deviations 1, for variables 1 and 2, and with correlations equal to 0.0, 0.3, 0.6, and 0.9

References

- [1] Balakrishnan, N. & Nevzorov, V.B. (2003). *A Primer on Statistical Distributions*, John Wiley & Sons, Hoboken.
- [2] Evans, M., Hastings, N. & Peacock, B. (2000). *Statistical Distributions*, 3rd Edition, John Wiley & Sons, New York.
- [3] Everitt, B.S. (2002). *Cambridge Dictionary of Statistics*, 2nd Edition, Cambridge University Press, Cambridge.

BRIAN S. EVERITT

Article originally published in Encyclopedia of Statistics in Behavioral Science (2005, John Wiley & Sons, Ltd). Minor revisions for this publication by Jeroen de Mast.

Collinearity

Collinearity refers to the intercorrelation among independent variables (or predictors). When more than two independent variables are highly correlated, collinearity is termed *multicollinearity*, but the two terms are interchangeable. Two or more independent variables are said to be collinear when they are perfectly correlated (i.e., $r = -1.0$ or $+1.0$). However, researchers also use the term *collinearity* when independent variables are highly, yet not perfectly, correlated (e.g., $r > 0.70$), a case formally termed *near collinearity* [1]. Collinearity is generally expressed as the multiple **correlation** among the set of i independent variables, or R_i .

Perfect collinearity violates the assumption of ordinary least-squares regression that there is no perfect correlation between any set of two independent variables. In the context of *multiple linear regression*, perfect collinearity renders regression weights of the collinear independent variables uninterpretable, and near collinearity results in biased parameter estimates and inflated standard errors associated with the parameter estimates [1].

Collinearity is obvious in multiple regression when the parameter estimates differ in magnitude and, possibly, in direction from their respective bivariate correlations and when the standard errors of the estimates are high. More formally, the *variance*

inflation factor (VIF) and *tolerance* diagnostics can help detect the presence of collinearity. VIF_i , reflecting the level of collinearity among a set of i independent variables, is defined as $1 / (1 - R_i^2)$, where R_i^2 are the squared multiple correlations among the i independent variables (see **Coefficient of Determination (R^2)**). A *VIF* value of 1 indicates no level of collinearity, and values greater than 1 indicate some level of collinearity. Tolerance is defined as $1 - R_i^2$ or $1/VIF$, where values equal to 1 indicate lack of collinearity, and 0 indicates perfect collinearity. Several remedies exist for reducing or eliminating the problems associated with collinearity, including dropping one or more collinear predictors, combining collinear predictors into a single composite, and employing ridge regression or principal components regression analysis (e.g., [2]).

References

- [1] Draper, N.R. & Smith, H. (1998). *Applied Regression Analysis*, 3rd Edition, John Wiley & Sons.
- [2] Netter, J., Wasserman, W. & Kutner, M.H. (1990). *Applied Linear Statistical Models: Regression, Analysis of Variance, and Experimental Designs*, 3rd Edition, Irwin Publishing, Boston.

GILAD CHEN

Article originally published in Encyclopedia of Statistics in Behavioral Science (2005, John Wiley & Sons, Ltd). Minor revisions for this publication by Jeroen de Mast.

Confidence Intervals

In textbooks, statistical packages, and statistics courses, confidence intervals are somewhat under-represented compared to **hypothesis testing**. A null hypothesis is a precise statement about one or more population parameters from which, given certain assumptions, a probability distribution for a relevant sample statistic can be derived (see **Hypothesis Testing**). As the name implies, a null hypothesis is usually a hypothesis that no real effect is present (two populations have identical means, two variables are not related in the population, etc.). A decision is made to reject or not reject such a null hypothesis on the basis of a significance test applied to appropriately collected sample data. A significant result is taken to mean that the observed data are not consistent with the null hypothesis, and that a real effect is present. The **significance level** represents the conditional probability that the process would lead to the rejection of the null hypothesis given that it is, in fact, true.

Critics have argued that rejecting or not rejecting a null hypothesis of, say, no difference in population means, does not constitute a full analysis of the data or advance knowledge very far. They have implored researchers to report other aspects of their data such as the size of any observed effects. In particular, they have advised researchers to calculate and report confidence intervals for effects of interest [1, 2]. However, it would be a mistake to believe that relying more on confidence intervals would make significance tests redundant. Confidence intervals can be viewed as a generalization of the null hypothesis significance test, but in order to understand this, the definition of a null hypothesis has to be broadened somewhat. It is, in fact, possible for a null hypothesis to specify an effect size other than zero, and such a hypothesis can be subjected to significance testing provided it allows the derivation of a relevant sampling distribution. To distinguish such null hypotheses from the more familiar null hypothesis of no effect, the term *nil hypothesis* was introduced for the latter [3]. While the significance test of the nil hypothesis indicates whether the data are consistent with a hypothesis that says there is no effect (or an effect of zero size), a confidence interval indicates all those effect sizes with which the data are and are not consistent.

The construction of a confidence interval will be introduced by considering the **estimation** of the mean of a normal population of unknown variance from the information contained in a simple random sample from that population. The formula for the confidence interval is a simple rearrangement of that of the single sample *t* test (see **Parametric Tests**):

$$\bar{X} + t_{0.025, n-1} \left(\frac{\hat{s}}{\sqrt{n}} \right) \leq \mu \leq \bar{X} + t_{0.975, n-1} \left(\frac{\hat{s}}{\sqrt{n}} \right) \quad (1)$$

where $\bar{X}$ is the sample mean, $\hat{s}$ is the unbiased sample standard deviation, n is the **sample size**, $(\hat{s}/\sqrt{n})$ is an estimate of the **standard error** of the mean, and $t_{0.025, n-1}$ and $t_{0.975, n-1}$ are the critical values of t with $n - 1$ **degrees of freedom** that cut off the bottom and top 2.5% of the distribution, respectively. This is the formula for a 95% confidence interval for the population mean of a **normal distribution** whose standard deviation is unknown. In repeated applications of the formula to simple random samples from normal populations, the constructed interval will contain the true population mean on 95% of occasions, and not contain it on 5% of occasions. Intervals for other levels of confidence are easily produced by substituting appropriate critical values. For example, if a 99% confidence interval was required, $t_{0.005, n-1}$ and $t_{0.995, n-1}$ would be employed.

The following example will serve to illustrate the calculation. A sample of 20 products was drawn at random from a manufacturing process and a quality measurement was applied to each. The sample mean was 283.09 and the unbiased sample standard deviation was 51.42. Assuming a normal distribution, the appropriate critical values of t with 19 degrees of freedom are -2.093 and $+2.093$. Substituting in equation (1), gives the 95% confidence interval for the mean of the population as $259.02 \leq \mu \leq 307.16$. It will be clear from the formula that the computed interval is centered on the sample mean and extends a number of estimated standard errors above and below that value. As sample size increases, the exact number of estimated standard errors above and below the mean approaches 1.96 as the t distribution approaches the standardized normal distribution. Of course, as sample size increases, the estimated standard error of the mean decreases and the confidence interval gets

2 Confidence Intervals

narrower, reflecting the greater accuracy of estimation resulting from a larger sample.

The Coverage Interpretation of Confidence Intervals

Ninety five percent of intervals calculated in the above way will contain the appropriate population mean, and 5% will not. This property is referred to as *coverage* and it is a property of the procedure rather than that of a particular interval. When it is claimed that an interval contains the estimated population mean with 95% confidence, what should be understood is that the procedure used to produce the interval will give 95% coverage in the long run. However, for any particular interval that has been calculated, it will be either one of the 95% that contain the population mean or one of the 5% that do not. There is no way of knowing which of these two is true. According to the frequentist view of probability from which the method derives [4], for any particular confidence interval the population is either included or not included in it. The probability of inclusion is therefore either 1 or 0. The confidence interval approach does not give the probability that the true population mean will be in the particular interval constructed, and the term *confidence* rather than *probability* is employed to reinforce the distinction. Despite the claims of the authors of many elementary textbooks [5], significance tests do not result in statements about the probability of the truth of a hypothesis, and confidence intervals, which derive from the same statistical theory, cannot do so either. The relationship between confidence intervals and significance testing is examined in the next section.

The Significance Test Interpretation of Confidence Intervals

The lower limit of the 95% confidence interval calculated above is 259.02. If the sample data were used to conduct a two-tailed, one-sample t test at the 0.05 significance level to test the null hypothesis that the true population mean was equal to 259.02, it would yield a t value that was just non-significant. However, any null hypothesis that specified that the population mean was a value less than 259.02 would be rejected by a test carried out

on the sample data. The situation is similar with the upper limit of the calculated interval, 307.16. The term *plausible* has been applied to values of a population parameter that are included in a confidence interval [6, 7]. According to this usage, the 95% confidence interval calculated above showed any value between 259.02 and 307.16 to be a “plausible” value for the mean of the population, as such values would not be rejected by a two-tailed test at the 0.05 level carried out on the obtained sample data.

One-Sided Confidence Intervals

In some circumstances, interest is focused on estimating only the highest (or lowest) plausible value for a population parameter. The distinction between two-sided intervals and one-sided intervals mirrors that between two-tailed and one-tailed significance tests. The formula for one of the one-sided 95% confidence intervals for the population mean, derived from the one-tailed single sample t test is as follows:

$$-\infty \leq \mu \leq \bar{X} + t_{0.950, n-1} \left(\frac{\hat{s}}{\sqrt{n}} \right) \quad (2)$$

Applying equation (2) to the example data above yields:

$$-\infty \leq \mu \leq 302.97 \quad (3)$$

In the long run, intervals calculated in this way will contain the true population value for 95% of appropriately collected samples.

Central and Noncentral Confidence Intervals

The two-sided and one-sided 95% confidence intervals calculated above represent the extremes of a continuum. With the two-sided interval, the 5% rejection region of the significance test was split equally between the two tails giving limits of 259.02 and 307.16. In the long run, such intervals are exactly as likely to miss the true population mean by having a lower bound that is too high as by having an upper bound that is too low. Intervals with this property are known as *central*. However, with the one-sided interval, which ranged from $-\infty$ to 302.97, it is impossible for the interval to lie above the population mean, but in the long run 5% of such intervals

will fall below the population mean. This is the most extreme case of a noncentral interval.

It is possible to divide the 5% rejection region in an infinite number of ways between the two tails. All that is required is that the tails sum to 0.05. If the lower rejection region was made to be 1% and the upper region made to be 4% by employing $t_{0.010, n-1}$ and $t_{0.960, n-1}$ as the critical values, then 95% of intervals calculated in this way would include the true population mean and 5% would not. Such intervals would be four times more likely to miss the population mean by being too low than by being too high.

While arguments for employing noncentral confidence intervals have been put forward in other disciplines [8], central confidence intervals are usually employed in industrial statistics. In a number of important applications, such as the construction of confidence intervals for population means (and their differences), central confidence intervals have the advantage that they are shorter than noncentral intervals and therefore give the required level of confidence for the shortest range of plausible values.

The Confidence Interval for the Difference between Two Population Means

If the intention is to compare means, then the appropriate approach is to construct a confidence interval for the difference in population means using a rearrangement of the formula for the two independent samples t tests (see **Parametric Tests**):

$$\begin{aligned}
 & (\bar{X}_1 - \bar{X}_2) \pm t_{0.025, n_1 + n_2 - 2} \\
 & \times \sqrt{\left(\frac{(n_1 - 1)\hat{s}_1^2 + (n_2 - 1)\hat{s}_2^2}{n_1 + n_2 - 2} \right) \left(\frac{1}{n_1} + \frac{1}{n_2} \right)} \\
 & \leq \mu_1 - \mu_2 \leq (\bar{X}_1 - \bar{X}_2) \pm t_{0.975, n_1 + n_2 - 2} \\
 & \times \sqrt{\left(\frac{(n_1 - 1)\hat{s}_1^2 + (n_2 - 1)\hat{s}_2^2}{n_1 + n_2 - 2} \right) \left(\frac{1}{n_1} + \frac{1}{n_2} \right)} \quad (4)
 \end{aligned}$$

where $\bar{X}_1$ and $\bar{X}_2$ are the two sample means, $\hat{s}_1$ and $\hat{s}_2$ are the unbiased sample standard deviations, and n_1 and n_2 are the two sample sizes. This is the formula for the 95% confidence interval for the difference between the means of two normal populations, which are assumed to have the

same (unknown) variance. When applied to data from two appropriately collected samples, over the course of repeated application, the upper and lower limits of the interval will contain the true difference in population means for 95% of those applications. The formula is for a central confidence interval that is equally likely to underestimate as it is to overestimate, the true difference in population means. Since the sampling distribution involved is symmetrical, this is the shortest interval providing 95% confidence.

If equation (4) is applied to the sample data that were introduced previously ($\bar{X}_1 = 283.09$, $\hat{s}_1 = 51.42$, and $n_1 = 20$) and to a second sample ($\bar{X}_2 = 328.40$, $\hat{s}_2 = 51.90$, and $n_2 = 21$), the 95% confidence interval for the difference in the two population means ranges from -77.96 to -12.66 (which is symmetrical about the observed difference between the sample means of -45.31). Since this interval does not contain zero, the difference in population means specified by the nil null hypothesis, it is readily apparent that the nil null hypothesis can be rejected at the 0.05 significance level. Therefore, the confidence interval provides all the information provided by a two-tailed test of the nil null hypothesis at the corresponding conventional significance level.

In addition, the confidence interval implies that a two-tailed independent samples t test carried out on the observed data would not lead to the rejection of any null hypothesis that specified the difference in population means was equal to a value in the calculated interval, at the 0.05 significance level. This information is not provided by the usual test of the nil null hypothesis, but is a unique contribution of the confidence interval approach. However, the test of the nil null hypothesis of no population mean difference yields an exact P value for the above samples of 0.008. So the nil null hypothesis can be rejected at a stricter significance level than the 0.05 implied by the 95% confidence interval. So while the confidence interval approach provides a test of all possible null hypotheses at a single conventional significance level, the significance testing approach provides a more detailed evaluation of the extent to which the data cast doubt on just one of these null hypotheses.

If, as in this example, all the “plausible” values for the difference in the means of two populations are of the same sign, then Tukey [9] suggests that it is appropriate to talk of a “confident direction” for

the effect. Here, the 95% confidence interval shows that the confident direction of the effect is that the first population mean is less than the second.

The Confidence Interval for the Difference between Two Population Means and Confidence Intervals for the Two Individual Population Means

The 95% confidence interval for the mean of the population from which the first sample was drawn has been previously shown to range from 259.03 to 307.16. The second sample yields a 95% confidence interval ranging from 304.77 to 352.02. Values between 304.77 and 307.16 are common to both intervals and therefore “plausible” values for the means of both populations. The overlap of the two intervals would be visually apparent if the intervals were presented graphically, and the viewer might be tempted to interpret the overlap as demonstrating the absence of a significant difference between the means of the two samples. However, the confidence interval for the difference between the population means did not contain zero, so the two sample means do differ significantly at the 0.05 level. This example serves to illustrate that caution is required in the comparison and interpretation of two or more confidence intervals. Confidence intervals and significance tests derived from the data from single samples are often based on different assumptions from those derived from the data from two (or more) samples. The former, then, cannot be used to determine the results of the latter directly.

Confidence Intervals Based on Approximations to a Normal Distribution

As sample sizes increase, the sampling distributions of many statistics tend toward a normal distribution with a standard error that is estimated with increasing precision. This is true, for example, of the mean of a random sample from some unknown nonnormal population, where an interval extending from 1.96 estimated standard errors below the observed sample mean to the same distance above the sample mean gives coverage closer and closer to 95% as sample size increases. A rough and ready 95% confidence interval in such cases would range from two

estimated standard errors below the sample statistic to two above. Intervals ranging from one estimated standard error below the sample statistic to one above would give approximately 68% coverage. Confidence intervals for population **correlation** coefficients are often calculated in this way, as Fisher’s z transformation of the sample Pearson correlation (r) tends toward a normal distribution of known variance as sample size increases. The procedure leads to intervals that are symmetrical about the transformed sample correlation (z_r), but the symmetry is lost when the limits of the interval are transformed back to r .

These are large sample approximations, however, and some caution is required in their interpretation. The use of the normal approximation to the binomial distribution in the construction of a confidence interval for a population proportion will be used to illustrate some of the issues.

Constructing Confidence Intervals for the Population Proportion

Under repeated random sampling from a population where the population proportion is π , the sample proportion, p , follows a discrete, binomial distribution (*see Probability Density Functions*), with mean π and standard deviation $\sqrt{\pi(1-\pi)/n}$. As n increases, provided π is not too close to zero or one, the distribution of p tends toward a normal distribution. In these circumstances, $\sqrt{p(1-p)/n}$ can provide a reasonably good estimate of $\sqrt{\pi(1-\pi)/n}$, and an approximate 95% confidence interval for π is provided by

$$p - 1.96\sqrt{\frac{p(1-p)}{n}} \leq \pi \leq p + 1.96\sqrt{\frac{p(1-p)}{n}} \quad (5)$$

In many circumstances, this may provide a serviceable confidence interval. However, the approach uses a continuous distribution to approximate a discrete one and makes no allowance for the fact that $\sqrt{p(1-p)/n}$ cannot be expected to be exactly equal to $\sqrt{\pi(1-\pi)/n}$. Particular applications of this approximation can result in intervals with an impossible limit (less than zero or greater than one) or of zero width (when p is zero or one).

A more satisfactory approach to the question of constructing a confidence interval for a population

proportion, particularly when the sample size is small, is to return to the relationship between the plausible parameters included in a confidence interval and statistical significance. According to this, the lower limit of the 95% confidence interval should be that value of the population proportion, π_L , which if treated as the null hypothesis, would just fail to be rejected in the upper tail of a two-tailed test at the 0.05 level on the observed data. Similarly, the upper limit of the 95% confidence interval should be that value of the population proportion, π_U , which if treated as the null hypothesis, would just fail to be rejected in the lower tail of a two-tailed test at the 0.05 level on the observed data. However, there are particular problems in conducting two-tailed tests on a discrete variable like the sample proportion [10]. As a result, the approach taken is to work in terms of one-tailed tests and to find the value of the population proportion π_L , which if treated as the null hypothesis, would just fail to be rejected by a one-tailed test (in the upper tail) at the 0.025 level on the observed data, and the value π_U , which if treated as the null hypothesis, would just fail to be rejected by a one-tailed test (in the lower tail) at the 0.025 level on the observed data [11].

The approach will be illustrated with a small sample example that leads to obvious inconsistencies when the normal approximation is applied. A simple random sample of products is drawn from a manufacturing process. Of the 20 products in the sample, two are defective, while 18 are free from defects. The task is to construct a 95% confidence interval for the proportion of products that is free from defects.

Using the normal approximation method, equation (5) would yield:

$$0.9 - 1.96\sqrt{\frac{0.9(1-0.9)}{20}} \leq \pi \leq 0.9 + 1.96\sqrt{\frac{0.9(1-0.9)}{20}} \quad (6)$$

The upper limit of this symmetrical interval is outside the range of the parameter. Even if it was treated as one, it is clear that if this were made the null hypothesis of a one-tailed test at the 0.025 level, then it must be rejected on the basis of the observed data, since if the process delivered 100% defect-free products, the sample could not contain any defective products, but it did contain two.

The exact approach proceeds as follows: That value of the population proportion, π_L , is found such that under this null hypothesis the sum of the probabilities of 18, 19, and 20 defect-free products in the sample of 20 would equal 0.025. Reference to an appropriate table [12] or statistical package (e.g., Minitab) yields a value of 0.683017. Then, that value of the population proportion, π_U , is found such that under this null hypothesis the sum of the probabilities of 0, 1, 2, ..., 18 defect-free products in a sample of 20 would be 0.025. The desired value is 0.987651.

So the exact 95% confidence interval for the population proportion of products is

$$0.683017 \leq \pi \leq 0.987651 \quad (7)$$

If a one-tailed test at the 0.025 level was conducted on any null hypothesis that specified that the population proportion was less than 0.683017, it would be rejected on the basis of the observed data. So also would any null hypothesis that specified that the population proportion was greater than 0.987651.

In contrast to the limits indicated by the approximate method, those produced by the exact method will always lie within the range of the parameter and will not be symmetrical around the sample proportion (except when this is exactly 0.5). The exact method yields intervals that are conservative in coverage [13, 14], and alternatives have been suggested. One approach questions whether it is appropriate to include the full probability of the observed outcome in the tail when computing π_L and π_U . Berry and Armitage [11] favor mid-P confidence intervals where only half the probability of the observed outcome is included. The Mid-P intervals will be rather narrower than the exact intervals, and will result in a rather less conservative coverage. However, in some circumstances the Mid-P coverage can fall below the intended level of confidence, which cannot occur with the exact method [14].

Confidence Intervals and the Bootstrap

While the central limit theorem might provide support for the use of the normal distribution in constructing approximate confidence intervals in some situations, there are other situations where sample sizes are too small to justify the process, or sampling distributions

are suspected or known to depart from normality. The **bootstrap** (“Overview of **resampling**, (*see* **Bootstrap and Jackknife, Overview**) bootstrapping and jackknifing” in the CIM section (has no id. yet)) [6, 15] is a numerical method designed to derive some idea of the sampling distribution of a statistic without recourse to assumptions of unknown or doubtful validity. The approach is to treat the collected sample of data as if it were the population of interest. A large number of random samples are drawn with replacement from this “population”, each sample being of the same size as the original sample. The statistic of interest is calculated for each of these “resamples” and an empirical sampling distribution is produced. If a large enough number of resamplings is performed, then a 95% confidence interval can be produced directly from these by identifying those values of the statistic that correspond to the 2.5th and 97.5th **percentiles**. On other occasions, the standard deviation of the statistic over the “resamples” is calculated and this is used as an estimate of the true standard error of the statistic in one of the standard formulas. Various ways of improving the process have been developed and subjected to some empirical testing. With small samples or with samples that just by chance are not very representative of their parent population, the method may provide rather unreliable information. Nonetheless, Efron and Tibshirani [6] provide evidence that it works well in many situations, and the methodology is becoming increasingly popular.

Bayesian Certainty Intervals

Confidence intervals derive from the frequentist approach to statistical inference that defines probabilities as the limits in the long run of relative frequencies. According to this view, the confidence interval is the random variable, not the population parameter [4, 8]. In consequence, when a 95% confidence interval is constructed, the frequentist position does not allow statements of the kind “The value of the population parameter lies in this interval with probability 0.95.” In contrast, the Bayesian approach to statistical inference (*see* **Probability Theory**) defines probability as a subjective or psychological variable [16]. According to this view, it is not only acceptable to claim that particular probabilities are associated with particular values of a parameter; such claims are the

very basis of the inference process. Bayesians use Bayes theorem to update their prior beliefs about the probability distribution of a parameter on the basis of new information. Sometimes, this updating process uses elements of the same statistical theory that is employed by a frequentist and the final result may appear to be identical to a confidence interval. However, the interpretation placed on the resulting interval could hardly be more different.

Suppose that a Bayesian was interested in estimating the mean score of some population on a particular measurement scale. The first part of the process would involve attempting to sketch a prior probability distribution for the population mean, showing what value seemed most probable, whether the distribution was symmetrical about this value, how spread out the distribution was, perhaps trying to specify the parameter values that contained the middle 50% of the distribution, the middle 90%, and so on. This process would result in a subjective prior distribution for the population mean. The Bayesian would then collect sample data, in very much the same way as a frequentist, and use Bayes theorem to combine the information from this sample with the information contained in the prior distribution to compute a posterior probability distribution. If the 2.5th and 97.5th percentiles of this posterior probability distribution are located, they form the lower and upper bounds on the Bayesian’s 95% certainty interval [16].

Certain parallels and distinctions can be drawn between the confidence interval of the frequentist and the certainty interval of the Bayesian. The center of the frequentist’s confidence interval would be the sample mean, but the center of the Bayesian’s certainty interval would be somewhere between the sample mean and the mean of the prior distribution. The smaller the variance of the prior distribution, the nearer the center of the certainty interval will be to that of the prior distribution. The width of the frequentist’s confidence interval would depend only on the sample standard deviation and the sample size, but the width of the Bayesian’s certainty interval would be narrower, as the prior distribution contributes a “virtual” sample size that can be combined with the actual sample size, to yield a smaller posterior estimate of the standard error as well as contribute to the degrees of freedom of any required critical values.

So, to the extent that the Bayesian had any views about the likely values of the population mean before

data were collected, these views will influence the location and width of the 95% certainty interval. Any interval so constructed is to be interpreted as a probability distribution for the population mean, permitting the Bayesian to make statements like, "The probability that μ lies between L and U is X." Of course, none of this is legitimate to the frequentist who distrusts the subjectivity of the Bayesian's prior distribution, its influence on the resulting interval, and the attribution of a probability distribution to a population parameter that can only have one value.

If the Bayesian admits ignorance of the parameter being estimated prior to the collection of the sample, then a uniform prior might be specified which says that any and all values of the population parameter are equally likely. In these circumstances, the certainty interval of the Bayesian and the confidence interval of the frequentist will be identical in numerical terms. However, an unbridgeable gulf will still exist between the interpretations that would be made of the interval.

References

- [1] Meehl, P.E. (1997). The problem is epistemology, not statistics: replace significance tests by confidence intervals and quantify accuracy of risky numerical predictions, in *What if there were no Significance Tests?* L.L. Harlow, S.A. Mulaik & J.H. Steiger, eds, Lawrence Erlbaum Associates, Mahwah.
- [2] Wilkinson, L., The Task Force on Statistical Inference. (1999). Statistical methods in psychology journals: guidelines and explanations, *American Psychologist* **54**, 594–604.
- [3] Cohen, J. (1994). The earth is round ($p < .05$), *American Psychologist* **49**, 997–1003.
- [4] Neyman, J. (1937). Outline of a theory of statistical estimation based on the classical theory of probability, *Philosophical Transactions of the Royal Society, A* **236**, 333–380.
- [5] Dracup, C. (1995). Hypothesis testing: what it really is, *Psychologist* **8**, 359–362.
- [6] Efron, B. & Tibshirani, R.J. (1993). *An Introduction to the Bootstrap*, Chapman & Hall, New York.
- [7] Smithson, M. (2000). *Statistics with Confidence*, Sage, London.
- [8] Kendall, M.G. & Stuart, A. (1979). *The Advanced Theory of Statistics. Vol. 2. Inference and Relationship*, 4th Edition, Griffin, London.
- [9] Tukey, J. (1991). The philosophy of multiple comparisons, *Statistical Science* **6**, 100–116.
- [10] Macdonald, R.R. (1998). Conditional and unconditional tests of association in 2×2 tables, *British Journal of Mathematical and Statistical Psychology* **51**, 191–204.
- [11] Berry, G. & Armitage, P. (1995). Mid-P confidence intervals: a brief review, *Statistician* **44**, 417–423.
- [12] Geigy Scientific Tables. (1982). *Vol. 2: Introduction to Statistics, Statistical Tables, Mathematical Formulae*, 8th Edition, Ciba-Geigy, Basle.
- [13] Neyman, J. (1935). On the problem of confidence intervals, *Annals of Mathematical Statistics* **6**, 111–116.
- [14] Vollset, S.E. (1993). Confidence intervals for a binomial proportion, *Statistics in Medicine* **12**, 809–824.
- [15] Davison, A.C. & Hinkley, D.V. (2003). *Bootstrap Methods and their Application*, Cambridge University Press, Cambridge.
- [16] Box, G.E.P. & Tiao, G.C. (1992). *Bayesian Inference in Statistical Analysis*, John Wiley & Sons.

CHRIS DRACUP

Article originally published in Encyclopedia of Statistics in Behavioral Science (2005, John Wiley & Sons, Ltd). Minor revisions for this publication by Jeroen de Mast.

Estimation

Introduction

In the simplest case, a data set consists of observations on a single variable, say real-valued observations. Suppose there are n such observations, denoted by $X_1, \dots, X_n$. For example, X_i could be a quality characteristic of product i , or the number of car accidents on day i , and so on. Suppose now that each observation follows the same probability law P . This means that the observations are relevant if one wants to predict the value of a new observation X (say, the quality of a new product, or the number of car accidents on a future day, etc.). Thus, a common underlying distribution P allows one to generalize the outcomes.

The emphasis in this paper is on the data and estimators derived from the data, and less on the estimation of population parameters describing a model for P . An estimator is any given function $T_n(X_1, \dots, X_n)$ of the data. Let us start with reviewing some common estimators.

The Empirical Distribution

The unknown P can be estimated from the data in the following way. Suppose first that we are interested in the probability that an observation falls in A , where A is a certain set chosen by the researcher. We denote this probability by $P(A)$. Now, from the frequentist point of view, a probability of an event is the limit of relative frequencies of occurrences of that event as the number of occasions of possible occurrences n grows without limit. So, it is natural to estimate $P(A)$ by the frequency of A that is, by

$$\begin{aligned} \hat{P}_n(A) &= \frac{\text{number of times an observation } X_i \text{ falls in } A}{\text{total number of observations}} \\ &= \frac{\text{number of } X_i \in A}{n} \end{aligned} \quad (1)$$

We now define the empirical distribution $\hat{P}_n$ as the probability law that assigns to a set A the probability $\hat{P}_n(A)$. We regard $\hat{P}_n$ as an estimator of the unknown P .

The Empirical Distribution Function

The distribution function of X is defined as

$$F(x) = P(X \leq x) \quad (2)$$

and the empirical distribution function is

$$\hat{F}_n(x) = \frac{\text{number of } X_i \leq x}{n} \quad (3)$$

Figure 1 plots the distribution function $F(x) = 1 - 1/x^2$, $x \geq 1$ (smooth curve) and the empirical distribution function $\hat{F}_n$ (stair function) of a sample from F with **sample size** $n = 200$.

Sample Moments and Sample Variance

The theoretical mean

$$\mu = E(X) \quad (4)$$

(E stands for *Expectation*), can be estimated by the sample average

$$\bar{X}_n = \frac{X_1 + \dots + X_n}{n} \quad (5)$$

Generally, for $j = 1, 2, \dots$ the j th sample moment

$$\hat{\mu}_{j,n} = \frac{X_1^j + \dots + X_n^j}{n} \quad (6)$$

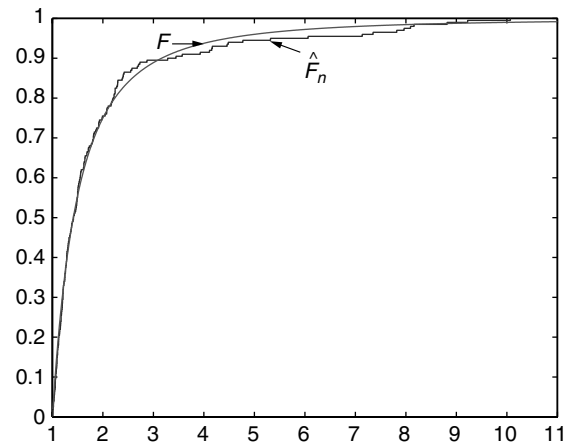


Figure 1 The empirical and theoretical distribution function

2 Estimation

is an estimator of the j th moment $E(X^j)$ of P (see **Moments**).

The sample variance

$$S_n^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X}_n)^2 \quad (7)$$

is an estimator of the variance $\sigma^2 = E(X - \mu)^2$.

Sample Median

The median of X is the value m that satisfies $F(m) = 1/2$ (assuming there is a unique solution). Its empirical version is any value $\hat{m}_n$ such that $\hat{F}_n(\hat{m}_n)$ is equal or as close as possible to $1/2$. In the above example $F(x) = 1 - 1/x^2$, so that the theoretical median is $m = \sqrt{2} = 1.4142$. In the ordered sample, the 100th observation is equal to 1.4166, and the 101th observation is equal to 1.4191. A common choice for the sample median is taking the average of these two values. This gives $\hat{m}_n = 1.4179$.

Parametric Models

The distribution P may be partly known beforehand. The unknown parts of P are called *parameters* of the model. For example, if the X_i are the results of pass/fail inspection on products (the binary case), we know that P allows only two possibilities, say 1 and 0 (pass = 1, fail = 0). There is only one parameter, the probability that the product passes $\theta = P(X = 1)$. Generally, in a parametric model, it is assumed that P is known up to a finite number of parameters $\theta = (\theta_1, \dots, \theta_d)$. We then write $P = P_\theta$. When there are infinitely many parameters (e.g., the case when P is completely unknown), the model is called *nonparametric*.

If $P = P_\theta$ is a parametric model, one can often apply the maximum-likelihood procedure to estimate θ (see **Maximum Likelihood**).

Example 1 The time one stands in line for a certain service is often modeled as exponentially distributed (see **Probability Density Functions**). The random variable X representing the waiting time then has a density of the form

$$f_\theta(x) = \theta e^{-\theta x}, \quad x > 0 \quad (8)$$

where the parameter θ is the so-called *intensity* (a large value of θ means that – on average – the waiting time is short), and the maximum-likelihood estimator of θ is

$$\hat{\theta}_n = \frac{1}{\bar{X}_n} \quad (9)$$

Example 2 In many cases, one assumes that X is normally distributed (see **Probability Density Functions; Normal Distribution**). In that case there are two parameters: the mean $\theta_1 = \mu$ and the variance $\theta_2 = \sigma^2$. The maximum-likelihood estimators of (μ, σ^2) are $(\hat{\mu}_n, \hat{\sigma}_n^2)$, where $\hat{\mu}_n = \bar{X}_n$ is the sample mean and $\hat{\sigma}_n^2 = \sum_{i=1}^n (X_i - \bar{X}_n)^2/n$.

The Method of Moments

Suppose that the parameter θ can be written as a given function of the **moments** $\mu_1, \mu_2, \dots$. The *method of moments estimator* replaces these moments by their empirical counterparts $\hat{\mu}_{n,1}, \hat{\mu}_{n,2}, \dots$.

Example 3 Vilfredo Pareto [1] noticed that the number of people whose income exceeds level x is often approximately proportional to $x^{-\theta}$, where θ is a parameter that differs from country to country. Therefore, as a model for the distribution of incomes, one may propose the Pareto density

$$f_\theta(x) = \frac{\theta}{x^{\theta+1}}, \quad x > 1 \quad (10)$$

When $\theta > 1$, one has $\theta = \mu/(\mu - 1)$. Hence, the method of moments estimator of θ is in this case $t_1(\hat{P}_n) = \bar{X}_n/(\bar{X}_n - 1)$. After some calculations, one finds that the maximum-likelihood estimator of θ is $t_2(\hat{P}_n) = (n / \sum_{i=1}^n \log X_i)$. Let us compare these on the simulated data in Figure 1. We generated in this simulation a sample from the Pareto distribution with $\theta = 2$. The sample average turns out to be $\bar{X}_n = 1.9669$, so that the methods of moments estimate is 2.0342. The maximum-likelihood estimate is 1.9790. Thus, the maximum-likelihood estimate is a little closer to the true θ than the methods of moments estimate.

Besides maximum likelihood and the method of moments, the method of least squares is a common estimation method.

Properties of Estimators

Let $T_n = T_n(X_1, \dots, X_n)$ be an estimator of the real-valued parameter θ . Then it is desirable that T_n is in some sense close to θ . A minimum requirement is that the estimator approaches θ as the sample size increases. This is called *consistency*. To be more precise, suppose the sample $X_1, \dots, X_n$ are the first n of an infinite sequence $X_1, X_2, \dots$ of independent copies of X . Then T_n is called *consistent* if (with probability one)

$$T_n \longrightarrow \theta \text{ as } n \longrightarrow \infty \quad (11)$$

Note that consistency of frequencies as estimators of probabilities, or means as estimators of expectations, follows from the (strong) law of **large numbers**.

The *bias* of an estimator T_n of θ is defined as its mean deviation from θ :

$$\text{bias}(T_n) = E(T_n) - \theta \quad (12)$$

We remark here that the distribution of $T_n = T_n(X_1, \dots, X_n)$ depends on P , and, hence, on θ . Therefore, the expectation $E(T_n)$ depends on θ as well. We indicate this by writing $E(T_n) = E_\theta(T_n)$. The estimator T_n is called *unbiased* if

$$E_\theta(T_n) = \theta \text{ for all possible values of } \theta \quad (13)$$

Example 4 Consider the estimators $S_n^2 = \sum_{i=1}^n (X_i - \bar{X})^2 / (n-1)$ and $\hat{\sigma}_n^2 = \sum_{i=1}^n (X_i - \bar{X})^2 / n$. Note that S_n^2 is larger than $\hat{\sigma}_n^2$, but that the difference is small when n is large. It can be shown that S_n^2 is an unbiased estimator of the variance $\sigma^2 = \text{var}(X)$. The estimator $\hat{\sigma}_n^2$ is biased: it underestimates the variance.

In many models, unbiased estimators do not exist. Moreover, it often heavily depends on the model under consideration, whether or not an estimator is unbiased. A weaker concept is *asymptotic unbiasedness* (see [2]).

The mean square error of T_n as estimator of θ is

$$MSE(T_n) = E(T_n - \theta)^2 \quad (14)$$

One may decompose the MSE as

$$MSE(T_n) = \text{bias}^2(T_n) + \text{var}(T_n) \quad (15)$$

where $\text{var}(T_n)$ is the variance of T_n (whose square root is named the standard error).

Bias, variance, and mean square error are often quite hard to compute, because they depend on the distribution of all n observations $X_1, \dots, X_n$. However, one may use certain approximations for large sample sizes n . Under regularity conditions, the maximum-likelihood estimator $\hat{\theta}_n$ of θ is asymptotically unbiased, with asymptotic variance $1/(nI(\theta))$, where $I(\theta)$ is the Fisher information in a single observation. Thus, maximum-likelihood estimators reach the minimum variance bound asymptotically.

Histograms

Our next aim is estimating the density $f(x)$ at a given point x . The density (see **Probability Density Function (PDF)**) is defined as the derivative of the distribution function F at x :

$$f(x) = \lim_{h \rightarrow 0} \frac{F(x+h) - F(x)}{h} = \lim_{h \rightarrow 0} \frac{P(x, x+h]}{h} \quad (16)$$

Here, $(x, x+h]$ is the interval with left endpoint x (not included) and right endpoint $x+h$ (included). Unfortunately, replacing P by $\hat{P}_n$ here does not work, as for h it is small enough, $\hat{P}_n(x, x+h]$ will be equal to zero. Therefore, instead of taking the limit as $h \rightarrow 0$, we fix h at a (small) positive value, called the *bandwidth*. The estimator of $f(x)$, thus, becomes

$$\hat{f}_n(x) = \frac{\hat{P}_n(x, x+h]}{h} = \frac{\text{number of } X_i \in (x, x+h]}{nh} \quad (17)$$

A plot of this estimator at points $x \in \{x_0, x_0+h, x_0+2h, \dots\}$ is called a *histogram* (see **Graphical Representation of Data**).

Example 3 continued Figure 2 shows the histogram, with bandwidth $h = 0.5$, for the sample of size $n = 200$ from the Pareto distribution with parameter $\theta = 2$. The solid line is the density of this distribution.

Minimum χ^2

Of course, for real (not simulated) data, the underlying distribution/density is not known. Let us explain in an example a procedure for checking whether certain model assumptions are reasonable. Suppose

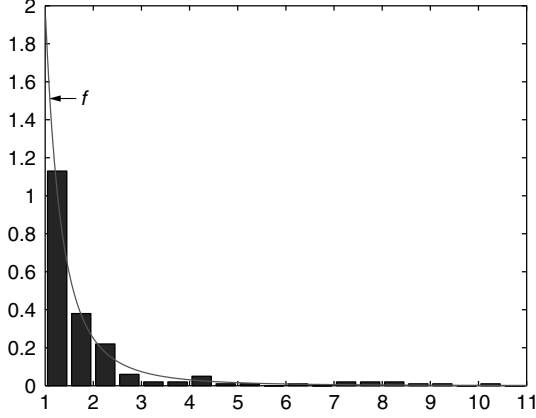


Figure 2 Histogram with bandwidth $h = 0.5$ and true density

that one wants to test whether data come from the exponential distribution with parameter θ equal to 1. We draw a histogram of the sample (sample size $n = 200$), with bandwidth $h = 1$ and 10 cells (see Figure 3). The cell counts are (151, 28, 4, 6, 1, 1, 4, 3, 1, 1). Thus, for example, the number of observations that falls in the first cell that is, those that have values between 0 and 1, is equal to 151. The cell probabilities are, therefore, (0.755, 0.140, 0.020, 0.030, 0.005, 0.005, 0.020, 0.015, 0.005, 0.005). Now, according to the exponential distribution, the probability that an observation falls in cell k is equal to $e^{-(k-1)} - e^{-k}$, for $k = 1, \dots, 10$. These cell probabilities are (0.6621, 0.2325, 0.0855, 0.0315, 0.0116, 0.0043, 0.0016, 0.0006, 0.0002, 0.0001). Because the probabilities of the last four cells are very small, we merge them together. This gives cell counts $(N_1, \dots, N_7) = (151, 28, 4, 6, 1, 1, 9)$ and cell probabilities $(\pi_1, \dots, \pi_7) = (0.6621, 0.2325, 0.0855, 0.0315, 0.0116, 0.0043, 0.0025)$. To check whether the cell frequencies differ significantly from the hypothesized cell probabilities, we calculate Pearson's χ^2 . It is defined as

$$\chi^2 = \frac{(N_1 - n\pi_1)^2}{n\pi_1} + \dots + \frac{(N_7 - n\pi_7)^2}{n\pi_7} \quad (18)$$

We write this as $\chi^2 = \chi^2(\text{exponential})$ to stress that the cell probabilities were calculated assuming the exponential distribution. Now, if the data are exponentially distributed, the χ^2 statistic is generally not too large. But what is large? Consulting a table of Pearson's χ^2 at the 5% **significance**

level gives the critical value $c = 12.59$. Here we use six **degrees of freedom**. This is because there are $m = 7$ cell probabilities, and there is the restriction $\pi_1 + \dots + \pi_m = 1$, so we estimated $m - 1 = 6$ parameters. After some calculations, one obtains $\chi^2(\text{exponential}) = 168.86$. This exceeds the critical value c , that is, $\chi^2(\text{exponential})$ is too large to support the assumption of the exponential distribution. In fact, the data considered here are the simulated sample from the Pareto distribution with parameter $\theta = 2$. We shifted this sample one unit to the left. The value of χ^2 for this (shifted) Pareto distribution is

$$\chi^2(\text{Pareto}) = 10.81 \quad (19)$$

This is below the critical value c , so that the test, indeed, does not reject the Pareto distribution. However, this comparison is not completely fair, as our decision to merge the last four cells was based on the exponential distribution, which has much lighter tails than the Pareto distribution.

In Figure 3, the histogram is shown, together with the densities of the exponential and Pareto distribution. Indeed, the Pareto distribution fits the data better in the sense that it puts more mass at small values.

Continuing with the test for the exponential distribution, we note that, in many situations, the intensity θ is not required to be fixed beforehand. One may use an estimator for θ and proceed as before, calculating χ^2 with the estimated value for θ . However, the critical values of the test then become smaller.

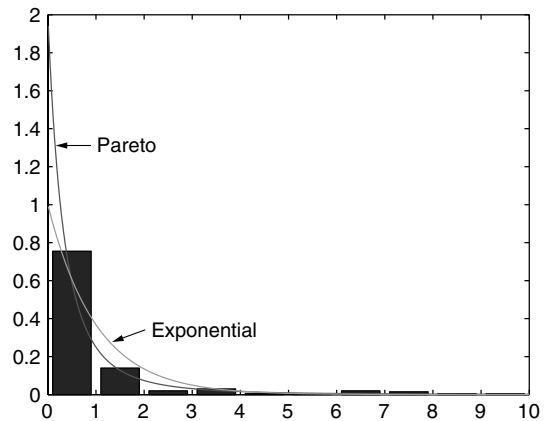


Figure 3 Histogram with bandwidth $h = 1$, exponential and Pareto density

This is because, clearly, estimating parameters using the sample means that the hypothesized distribution is pulled toward the sample. Moreover, when using, for example, maximum-likelihood estimators of the parameters, critical values will in fact be hard to compute. The minimum χ^2 estimator overcomes this problem. Let $\pi_k(\vartheta)$ denote the cell probabilities when the parameter value is ϑ , that is, in the exponential case $\pi_k(\vartheta) = e^{-\vartheta(k-1)} - e^{-\vartheta k}$, $k = 1, \dots, m-1$, and $\pi_m(\vartheta) = 1 - \sum_{k=1}^{m-1} \pi_k(\vartheta)$. The minimum χ^2 estimator $\hat{\theta}_n$ is now the minimizer over ϑ of

$$\left\{ \frac{(N_1 - n\pi_1(\vartheta))^2}{n\pi_1(\vartheta)} + \dots + \frac{(N_m - n\pi_m(\vartheta))^2}{n\pi_m(\vartheta)} \right\} \quad (20)$$

The χ^2 test with this estimator for θ now has $m-2$ degrees of freedom. Generally, the number of degrees of freedom is $m-1-d$, where d is the number of estimated free parameters when calculating cell probabilities. The critical values of the test can be found in a χ^2 table.

Sufficiency

A goal of statistical analysis is generally to summarize the (large) data set into a small number of characteristics. The sample mean and sample variance are such summarizing statistics, but so is, for example, the sample median, and so on. The question arises, to what extent one can summarize data without throwing away information. For example, suppose you are given the empirical distribution function $\hat{F}_n$, and you are asked to reconstruct the original data $X_1, \dots, X_n$. This is not possible since the ordering of the data is lost. However, the index i of X_i is just a label: it contains no information about the distribution P of X_i (assuming that each observation X_i comes from the same distribution, and the observations are independent). We say that the empirical distribution $\hat{F}_n$ is *sufficient*. Generally, a statistic

$T_n = T_n(X_1, \dots, X_n)$ is called *sufficient* for P if the distribution of the data given the value of T_n does not depend on P . For example, it can be shown that when P is the exponential distribution with unknown intensity, then the sample mean is sufficient. When P is the **normal distribution** with unknown mean and variance, then the sample mean and sample variance are sufficient. Cell counts are not sufficient when, for example, P is a continuous distribution. This is because, if one only considers the cell counts, one throws away information on the distribution within a cell. Indeed, when one compares Figures 2 and 3 (recall that in Figure 3 we shifted the sample one unit to the left), one sees that, by using just 10 cells instead of 20, the strong decrease in the second half of the first cell is no longer visible.

Sufficiency depends very heavily on the model for P . Clearly, when one decides to ignore information because of a sufficiency argument, one may be ignoring evidence that the model's assumptions may be wrong. Sufficiency arguments should be treated with caution.

References

- [1] Pareto, V. (1897). *Course d'Économie Politique*, Rouge, Lausanne et Paris.
- [2] Bickel, P.J. & Doksum, K.A. (2001). *Mathematical Statistics*, 2nd Edition, Holden-Day, San Francisco.

Related Article

Bias of an Estimator.

SARA A. VAN DE GEER

Article originally published in Encyclopedia of Statistics in Behavioral Science (2005, John Wiley & Sons, Ltd). Minor revisions for this publication by Jeroen de Mast

Expectation

The expectation or expected value of a random variable is also referred to as its *mean* or *average value*. The terms *expectation*, *expected value*, and *mean* can be used interchangeably. The expectation can be thought of as a measure of the “center” or “location” of the probability distribution of the random variable. Two other measures of the center or location are the median and the mode. If the random variable is denoted by the letter X , the expectation of X is usually denoted by $E(X)$, read “ E of X ”.

The form of the definition of the expectation of a random variable is slightly different depending on the nature of the probability distribution of the random variable (see **Probability Density Functions**). The expectation of a discrete random variable is defined as the sum of each value of the random variable weighted by the probability that the random variable is equal to that value. Thus,

$$E(X) = \sum_x xf(x) \quad (1)$$

where $f(x)$ denotes the probability distribution of the discrete random variable, X . So for example, suppose that the random variable X takes on five discrete values, -2 , -1 , 0 , 1 , and 2 . Also suppose that the probability that the random variable takes on the value -2 (i.e., $X = -2$) is equal to 0.1 , and that $P(X = -1) = 0.2$, $P(X = 0) = 0.3$, $P(X = 1) = 0.25$, and $P(X = 2) = 0.15$. Then,

$$E(X) = (-2 \times 0.1) + (-1 \times 0.2) + (0 \times 0.3) + (1 \times 0.25) + (2 \times 0.15) = 0.15 \quad (2)$$

Note that the individual probabilities sum to 1 . The expectation is meant to summarize a “typical” observation of the random variable but, as can be seen from the example above, the expectation of a random variable is not necessarily equal to one of the values of that random variable. It can also be very sensitive to small changes in the probability assigned to a large value of X .

It is possible that the expectation of a random variable does not exist. In the case of a discrete random variable, this occurs when the sum defining the expectation does not converge or is not equal to infinity. Consider the *St. Petersburg paradox* discovered by

the Swiss eighteenth-century mathematician, Daniel Bernoulli [1]. A fair coin is tossed repeatedly until it shows tails. The total number of tosses determines the prize, which equals $\$2^n$, where n is the number of tosses. The expected value of the prize for each toss is $\$1$ ($2 \times 0.5 + 0 \times 0.5$). The expected value of the prize for the game is the sum of the expected values for each toss. As there are an infinite number of tosses, the expected value of the prize for the game is an infinite number of dollars.

The expectation of a continuous random variable is defined as the integral of the individual values of the random variable weighted according to their probability distribution. Thus,

$$E(X) = \int_{-\infty}^{\infty} xf(x) dx \quad (3)$$

where $f(x)$ is the probability distribution (see **Probability Density Function (PDF)**) of the continuous random variable, X . This achieves the same end for a continuous random variable as equation (1) did for a discrete random variable; it sums each value of the random variable weighted by the probability that the random variable is equal to that value.

Again it is possible that the expectation of the continuous random variable does not exist. This occurs when the integral does not converge or equals infinity. A classic example of a random variable whose expected value does not exist is a Cauchy random variable.

An expectation can also be calculated for the function of a random variable. This is done by applying equation (1) or (3) to the distribution of a function of the random variable. For example, a function of the random variable X is $(X - \mu_X)^2$, where μ_X denotes the mean of the random variable X . The expectation of this function is referred to as the *variance* (see **Variance**) of X . The mean and variance are related to the first and second *moment* (see **Moments**).

Expectations have a number of useful properties.

1. The expectation of a constant is equal to that constant.
2. The expectation of a constant multiplied by a random variable is equal to the constant multiplied by the expectation of the random variable.
3. The expectation of the sum of a group of random variables is equal to the sum of their individual expectations.

2 Expectation

4. The expectation of the product of a group of random variables is equal to the product of their individual expectations if the random variables are independent.

More information on the topic of expectation is given in [2–4].

- [2] Casella, G. & Berger, R.L. (1990). *Statistical Inference*, Duxbury Press, Belmont.
- [3] DeGroot, M.H. (1986). *Probability and Statistics*, 2nd Edition, Addison-Wesley, Reading.
- [4] Mood, A.M., Graybill, F.A. & Boes, D.C. (1974). *Introduction to the Theory of Statistics*, McGraw-Hill, Singapore.

REBECCA WALWYN

References

- [1] Bernoulli, D. (1738). Exposition of a new theory on the measurement of risk, *Econometrica* **22**, 23–36.

Exploratory Data Analysis

For most statisticians in the late 1960s and early 1970s, the really hot topics in their subject were definitely not simple **descriptive statistics** or basic plots. Instead, the fireworks came from the continuing debates over the nature of statistical inference and the foundations of probability, followed closely by the modeling of increasingly complex experimental designs within a general *analysis of variance* setting, an area that was boosted by the growing availability of cheap and powerful computing facilities. However, in an operation reminiscent of the *samizdat* in the latter days of the Soviet Union (which secretly circulated the banned novels of Aleksandr Solzhenitsyn, among other things), mimeographed copies of a revolutionary recasting of statistics began their hand-to-hand journey around a selected number of universities from early 1970 onward.

This heavily thumbed, three-volume text was not the work of a dissident junior researcher impatient with the conservatism of his/her elders, but was by one of America's most distinguished mathematicians, John Tukey. What Tukey did was to reestablish the relationship between these neglected descriptive measures (and plots) and many of the currently active areas in statistics, including computing, modeling, and inference. In doing so, he also resuscitated interest in the development of new summary measures and displays, and in the almost moribund topic of regression. The central notion behind what he termed *exploratory data analysis*, or EDA, is simple but fruitful: knowledge about the world comes from many wells and by many routes. One legitimate and productive way is to treat data not as being merely supportive of already existing knowledge, but primarily as a source of new ideas and hypotheses. For Tukey, like a detective looking for the criminal, the world was full of potentially valuable material that could yield new insights into what might be the case. EDA is the active reworking of this material with the aim of extracting new meanings from it: "The greatest value of a picture is when it *forces* us to notice what we never expected to see" [6, p. vi]. In addition, while EDA has very little to say about how data might be gathered (although there is a subdivision of the movement whose concentration on model evaluation could be said to have *indirect* implications for **data**

collection), the movement is very concerned that the tools to extract meaning from data do so with minimum ambiguity. What this indicates, using a slightly more technical term, is that EDA is concerned with *robustness*, that is, with the study of all the factors that disturb or distort conclusions drawn from data, and how to draw their teeth (see also [1] for an introduction to EDA, which nicely mixes the elementary with the sophisticated aspects of the topic).

A simple example here is the use of the median *versus* the arithmetic mean as a measure of *location*, the former is said to be robust and resistant to discrepant values or *outliers* because it is based only on the *position* of the middle number in the (ordered) sample, while the latter uses all the numbers and their values, however unrepresentative, in its calculation. Further, a *trimmed mean*, calculated from an *ordered* sample where, say, 10% of the top and bottom numbers have been removed, is also more robust and resistant than the raw arithmetic mean based on all the readings, since trimming automatically removes any outliers, which by definition, lie in the upper and lower tails of such a sample, and are few in number. In general, robust-resistant measures are preferable to nonrobust ones because, while both give pretty much the same answer in samples where there are no outliers, the former yields a less distorted measure in samples where outliers are present. Tukey and some of his colleagues [2] set themselves the task of investigating over 35 separate measures of location for their robustness. Although the median and the percentage trimmed mean came out well, Tukey also promoted the computer-dependent *biweight* as his preferred measure of location. Comparable arguments can also be made for choosing a robust measure of *spread* (or *scale*) over a nonrobust one, with the *midspread* (or interquartile range), that is, the middle 50% of a sample (or the difference between the upper and lower quartiles (see **Quartiles**) or *hinges*, to use Tukey's term), being more acceptable than the variance or the standard deviation. Since the experienced data explorer is unlikely to know beforehand what to expect with a given *batch* (Tukey's neologism for a *sample*), descriptive tools that are able to cope with messy data sets are preferable to the usual, but more outlier-prone, alternatives.

Plots that use robust rather than outlier-sensitive measures have also been part of the EDA doctrine; thus the ubiquitous box-and-whisker plot is based on both the median and the midspread, while the latter

measure helps determine the whisker length, which can then be useful in identifying potential outliers. Similar novelties such as the stem and leaf plot are also designed for the median rather than the mean to be read off easily, as can the upper and lower quartiles, and hence the midspread. Rather more exotic birds such as *hanging or suspended rootograms* and *half-normal plots* (all used for checking the Gaussian nature of the data) were either Tukey's own creation or that of workers influenced by him, for instance, Daniel and Wood [3], whose classic study of data modeling owed much to EDA. Tukey did not neglect earlier and simpler displays such as scatterplots (*see Graphical Representation of Data*), although, true to his radical program, these became transformed into windows to data shape and symmetry, robustness, and even robust differences of location and spread. A subset of these displays, residual and leverage plots for instance, also played a key role in the revival of interest in regression, particularly in the area termed *regression diagnostics*. In addition, while EDA has concentrated on relatively simple data sets and data structures, there are less well known but still provocative incursions into the partitioning of multiway and multifactor tables [4], including Tukey's rethinking of the analysis of variance (*see* [5]).

While EDA offered a flexibility and adaptability to knowledge building that was missing from the more formal processes of statistical inference, this was achieved with a loss of certainty about the knowledge gained – provided, of course, that one believes that more formal methods do generate truth, or your money back. Tukey was less sure on this latter point in that while formal methods of inference are bolted on to EDA in the form of confirmatory data analysis, or CDA (which favored the Bayesian route – *see* Mosteller and Tukey's early account of EDA and CDA in [6]), he never expended as much effort on CDA as he did on EDA, although he often argued that both should run in harness, with one complementing the other. Thus, for Tukey (and EDA), truth gained by such an epistemologically *inductive* method was always going to be partial, local, and relative, and liable to fundamental challenge and change. On the other hand, without taking such risks, nothing new could emerge. What seemed to be on offer, therefore, was an *entrepreneurial* view of statistics, and *science*, as a permanent revolution.

Although EDA broke cover more than 25 years ago with the simultaneous publication of the two

texts on EDA and regression [7, 8], it is still too early to say what has been the real impact of EDA on data analysis. On the one hand, many of EDA's graphical novelties have been incorporated into most modern texts and computer programs in statistics, as have the raft of robust measures on offer. On the other, the risky approach to statistics adopted by EDA has been less popular, with deduction-based inference taken to be the only way to discover the world. However, while one could point to underlying and often contradictory cultural movements in scientific belief and practice to account for the less than wholehearted embrace of Tukey's ideas, the future of statistics lies with EDA.

References

- [1] Velleman, P.F. & Hoaglin, D.C. (1981). *Applications, Basics and Computing of Exploratory Data Analysis*, Duxbury Press, Boston.
- [2] Andrews, D.F., Bickel, P.J., Hampel, F.R., Huber, P.J., Rogers, W.H. & Tukey, J.W. (1972). *Robust Estimates of Location: Survey and Advances*, Princeton University Press, Princeton.
- [3] Daniel, C. & Wood, F.S. (1980). *Fitting Equations to Data*, 2nd Edition, John Wiley & Sons, New York.
- [4] Hoaglin, D.C., Mosteller, F. & Tukey, J.W. (eds) (1985). *Exploring Data Tables, Trends and Shapes*, John Wiley & Sons, New York.
- [5] Hoaglin, D.C., Mosteller, F. & Tukey, J.W. (1992). *Fundamentals of Exploratory Analysis of Variance*, John Wiley & Sons, New York.
- [6] Mosteller, F. & Tukey, J.W. (1968). Data analysis including statistics, in *Handbook of Social Psychology*, G. Lindzey & E. Aronson, eds, Addison-Wesley, Reading.
- [7] Tukey, J.W. (1977). *Exploratory Data Analysis*, Addison-Wesley, Reading.
- [8] Mosteller, F. & Tukey, J.W. (1977). *Data Analysis and Regression: A Second Course in Statistics*, Addison-Wesley, Reading.

SANDY LOVIE

Article originally published in Encyclopedia of Statistics in Behavioral Science (2005, John Wiley & Sons, Ltd). Minor revisions for this publication by Jeroen de Mast.

Heteroscedasticity

Introduction

Homoscedasticity is the condition in which a random variable or its observed values have the same degree of variation for all sampling units in the mathematical model or data set. By “constant variation” we usually mean a common variance for all sampling units. When that condition is not met, the random variable or data are called *heteroscedastic*. The origin of these terms is from the Greek word $\sigma\kappa\epsilon\delta\acute{\alpha}\nu\nu\mu\iota$, “to scatter or disperse”, and the Greek words for “same” and “different”.

Linear Regression Analysis

Homoscedasticity is an essential assumption in the linear regression model and its fit by *least squares*. For the simplest case of a single dependent variable Y and one independent variable X , the model is

$$Y_i = \alpha + \beta X_i + e_i, \quad i = 1, \dots, N \quad (1)$$

The e_i are random variables with $E(e_i) = 0$, and if the homoscedasticity condition holds, $\text{var}(e_1) = \dots = \text{var}(e_N) = \sigma^2$. If the random disturbance terms are heteroscedastic and have different variances $\sigma_i^2 = \text{var}(e_i)$, the ordinary least-squares estimators of α and β , while still unbiased, no longer have the Gauss–Markov property of minimum variance. If the individual variances are known up to a proportionality constant, weighted least-squares fit of the linear model can be obtained by using the scaled data

$$\frac{Y_i}{\sigma_i} = \frac{\alpha}{\sigma_i} + \frac{\beta X_i}{\sigma_i} + \frac{e_i}{\sigma_i} \quad (2)$$

The usual ordinary least-squares **estimation** process is applied to the scaled values, with the intercept also scaled by the known standard deviations σ_i .

The weighted least-squares estimators are easily displayed by the matrix form of the multiple regression model with an intercept and p independent variables:

$$\mathbf{Y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{e} \quad (3)$$

$\mathbf{Y}$ is the $N \times 1$ vector of observations on the dependent variable. $\mathbf{X}$ is the $N \times (p + 1)$ matrix of predictor variable values of full rank $p + 1$, with a

first column of ones for the intercept term. $\boldsymbol{\beta}$ is the $(p + 1) \times 1$ parameter vector with the intercept α in the first position. The $N \times 1$ vector of random disturbances $\mathbf{e}$ has $E(\mathbf{e}) = \mathbf{0}$ and diagonal **covariance** matrix

$$\boldsymbol{\Sigma} = \text{cov}(\mathbf{e}, \mathbf{e}') = \begin{bmatrix} \sigma_1^2 & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & \sigma_N^2 \end{bmatrix} \quad (4)$$

The weighted least-squares estimator of $\boldsymbol{\beta}$ is

$$\hat{\boldsymbol{\beta}}_w = (\mathbf{X}'\boldsymbol{\Sigma}^{-1}\mathbf{X})^{-1}\mathbf{X}'\boldsymbol{\Sigma}^{-1}\mathbf{Y} \quad (5)$$

where $\boldsymbol{\Sigma}^{-1}$ is the $N \times N$ diagonal matrix, the diagonal elements of which are the reciprocals $1/\sigma_i^2$ of the corresponding elements in $\boldsymbol{\Sigma}$. Unfortunately, since the variances of the disturbance terms are rarely known, these results are mainly only of theoretical interest.

Two-Sample t Test

The hypothesis that two independent **normal distribution** means are equal arises frequently in practical data analysis. The Student–Fisher t test of that hypothesis requires that the two populations have a common unknown variance (*see Parametric Tests*). That assumption is easily tested by the ratio $F = s_1^2/s_2^2$ of the independent sample variances (the F test; *see Parametric Tests*). The standard procedure of the t test assumes that $\sigma_1^2 = \sigma_2^2$. Hsu [1] calculated the true type I error probabilities for selected ratios of the population variances $\theta = \sigma_1^2/\sigma_2^2$ and sample sizes $R = N_1/N_2$ for a 0.05 level test (*see Power and Significance Level*). If θ is less than 1 and R is larger than 1, the true type I error rate will be larger than the nominal 0.05 value. If θ and R are both larger than 1, the true type I probability will be less than 0.05. These and other properties of the t test in the presence of unequal variances have been described by Scheffé [2].

The test of the equality of two means of independent normal populations is called the *Behrens–Fisher problem*, after its original investigators. Welch [3, 4] has proposed an alternative test with an approximate t distribution when the population variances are unequal.

The Analysis of Variance

Box [5, 6] has investigated the effect of heteroscedastic disturbance terms on the one- and two-way *analysis of variance* (see **Analysis of Variance**) type I error rates. As in the two-sample *t* test, the effect of unequal variances in the one-way layout is exacerbated by unequal sample sizes. When the sample sizes are equal, the true type I error rates are only slightly higher than the nominal 0.05 value. If the treatments with the smaller variances have larger sample sizes, the type I error rate may be much greater than 0.05. If the opposite condition holds, or if the ratios of the variances roughly follow those of the sample sizes, the true type I error rate may be smaller than the nominal 0.05. For the two-way layout with a single observation in each cell, Box showed that unequal column variances led to row test type I error rates slightly below the nominal 0.05 value, and to column test type I error rates slightly above 0.05. Further remarks on heteroscedasticity in the analysis of variance have been given by Scheffé [2].

References

- [1] Hsu, P.L. (1938). Contribution to the theory of student's *t* test as applied to the problem of two samples, *Statistical Research Memories* **2**, 1–24.

- [2] Scheffé, H. (1959). *The Analysis of Variance*, John Wiley & Sons, New York.
- [3] Welch, B.L. (1937). The significance of the difference between two means when the population variances are unequal, *Biometrika* **29**, 350–362.
- [4] Welch, B.L. (1947). The generalization of “Student's” problem when several population variances are involved, *Biometrika* **34**, 28–35.
- [5] Box, G.E.P. (1954). Some theorems on quadratic forms applied in the study of analysis of variance problems, I. effect of inequality of variance in the one-way classification, *Annals of Mathematical Statistics* **25**, 290–302.
- [6] Box, G.E.P. (1954). Some theorems on quadratic forms applied in the study of analysis of variance problems, II. Effects of inequality of variance and of correlation between errors in the two-way classification, *Annals of Mathematical Statistics* **25**, 484–498.

Related Articles

Homogeneity of Variances.

DONALD F. MORRISON

Article originally published in Encyclopedia of Statistics in Behavioral Science (2005, John Wiley & Sons, Ltd). Minor revisions for this publication by Jeroen de Mast.

Laws of Large Numbers

Introduction

Laws of large numbers concern the behavior of the average of n random variables, $S_n = 1/n \sum_{i=1}^n X_i$, as n grows large. This topic has an interesting early history dating back to the work of Bernoulli and Poisson in the eighteenth and nineteenth centuries; see [1] and [2] for details. For a fairly straightforward modern introduction, see [3]. There are two types of laws, namely, *weak* and *strong*, depending on the mode of convergence of the average S_n to its limit: *weak* refers to convergence in probability and *strong* denotes convergence with probability 1. The difference between these modes is perhaps best appreciated by those with a solid grounding in measure theory; see [4] or [5], for example. There are several versions of the law depending on assumptions regarding (a) the independence of the X_i and (b) whether the X_i follow the same probability distribution. Here, we consider the simplest case in which the X_i are assumed to be mutually independent and follow the same probability distribution with a finite mean μ [5, 6]. More general versions exist in which either, or both, of the assumptions (a) and (b) are relaxed, but then, other conditions are imposed. For example, if the random variables have different means and different variances, but certain moment conditions hold, then the Kolmogorov's strong law is obtained [5]; if the random variables form a stationary stochastic process, then a strong law is available (the "ergodic" theorem [7]); if both assumptions (a) and (b) are weakened while the random variables have a martingale structure, and a moment condition is satisfied, then a strong law is available [5, 8].

The weak law states that S_n converges in probability to μ as $n \rightarrow \infty$; put more precisely, this law states that for any small positive number ϵ , $\Pr\{|S_n - \mu| < \epsilon\} \rightarrow 1$ as $n \rightarrow \infty$. The strong law states that S_n converges to μ with probability 1 as $n \rightarrow \infty$; put more precisely, this law states that for any small positive number ϵ , $\lim_{n \rightarrow \infty} \Pr\{|S_n - \mu| < \epsilon\} = 1$. We now consider three applications of the law.

Application 1 The frequentist concept of probability is based on the notion of the stabilization

of relative frequency. Suppose that n independent experiments are conducted under identical conditions, and suppose that, in each experiment, the event A either occurs, with probability π , or does not occur, with probability $1 - \pi$. In the i th experiment, let X_i take the value 1 if A occurs, and 0 if A does not occur ($i = 1, \dots, n$). Then S_n , the proportion of times that the event A occurs, is the relative frequency of the event A in the first n experiments and by the strong law, as $n \rightarrow \infty$, S_n converges with probability 1 to π , the true probability of the event A .

Application 2 Suppose that it is of interest to estimate the mean and variance of a quality characteristic of a production process, and that measurements are available from a random sample of n products from the process. Let X_i be a random variable representing the quality characteristic for the i th product and suppose that the X_i have a finite mean μ and variance σ^2 ($i = 1, \dots, n$). Consider the usual unbiased estimators S_n of μ and $T_n = n/(n-1)\{1/n \sum_{i=1}^n X_i^2 - \bar{X}^2\}$ of σ^2 . Then, from the strong law, it follows that S_n converges with probability 1 to the population mean μ , that $1/n \sum_{i=1}^n X_i^2$ converges with probability 1 to the expectation of the X_i^2 (which is $\mu^2 + \sigma^2$), and, therefore, that T_n converges with probability 1 to the population variance σ^2 . These results provide theoretical support for the strong consistency of statistical estimation for the sample estimators S_n and T_n (see **Estimation**).

References

- [1] Hacking, I. (1990). *The Taming of Chance*, Cambridge University Press, Cambridge.
- [2] Stigler, S.M. (1986). *The History of Statistics: The Measurement of Uncertainty Before 1900*, The Belknap Press of Harvard University Press, Cambridge.
- [3] Isaac, R. (1995). *The Pleasures of Probability*, Springer-Verlag, New York.
- [4] Chung, K.L. (2001). *A Course in Probability Theory*, 3rd Edition, Academic Press, San Diego.
- [5] Sen, P.K. & Singer, J.M. (1993). *Large Sample Methods in Statistics: An Introduction with Applications*, Chapman & Hall, London.
- [6] Feller, W. (1968). *An Introduction to Probability and its Applications*, 3rd Edition, John Wiley & Sons, New York, Vol. I.
- [7] Breiman, L. (1968). *Probability*, Addison-Wesley, Reading.

2 **Laws of Large Numbers**

- [8] Feller, W. (1971). *An Introduction to Probability and Its Applications*, 2nd Edition, John Wiley & Sons, New York, Vol. II.

Article originally published in Encyclopedia of Statistics in Behavioral Science (2005, John Wiley & Sons, Ltd). Minor revisions for this publication by Jeroen de Mast.

JIM W. KAY

Least-Squares Estimation

Introduction

We present an explicit expression for the least-squares estimator $\hat{\beta}$ of β in the linear regression model, and also give its **covariance** matrix. The calculations are illustrated with a numerical example. The F statistic for testing linear hypotheses on β is also given. We conclude with some extensions.

The method of least-squares estimates parameters by minimizing the squared discrepancies between observed data, on the one hand, and their expected values on the other. We will study the method in the context of a regression problem, where the variation in one variable, called the *response variable*, Y , can be partly explained by the variation in the other variables, called *covariables*, X . For example, variation in exam results, Y , are mainly caused by variation in abilities and diligence, X , of the students, or variation in survival times, Y , are primarily due to variations in environmental conditions X . Given the value of X , the best prediction of Y (in terms of mean square error – see **Estimation**) is the mean $f(X)$ of Y given X . We say that Y is a function of X plus noise:

$$Y = f(X) + \text{noise} \quad (1)$$

The function f is called a *regression function*. It is to be estimated from n sampled covariables and their responses $(x_1, y_1), \dots, (x_n, y_n)$.

Suppose f is known up to a finite number $p \leq n$ of parameters $\beta = (\beta_1, \dots, \beta_p)'$, that is, $f = f_\beta$. We estimate β by the value $\hat{\beta}$ that gives the best fit to the data. The least-squares estimator, denoted by $\hat{\beta}$, is that value of b that minimizes

$$\sum_{i=1}^n (y_i - f_b(x_i))^2 \quad (2)$$

over all possible values of b .

The least-squares criterion is a computationally convenient measure of fit. It corresponds to *maximum-likelihood estimation* when the noise is normally distributed with equal variances. Other measures of fit are sometimes used, for example least-absolute deviations, which is more robust against **outliers**.

Linear Regression

Consider the case where f_β is a linear function of β , that is,

$$f_\beta(\mathbf{X}) = X_1\beta_1 + \dots + X_p\beta_p \quad (3)$$

Here $(X_1, \dots, X_p)$ stand for the observed variables used in $f_\beta(X)$.

To write down the least-squares estimator for the linear regression model, it will be convenient to use matrix notation. Let $\mathbf{y} = (y_1, \dots, y_n)'$ and let $\mathbf{X}$ be the $n \times p$ data matrix of the n observations on the p variables

$$\mathbf{X} = \begin{pmatrix} \mathbf{x}_{1,1} & \dots & \mathbf{x}_{1,p} \\ \vdots & \dots & \vdots \\ \mathbf{x}_{n,1} & \dots & \mathbf{x}_{n,p} \end{pmatrix} = (\mathbf{x}_1 \dots \mathbf{x}_p) \quad (4)$$

where $\mathbf{x}_j$ is the column vector containing the n observations on variable j , $j = 1, \dots, p$. Denote the squared length of an n -dimensional vector $\mathbf{v}$ by $\|\mathbf{v}\|^2 = \mathbf{v}'\mathbf{v} = \sum_{i=1}^n v_i^2$. Then equation (2) can be written as

$$\|\mathbf{y} - \mathbf{X}\mathbf{b}\|^2 \quad (5)$$

which is the squared distance between the vector $\mathbf{y}$ and the linear combination $\mathbf{b}$ of the columns of the matrix $\mathbf{X}$. The distance is minimized by taking the *projection* of $\mathbf{y}$ on the space spanned by the columns of $\mathbf{X}$ (see Figure 1).

Suppose now that $\mathbf{X}$ has full-column rank, that is, no column in $\mathbf{X}$ can be written as a linear

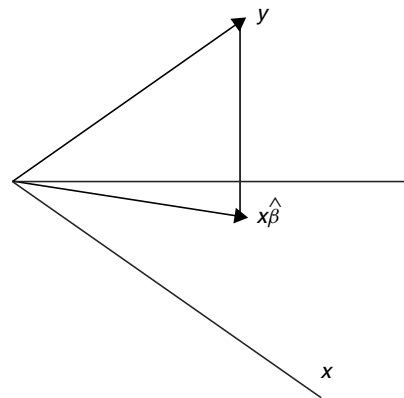


Figure 1 The projection of the vector $\mathbf{y}$ on the plane spanned by $\mathbf{x}$

2 Least-Squares Estimation

combination of the other columns. Then, the least-squares estimator $\hat{\beta}$ is given by

$$\hat{\beta} = (\mathbf{X}'\mathbf{X})^{-1} \mathbf{X}'\mathbf{y} \quad (6)$$

The Variance of the Least-Squares Estimator

To construct **confidence intervals** for the components of $\hat{\beta}$, or linear combinations of these components, one needs an estimator of the covariance matrix of $\hat{\beta}$. Now, it can be shown that, given $\mathbf{X}$, the covariance matrix of the estimator $\hat{\beta}$ is equal to

$$(\mathbf{X}'\mathbf{X})^{-1} \sigma^2 \quad (7)$$

where σ^2 is the variance of the noise. As an estimator of σ^2 , we take

$$\hat{\sigma}^2 = \frac{1}{n-p} \|\mathbf{y} - \mathbf{X}\hat{\beta}\|^2 = \frac{1}{n-p} \sum_{i=1}^n \hat{e}_i^2 \quad (8)$$

where the $\hat{e}_i$ are the **residuals**

$$\hat{e}_i = y_i - \mathbf{x}_{i,1}\hat{\beta}_1 - \cdots - \mathbf{x}_{i,p}\hat{\beta}_p \quad (9)$$

The covariance matrix of $\hat{\beta}$ can, therefore, be estimated by

$$(\mathbf{X}'\mathbf{X})^{-1} \hat{\sigma}^2 \quad (10)$$

For example, the estimate of the variance of $\hat{\beta}_j$ is

$$\text{var}(\hat{\beta}_j) = \tau_j^2 \hat{\sigma}^2 \quad (11)$$

where τ_j^2 is the j th element on the diagonal of $(\mathbf{X}'\mathbf{X})^{-1}$. A confidence interval for β_j is now obtained by taking the least-squares estimator $\hat{\beta}_j \pm$ a margin:

$$\hat{\beta}_j \pm c \sqrt{\text{var}(\hat{\beta}_j)} \quad (12)$$

where c depends on the chosen confidence level. For a 95% confidence interval, the value $c = 1.96$ is a good approximation when n is large. For smaller values of n , one usually takes a more conservative c using the tables for the student distribution with $n - p$ degrees of freedom.

Numerical Example

Consider a regression with constant, linear, and quadratic terms:

$$f_{\beta}(X) = \beta_1 + X\beta_2 + X^2\beta_3 \quad (13)$$

We take $n = 100$ and $x_i = i/n$, $i = 1, \dots, n$. The matrix $\mathbf{X}$ is now

$$\mathbf{X} = \begin{pmatrix} 1 & x_1 & x_1^2 \\ \vdots & \vdots & \vdots \\ 1 & x_n & x_n^2 \end{pmatrix} \quad (14)$$

This gives

$$\mathbf{X}'\mathbf{X} = \begin{pmatrix} 100 & 50.5 & 33.8350 \\ 50.5 & 33.8350 & 25.5025 \\ 33.8350 & 25.5025 & 20.5033 \end{pmatrix} \quad (15)$$

$$(\mathbf{X}'\mathbf{X})^{-1} = \begin{pmatrix} 0.0937 & -0.3729 & 0.3092 \\ -0.3729 & 1.9571 & -1.8189 \\ 0.3092 & -1.8189 & 1.8009 \end{pmatrix} \quad (16)$$

We simulated n independent standard normal random variables $e_1, \dots, e_n$, and calculated for $i = 1, \dots, n$,

$$y_i = 1 - 3x_i + e_i \quad (17)$$

Thus, in this example, the parameters are

$$\begin{pmatrix} \beta_1 \\ \beta_2 \\ \beta_3 \end{pmatrix} = \begin{pmatrix} 1 \\ -3 \\ 0 \end{pmatrix} \quad (18)$$

Moreover, $\sigma^2 = 1$. Because this is a simulation, these values are known.

To calculate the least-squares estimator, we need the values of $\mathbf{X}'\mathbf{y}$, which, in this case, turn out to be

$$\mathbf{X}'\mathbf{y} = \begin{pmatrix} -64.2007 \\ -52.6743 \\ -42.2025 \end{pmatrix} \quad (19)$$

The least-squares estimate is thus

$$\hat{\beta} = \begin{pmatrix} 0.5778 \\ -2.3856 \\ -0.0446 \end{pmatrix} \quad (20)$$

From the data, we also calculated the estimated variance of the noise, and found the value

$$\hat{\sigma}^2 = 0.883 \quad (21)$$

The data are represented in Figure 2. The dashed line is the true regression $f_{\beta}(x)$. The solid line is the estimated regression $f_{\hat{\beta}}(x)$.

The estimated regression is barely distinguishable from a straight line. Indeed, the value $\hat{\beta}_3 = -0.0446$

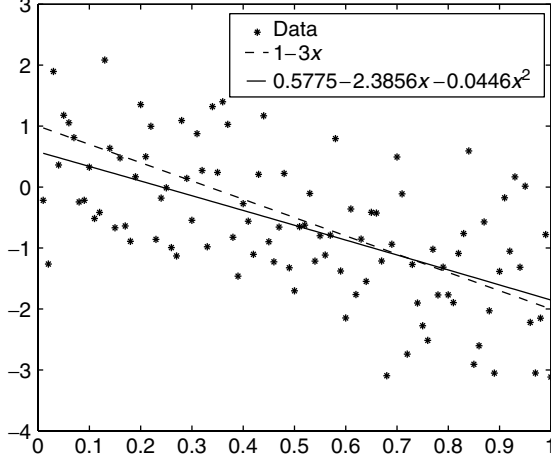


Figure 2 Observed data, true regression (dashed line), and least-squares estimate (solid line)

of the quadratic term is small. The estimated variance of $\hat{\beta}_3$ is

$$\text{var}(\hat{\beta}_3) = 1.8009 \times 0.883 = 1.5902 \quad (22)$$

Using $c = 1.96$ in equation (12), we find the confidence interval

$$\beta_3 \in -0.0446 \pm 1.96\sqrt{1.5902} = [-2.5162, 2.470] \quad (23)$$

Thus, β_3 is not significantly different from zero at the 5% level, and, hence, we do not reject the hypothesis $H_0 : \beta_3 = 0$.

Below, we will consider general test statistics for testing hypotheses on β . In this particular case, the test statistic takes the form

$$T^2 = \frac{\hat{\beta}_3^2}{\text{var}(\hat{\beta}_3)} = 0.0012 \quad (24)$$

Using this test statistic is equivalent to the above method based on the confidence interval. Indeed, as $T^2 < (1.96)^2$, we do not reject the hypothesis $H_0 : \beta_3 = 0$.

Under the hypothesis $H_0 : \beta_3 = 0$, we use the least-squares estimator

$$\begin{pmatrix} \hat{\beta}_{1,0} \\ \hat{\beta}_{2,0} \end{pmatrix} = (\mathbf{X}_0' \mathbf{X}_0)^{-1} \mathbf{X}_0' \mathbf{y} = \begin{pmatrix} 0.5854 \\ -2.4306 \end{pmatrix} \quad (25)$$

Here,

$$\mathbf{X}_0 = \begin{pmatrix} 1 & x_1 \\ \vdots & \vdots \\ 1 & x_n \end{pmatrix} \quad (26)$$

It is important to note that setting β_3 to zero changes the values of the least-squares estimates of β_1 and β_2 :

$$\begin{pmatrix} \hat{\beta}_{1,0} \\ \hat{\beta}_{2,0} \end{pmatrix} \neq \begin{pmatrix} \hat{\beta}_1 \\ \hat{\beta}_2 \end{pmatrix} \quad (27)$$

This is because $\hat{\beta}_3$ is correlated with $\hat{\beta}_1$ and $\hat{\beta}_2$. One may verify that the **correlation** matrix of $\hat{\beta}$ is

$$\begin{pmatrix} 1 & -0.8708 & 0.7529 \\ -0.8708 & 1 & -0.9689 \\ 0.7529 & -0.9689 & 1 \end{pmatrix} \quad (28)$$

Testing Linear Hypotheses

The testing problem considered in the numerical example is a special case of testing a linear hypothesis $H_0 : A\beta = 0$, where A is some $r \times p$ matrix. As another example of such a hypothesis, suppose we want to test whether two coefficients are equal, say $H_0 : \beta_1 = \beta_2$. This means there is one restriction $r = 1$, and we can take A as the $1 \times p$ row vector

$$A = (1, -1, 0, \dots, 0) \quad (29)$$

In general, we assume that there are no linear dependencies in the r restrictions $A\beta = 0$. To test the linear hypothesis, we use the statistic

$$T^2 = \frac{\|\mathbf{X}\hat{\beta}_0 - \mathbf{X}\hat{\beta}\|^2/r}{\hat{\sigma}^2} \quad (30)$$

where $\hat{\beta}_0$ is the least-squares estimator under $H_0 : A\beta = 0$. In the numerical example, this statistic takes the form given in equation (24). When the noise is normally distributed, critical values can be found in a table for the F distribution with r and $n - p$ **degrees of freedom**. For large n , approximate critical values are in the table of the χ^2 distribution with r degrees of freedom.

Some Extensions

Weighted Least-Squares

In many cases, the variance σ_i^2 of the noise at measurement i depends on x_i . Observations in which

4 Least-Squares Estimation

σ_i^2 is large are less accurate, and, hence, should play a smaller role in the estimation of β . The *weighted* least-squares estimator is the value of b that minimizes the criterion

$$\sum_{i=1}^n \frac{(y_i - f_b(x_i))^2}{\sigma_i^2} \quad (31)$$

over all possible b . In the linear case, this criterion is numerically of the same form, as we can make the change of variables $\tilde{y}_i = y_i/\sigma_i$ and $\tilde{\mathbf{x}}_{i,j} = \mathbf{x}_{i,j}/\sigma_i$.

Nonlinear Regression

When f_β is a nonlinear function of β , one usually needs iterative algorithms to find the least-squares estimator. The variance can then be approximated

as in the linear case, with $\dot{f}_\beta(x_i)$ taking the role of the rows of $\mathbf{X}$. Here, $\dot{f}_\beta(x_i) = \partial f_\beta(x_i)/\partial \beta$ is the row vector of derivatives of $f_\beta(x_i)$. For more details, see e.g. [1].

Reference

- [1] Seber, G.A.F. & Wild, C.J. (2003). *Nonlinear Regression*, John Wiley & Sons, New York.

SARA A. VAN DE GEER

Article originally published in Encyclopedia of Statistics in Behavioral Science (2005, John Wiley & Sons, Ltd). Minor revisions for this publication by Jeroen de Mast.

Moments

Introduction

Moments are an important class of *expectations* used to describe probability distributions. Together, the entire set of moments of a random variable will generally determine its probability distribution exactly.

There are three main types of moments:

1. raw moments,
2. central moments, and
3. factorial moments.

Raw Moments

Where a random variable is denoted by the letter X , and k is any positive integer, the k th raw moment of X is defined as $E(X^k)$, the expectation of the random variable X raised to the power k . Raw moments are usually denoted by μ'_k where $\mu'_k = E(X^k)$, if that expectation exists. The first raw moment of X is $\mu'_1 = E(X)$, also referred to as the *mean of X* . The second raw moment of X is $\mu'_2 = E(X^2)$, the third $\mu'_3 = E(X^3)$, and so on. If the k th moment of X exists, then all moments of lower order also exist. Therefore, if the $E(X^2)$ exists, it follows that $E(X)$ exists.

Central Moments

Where X again denotes a random variable and k is any positive integer, the k th central moment of X is defined as $E[(X - c)^k]$, the expectation of X minus a constant, all raised to the power k . Where the constant is the mean of the random variable, this is referred to as the *k th central moment around the mean*. Central moments around the mean are usually denoted by μ_k where $\mu_k = E[(X - \mu_X)^k]$. The first central moment is equal to zero as $\mu_1 = E[(X - \mu_X)] = E(X) - E(\mu_X) = \mu_X - \mu_X = 0$. In fact, if the probability distribution is symmetrical around the mean (e.g., the **normal distribution**) all odd central moments of X around the mean are equal to zero, provided they exist. The most important central moment is the second central moment of X around the mean. This

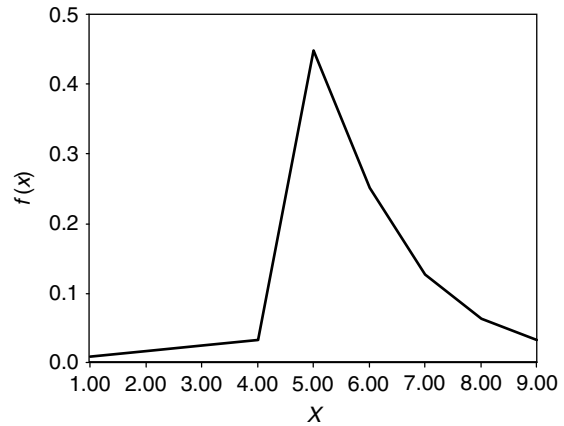


Figure 1 An example of an asymmetrical probability distribution where the third central moment around the mean is equal to zero

is $\mu_2 = E[(X - \mu_X)^2]$, the *variance* (see **Variance**) of X .

The third central moment about the mean, $\mu_3 = E[(X - \mu_X)^3]$, is sometimes used as a measure of asymmetry or **skewness**. As an odd central moment around the mean, μ_3 is equal to zero if the probability distribution is symmetrical. If the distribution is negatively skewed, the third central moment about the mean is negative, and if it is positively skewed, the third central moment around the mean is positive. Thus, knowledge of the shape of the distribution provides information about the value of μ_3 . Knowledge of μ_3 does not necessarily provide information about the shape of the distribution, however. A value of zero may not indicate that the distribution is symmetrical. As an illustration of this, μ_3 is approximately equal to zero for the distribution depicted in Figure 1, but it is not symmetrical. The third central moment is therefore not used much in practice.

The fourth central moment about the mean, $\mu_4 = E[(X - \mu_X)^4]$, is sometimes used as a measure of excess or **kurtosis**. This is the degree of flatness of the distribution near its center. The coefficient of kurtosis ($\mu_4/\sigma^4 - 3$) is sometimes used to compare an observed distribution to that of a normal curve. Positive values are thought to be indicative of a distribution that is more peaked around its center than that of a normal curve, and negative values are thought to be indicative of a distribution that is more flat around

its center than that of a normal curve. However, as was the case for the third central moment around the mean, the coefficient of kurtosis does not always indicate what it is supposed to.

Factorial Moments

Finally, where X denotes a random variable and k is any positive integer, the k th factorial moment of X is defined as the following expectation $E[X(X-1)\dots(X-k+1)]$. The first factorial moment of X is therefore $E(X)$, the second factorial moment of X is $E[(X-1+1)(X-2+1)] = E(X^2 - X)$, and so on. Factorial moments are easier to calculate than raw moments for some random variables (usually discrete). As raw moments can be obtained from factorial moments and *vice versa*, it is sometimes easier to obtain the raw moments for a random variable from its factorial moments.

Moment Generating Function

For each type of moment, there is a function that can be used to generate all of the moments of a random variable or probability distribution. This is referred to as the *moment generating function* and denoted by mgf, $m_X(t)$ or $m(t)$. In practice, however,

it is often easier to calculate moments directly. The main use of the mgf is therefore in characterizing a distribution and for theoretical purposes. For instance, if a mgf of a random variable exists, then this mgf uniquely determines the corresponding distribution function. As such, it can be shown that if the mgfs of two random variables both exist and are equal for all values of t in an interval around zero, then the two cumulative distribution functions are equal. However, existence of all moments is not equivalent to existence of the mgf.

More information on the topic of moments and mgfs is given in [1-3].

References

- [1] Casella, G. & Berger, R.L. (1990). *Statistical Inference*, Duxbury Press, Belmont.
- [2] DeGroot, M.H. (1986). *Probability and Statistics*, 2nd Edition, Addison-Wesley, Reading.
- [3] Mood, A.M., Graybill, F.A. & Boes, D.C. (1974). *Introduction to the Theory of Statistics*, McGraw-Hill, Singapore.

REBECCA WALWYN

Article originally published in Encyclopedia of Statistics in Behavioral Science (2005, John Wiley & Sons, Ltd). Minor revisions for this publication by Jeroen de Mast.

Probability Plots

Introduction

Probability plots are often used to compare two empirical distributions (*see Estimation*), but they are particularly useful to evaluate the shape of an empirical distribution against a theoretical distribution. For example, we can use a probability plot to examine the degree to which an empirical distribution differs from a **normal distribution**.

There are two major forms of probability plots, known as *probability–probability (P–P) plots* and *quantile–quantile (Q–Q) plots*. They both serve essentially the same purpose, but plot different functions of the data.

An Example

It is easiest to see the difference between P–P and Q–Q plots with a simple example. The data displayed in Table 1 are derived from data collected by Compas (personal communication) on the effects of stress on

cancer patients. Each of the 41 patients completed the externalizing behavior scale of the Brief Symptom Inventory. We wish to judge the normality of the sample data. A histogram for these data is presented in Figure 1 with a normal distribution superimposed. The sample mean is 12.10, and the sample standard deviation is 10.32.

P–P Plots

A P–P plot displays the cumulative probability of the obtained data on the x axis and the corresponding expected cumulative probability for the reference distribution (in this case, the normal distribution) on the y axis. For example, a raw score of 15 for our data had a cumulative probability of 0.756. It also had a z score, given the mean and standard deviation of the sample, of 0.28. For a normal distribution, we expect to find 61% of the distribution falling at or below $z = 0.28$. So, for this observation, we had considerably more scores (75.6%) falling at or below 15 than we would have expected if the distribution were normal (61%). From Table 1, you

Table 1 Frequency distribution of externalizing scores with normal probabilities and **quantiles**

Score	Frequency	Percent	Cumulative percent	z	CDF normal	Normal quantile
0	3	7.3	7.3	−1.17	0.12	−3.01
1	2	4.9	12.2	−1.07	0.14	−0.03
2	2	4.9	17.1	−0.98	0.16	2.19
4	1	2.4	19.5	−0.78	0.22	3.13
5	3	7.3	26.8	−0.69	0.25	5.61
6	2	4.9	31.7	−0.59	0.28	7.08
7	1	2.4	34.1	−0.49	0.31	7.77
8	1	2.4	36.6	−0.40	0.35	8.46
9	6	14.6	51.2	−0.30	0.38	12.31
10	2	4.9	56.1	−0.20	0.42	13.58
11	3	7.3	63.4	−0.1	0.46	15.54
12	2	4.9	68.3	−0.01	0.50	16.92
14	2	4.9	73.2	0.18	0.57	18.39
15	1	2.4	75.6	0.28	0.61	19.16
16	2	4.9	80.5	0.38	0.65	20.87
19	1	2.4	82.9	0.67	0.75	21.81
24	1	2.4	85.4	1.15	0.88	22.88
25	1	2.4	87.8	1.25	0.89	24.03
28	1	2.4	90.2	1.54	0.94	25.35
29	1	2.4	92.7	1.64	0.95	27.01
31	1	2.4	95.1	1.83	0.97	29.08
40	1	2.4	97.6	2.70	0.99	32.41
42	1	2.4	100.0	2.90	1.00	36.02

2 Probability Plots

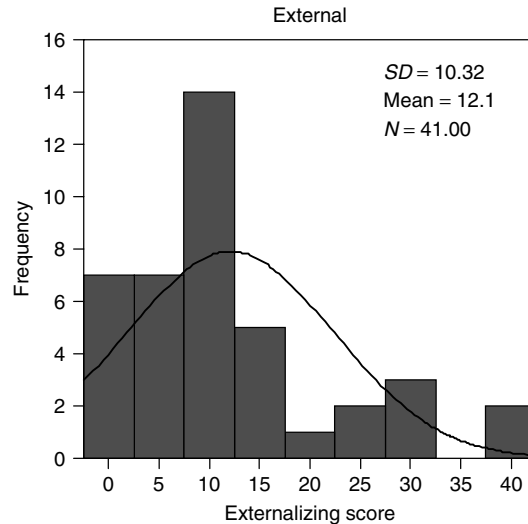


Figure 1 Histogram of externalizing scores with normal curve superimposed

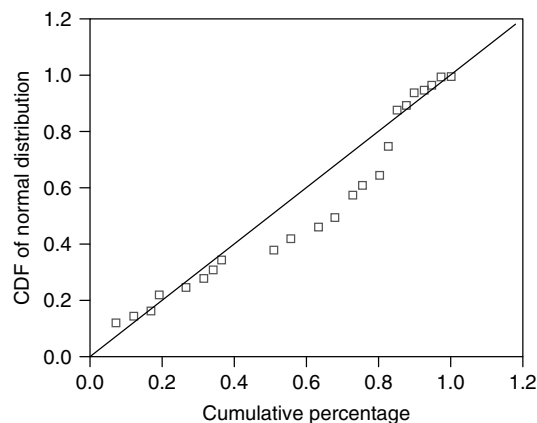


Figure 2 P-P plot of obtained distribution against the normal distribution

can see the results of similar calculations for the full data set. These are displayed in columns 4 and 6.

If we plot the obtained cumulative percentage on the x axis and the expected cumulative percentage on the y axis, we obtain the results shown in Figure 2. Here, a line drawn at 45° is superimposed. If the data had been exactly normally distributed, the points would have fallen on that line. It is clear from the figure that this was not the case. The points deviate noticeably from the line, and thus, from normality.

Q-Q Plots

A Q-Q plot resembles a P-P plot in many ways, but it displays the observed values of the data on the x axis and the corresponding expected values from a normal distribution on the y axis. The expected values are based on the *quantiles* of the distribution. To take our example of a score of 15 again, we know that 75.6% of the observations fell at or below 15, placing 15 at the 75.6 percentile. If we had a normal distribution with a mean of 12.10 and a standard deviation of 10.32, we would

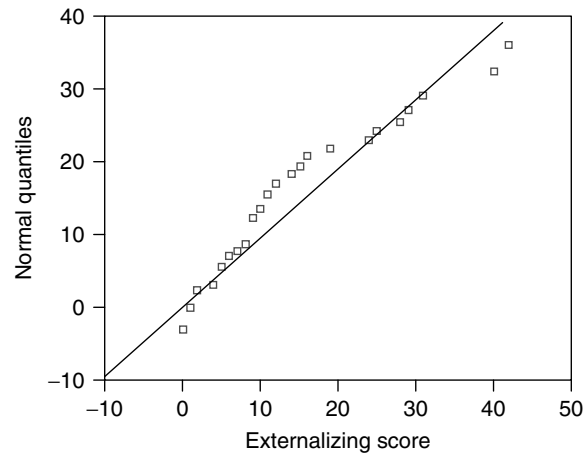


Figure 3 Q–Q plot of obtained distribution against the normal distribution

expect the 75.6 percentile to fall at a score of 19.16. We can make similar calculations for the remaining data, and these are shown in column 7 of Table 1. (The final cumulative percentage has been treated as 99.9 instead of 100 because the normal distribution runs to infinity and the quantile for 100% would be undefined.)

The plot of these values can be seen in Figure 3, where, again, a straight line is superimposed representing the results that we would have if the data were perfectly normally distributed. Again, we can see significant departure from normality.

Related Article

Half-Normal Plot.

DAVID C. HOWELL

Article originally published in Encyclopedia of Statistics in Behavioral Science (2005, John Wiley & Sons, Ltd). Minor revisions for this publication by Jeroen de Mast.

Quantiles

A quantile of a distribution is the value such that a proportion p of the population values are less than or equal to it. In other words, the p th quantile is the value that cuts off a certain proportion p of the area under the probability distribution function (see **Cumulative Distribution Function (CDF)**). It is sometimes known as a *theoretical quantile* or a *population quantile* and is denoted by $Q_i(p)$. Quantiles for common distributions are easily obtained using routines for the inverse cumulative distribution function (inverse **cdf**) found in many popular statistical packages.

For example, the 0.5 quantile (which we also recognize as the median or 50th percentile) of a standard **normal distribution** is $P(x \leq x_{0.5}) = \Phi^{-1}(0.5) = 0$, where $\Phi^{-1}(\cdot)$ denotes the inverse cdf of that distribution. Equally easily, we locate the 0.5 quantile for an exponential distribution (see **Probability Density Functions**) with mean of 1 as 0.6931. Other useful quantiles such as the lower and upper **quartiles** (0.25 and 0.75 quantiles) can be obtained in a similar way.

However, we are more likely to be concerned with quantiles associated with data. Here, the *empirical quantile* $Q(p)$ is that point on the data scale that splits the data (the empirical distribution; see **Estimation**) into two parts so that a proportion p of the observations fall below it and the rest lie above. Although each data point is itself some quantile, obtaining a *specific* quantile requires a more pragmatic approach.

Consider the following 10 observations, ordered from the smallest to the largest.

274, 302, 334, 430, 489, 703, 978, 1656, 1697, 2745

Imagine that we are interested in the 0.5 and 0.85 quantiles for these data. With 10 data points, the point $Q(0.5)$ on the data scale that has half of the observations below it and half above does not actually correspond to a member of the sample but lies *somewhere* between the middle two values 489 and 703. Also, we have no way of splitting off exactly 0.85 of the data. So, where do we go from here?

Several different remedies have been proposed, each of which essentially entails a modest redefinition of quantile. For most practical purposes, any differences between the estimates obtained from the

various versions will be negligible. The only method that we shall consider here is the one that is also used by Minitab and SPSS (but see [1] for a description of an alternative favored by the *exploratory data analysis* (EDA) community).

Suppose that we have n ordered observations x_i ($i = 1-n$). These split the data scale into $n + 1$ segments: one below the smallest observation, $n - 1$ between each adjacent pair of values and one above the largest. The proportion p of the distribution that lies below the i th observation is then estimated by $i/(n + 1)$. Setting this equal to p gives $i = p(n + 1)$. If i is an integer, the i th observation is taken to be $Q(p)$. Otherwise, we take the integer part of i , say j , and look for $Q(p)$ between the j th and $j + 1$ th observations. Assuming simple linear interpolation, $Q(p)$ lies a fraction $(i - j)$ of the way from x_j to x_{j+1} and is estimated by

$$Q(p) = x_j + (x_{j+1} - x_j) \times (i - j) \quad (1)$$

A point to note is that we cannot determine quantiles in the tails of the distribution for $p < 1/(n + 1)$ or $p > n/(n + 1)$ since these take us outside the range of the data. If extrapolation is unavoidable, it is safest to define $Q(p)$ for all P values in these two regions as x_1 and x_n , respectively.

Returning to the 10 data points given before, we note that for the 0.5 quantile, $i = 0.5 \times 11 = 5.5$, which is not an integer so $Q(0.5) = 489 + (703 - 489) \times (5.5 - 5) = 596$, which is exactly half way between the fifth and sixth observations. For the 0.85 quantile, $i = 0.85 \times 11 = 9.35$, so the required value will lie 0.35 of the way between the 9th and 10th observations, and is estimated by $Q(0.85) = 1697 + (2745 - 1697) \times (9.35 - 9) = 2063.8$.

From a plot of the empirical quantiles $Q(p)$, that is, the data, against proportion p as in Figure 1, we can get some feeling for the distributional properties of our sample (see also **Probability Plots**). For example, we can read off rough values of important quantiles such as the median and quartiles. Moreover, the increasing steepness of the slope on the right-hand side indicates a lower density of data points in that region and thereby alerts us to **skewness** and possible **outliers**. Note that we have also defined $Q(p) = x_1$ when $p < 1/11$ and $Q(p) = x_{10}$ when $p > 10/11$. Clearly, extrapolation would not be advisable, especially at the upper end.

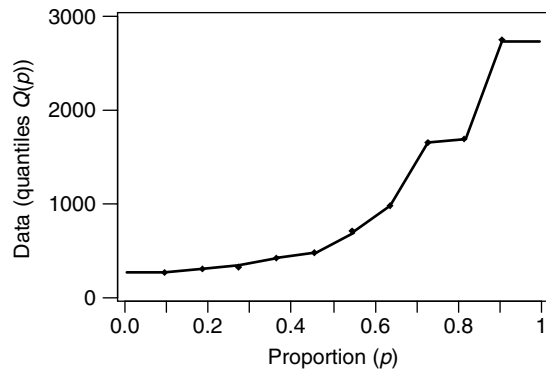


Figure 1 Quantile plot for Necker cube reversal times

It is possible to obtain *confidence intervals* for quantiles, and information on how these are constructed can be found, for example in [2]. Routines for these procedures are available within certain statistical packages such as SAS and S-PLUS.

Quantiles play a central role within *exploratory data analysis*. The median ($Q(0.5)$), upper and lower quartiles ($Q(0.75)$ and $Q(0.25)$), and the maximum and minimum values, which constitute the so-called five-number summary, are the basic elements of a box plot. An empirical quantile–quantile plot (see **Probability Plots**) is a scatterplot of the paired empirical quantiles of two samples of data, while the symmetry plot is a scatterplot of the paired empirical quantiles from a single sample, which has

been split and folded at the median. Of course, for these latter plots there is no need for the P values to be explicitly identified.

Finally, it is worth noting that the term *quantile* is a general one referring to the class of statistics, which split a distribution according to a proportion p that can be *any* value between 0 and 1. As we have seen, well-known measures that have specific P values, such as the median, quartile, and percentile (on Galton's original idea of splitting by 100ths [3]), belong to this group. Nowadays, however, quantile and percentile are used almost interchangeably: the only difference is whether p is expressed as a decimal fraction or as a percentage.

References

- [1] Chambers, J.M., Cleveland, W.S., Kleiner, B. & Tukey, P.A. (1983). *Graphical Methods for Data Analysis*, Duxbury, Boston.
- [2] Conover, W.J. (1999). *Practical Nonparametric Statistics*, 3rd Edition, John Wiley & Sons, New York.
- [3] Galton, F. (1885). Anthropometric percentiles, *Nature* **31**, 223–225.

PAT LOVIE

Article originally published in Encyclopedia of Statistics in Behavioral Science (2005, John Wiley & Sons, Ltd). Minor revisions for this publication by Jeroen de Mast.

Time Series Analysis

Time Series Data

A time series is a set of observations x_t , each one associated with a particular time t , and usually displayed in a *time series plot* of x_t as a function of t . The set of times T at which observations are recorded may be a discrete set, as is the case when the observations are recorded at uniformly spaced times (e.g., daily rainfall, hourly temperature, annual income, etc.), or it may be a continuous interval, as when the data is recorded continuously (e.g., by a seismograph or electrocardiograph). A very large number of practical problems involve observations that are made at uniformly spaced times. For this reason, this article focuses on this case, indicating briefly how missing values and irregularly spaced data can be handled.

Example 1 Figure 1 shows the number of accidental deaths recorded monthly in the United States for the years 1973–1978. The graph strongly suggests (as is usually the case for monthly data) the presence of a periodic component with period 12 corresponding to the seasonal cycle of 12 months, as well as a smooth trend accounting for the relatively slow change in level of the series, and a random component accounting for irregular deviations from a deterministic model involving trend and seasonal components only.

Example 2 Figure 2 shows the closing value in Australian dollars of the Australian All-Ordinaries index (an average of 100 stocks sold on the Australian Stock Exchange) on 521 successive trading days, ending on July 18, 1994. It displays irregular variation around a rather strong trend. Figure 3 shows the daily percentage changes in closing value of the index for each of the 520 days ending on July 18, 1994. The trend apparent in Figure 2 has virtually disappeared, and the series appears to be varying randomly around a mean value close to zero.

Objectives

The objectives of time series analysis are many and varied, depending on the particular field of application. From the observations $x_1, \dots, x_n$, we

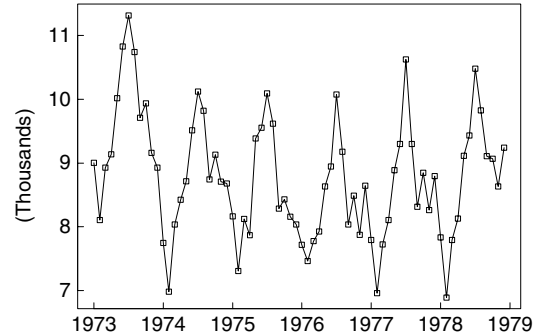


Figure 1 Monthly accidental deaths in the United States from 1973 through 1978

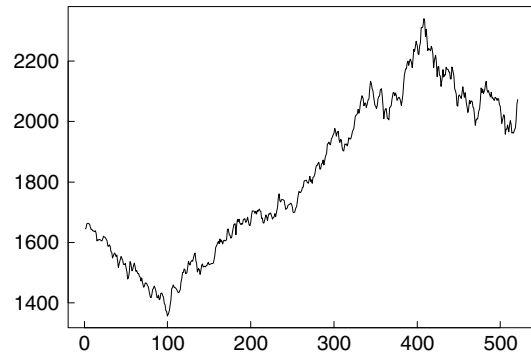


Figure 2 The Australian All-Ordinaries index of stock prices on 521 successive trading days up to July 18, 1994

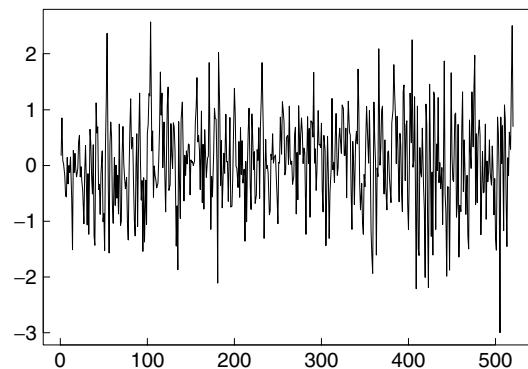


Figure 3 The daily percentage changes in the All-Ordinaries index over the same period

may wish to make inferences about the way in which the data are generated, to predict future values of

2 Time Series Analysis

the series (forecasting), to detect a “signal” hidden in noisy data, or simply to find a compact description of the available observations. In industrial statistics, time series analysis plays an important role in process monitoring and adjustment (see [1]).

In order to achieve these goals, it is necessary to postulate a mathematical model (or family of models), according to which we suppose that the data are generated. Once an appropriate family has been selected, we then select a *specific* model by estimating model parameters and checking the resulting model for goodness of fit to the data. Once we are satisfied that the selected model provides a good representation of the data, we use it to address questions of interest. It is rarely the case that there is a “true” mathematical model underlying empirical data, however, systematic procedures have been developed for selecting the best model, according to clearly specified criteria, within a broad class of candidates.

Time Series Models

As indicated above, the graph of the accidental deaths series in Figure 1 suggests representing x_t as the sum of a slowly varying *trend* component, a period-12 *seasonal* component, and a *random* component that accounts for the irregular deviations from the sum of the other two components. In order to take account of the randomness, we suppose that for each t , the observation x_t is just one of many possible values of a random variable X_t that we *might* have observed. This leads to the following *classical decomposition model* for the accidental deaths data,

$$X_t = m_t + s_t + Y_t, \quad t = 1, 2, 3, \dots \quad (1)$$

where the sequence $\{m_t\}$ is the *trend* component describing the long-term movement in the level of the series, $\{s_t\}$ is a *seasonal* component with known period (in this case, 12), and $\{Y_t\}$ is a sequence of random variables with mean zero, referred to as the *random component*. If we can characterize m_t , s_t , and Y_t in simple terms and in such a way that the model (1) provides a good representation of the data, then we can proceed to use the model to make forecasts or to address other questions related to the series. Completing the specification of the model by estimating the trend and seasonal components and characterizing the random component is a major part of time series

analysis. The model (1) is sometimes referred to as an *additive decomposition model*. Provided the observations are all positive, the *multiplicative model*,

$$X_t = m_t s_t Y_t \quad (2)$$

can be reduced to an additive model by taking logarithms of each side to get an additive model for the logarithms of the data.

The general form of the additive model (1) supposes that the seasonal component s_t has known period d (12 for monthly data, 4 for quarterly data, etc.), and satisfies the conditions

$$s_{t+d} = s_t \text{ and } \sum_{t=1}^d s_t = 0 \quad (3)$$

while $\{Y_t\}$ is a *weakly stationary sequence of random variables*, that is, a sequence of random variables satisfying the conditions,

$$E(Y_t) = \mu, \quad E(Y_t^2) < \infty \quad (4)$$

and

$$\text{Cov}(Y_{t+h}, Y_t) = \gamma(h) \text{ for all } t \quad (5)$$

with $\mu = 0$. The function γ is called the *autocovariance function* of the sequence $\{Y_t\}$, and the value $\gamma(h)$ is the *autocovariance at lag h* . In the special case when the random variables Y_t are independent and identically distributed, the model (1) is a classical regression model and $\gamma(h) = 0$ for all $h \neq 0$. However, in time series analysis, it is the dependence between Y_{t+h} and Y_t that is of special interest, and it allows the possibility of using past observations to obtain forecasts of future values that are better in some average sense than just using the expected value of the series. A measure of this dependence is provided by the autocovariance function. A more convenient measure (since it is independent of the origin and scale of measurement of Y_t) is the *autocorrelation function* (see **Autocorrelation Function**),

$$\rho(h) = \frac{\gamma(h)}{\gamma(0)} \quad (6)$$

From observed values $y_1, \dots, y_n$ of a weakly stationary sequence of random variables $\{Y_t\}$, good estimators of the mean $\mu = E(Y_t)$ and the autocovariance function $\gamma(h)$ are the *sample mean* and *sample autocovariance function*,

$$\hat{\mu} = \frac{1}{n} \sum_{i=1}^n y_i \quad (7)$$

and

$$\hat{\gamma}(h) = \frac{1}{n} \sum_{i=1}^{n-|h|} (y_{i+|h|} - \hat{\mu})(y_i - \hat{\mu}), \quad -n < h < n \quad (8)$$

respectively. The autocorrelation function of $\{Y_t\}$ is estimated by the *sample autocorrelation function*,

$$\hat{\rho}(h) = \frac{\hat{\gamma}(h)}{\hat{\gamma}(0)} \quad (9)$$

Elementary techniques for estimating m_t and s_t can be found in many texts on time series analysis (e.g. [2]). Once estimators $\hat{m}_t$ of m_t and $\hat{s}_t$ of s_t have been obtained, they can be subtracted from the observations to yield the **residuals**,

$$y_t = x_t - \hat{m}_t - \hat{s}_t \quad (10)$$

A stationary time series model can then be fitted to the residual series to complete the specification of the model. The model is usually chosen from the class of *autoregressive moving average* (or **ARMA**) *processes*, defined below in the section titled “ARMA processes”.

Instead of estimating and subtracting the trend and seasonal components to generate a sequence of residuals, an alternative approach, developed by Box *et al.* [3], is to apply *difference operators* to the original series to remove trend and seasonality. The *backward shift operator* B is an operator that, when applied to X_t , gives X_{t-1} . Thus,

$$BX_t = X_{t-1}, \quad B^j X_t = X_{t-j}, \quad j = 2, 3, \dots \quad (11)$$

The *lag-1 difference operator* is the operator $\nabla = (1 - B)$. Thus,

$$\nabla X_t = (1 - B)X_t = X_t - X_{t-1} \quad (12)$$

When applied to a polynomial trend of degree p , the operator ∇ reduces it to a polynomial of degree $p - 1$. The operator ∇^p , denoting p successive applications of ∇ , therefore, reduces any polynomial trend of degree p to a constant. Usually, a small number of applications of ∇ is sufficient to eliminate trends encountered in practice. Application of the *lag- d difference operator*, $\nabla_d = (1 - B^d)$ (not to be confused with ∇^d) to X_t gives

$$\nabla_d X_t = (1 - B^d)X_t = X_t - X_{t-d} \quad (13)$$

eliminating any seasonal component with period d . In the Box–Jenkins approach to time series modeling, the operators ∇ and ∇_d are applied as many times as is necessary to eliminate trend and seasonality, and the sample mean of the differenced data is subtracted to generate a sequence of residuals y_t , which are then modeled with a suitably chosen ARMA process in the same way as the residuals (equation (10)).

Figure 4 shows the effect of applying the operator ∇_{12} to the accidental deaths series of Figure 1. The seasonal component is no longer apparent, but there is still an approximately linear trend. Further application of the operator ∇ yields the series shown in Figure 5 with relatively constant level. This new series is a good candidate for representation by a stationary time series model.

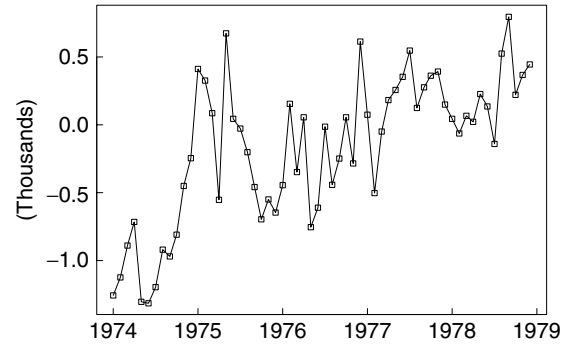


Figure 4 The differenced series $\{\nabla_{12}x_t, t = 13, \dots, 72\}$ derived from the monthly accidental deaths $\{x_1, \dots, x_{72}\}$ shown in Figure 1

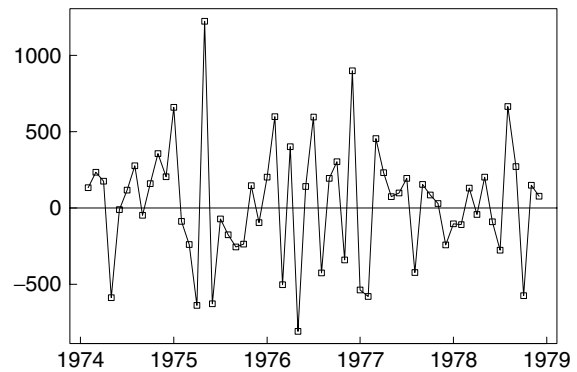


Figure 5 The differenced series $\{\nabla\nabla_{12}x_t, t = 14, \dots, 72\}$ derived from the monthly accidental deaths $\{x_1, \dots, x_{72}\}$ shown in Figure 1

The daily percentage returns on the Australian All-Ordinaries index shown in Figure 2 already show no sign of trend or seasonality, and can be modeled as a stationary sequence without preliminary detrending or deseasonalizing.

In cases where the variability of the observed data appears to change with the level of the data, a preliminary transformation, prior to detrending and deseasonalizing, may be required to stabilize the variability. For this purpose, a member of the family of *Box–Cox transformations* (see **Box and Cox Transformation**; [2]) is frequently used.

ARMA Processes

For modeling the residuals $\{y_t\}$ (found as described above), a very useful parametric family of zero-mean stationary sequences is furnished by the *ARMA processes*. The $\text{ARMA}(p, q)$ process $\{Y_t\}$ with autoregressive coefficients, $\phi_1, \dots, \phi_p$, moving average coefficients, $\theta_1, \dots, \theta_q$, and white noise variance σ^2 , is defined as a weakly stationary solution of the difference equations,

$$(1 - \phi_1 B - \dots - \phi_p B^p)Y_t = (1 + \theta_1 + \dots + \theta_q B^q)Z_t, \quad t = 0, \pm 1, \pm 2, \dots \quad (14)$$

where B is the backward shift operator, the polynomials $\phi(z) = 1 - \phi_1 z - \dots - \phi_p z^p$ and $\theta(z) = 1 + \theta_1 z + \dots + \theta_q z^q$ have no common factors, and $\{Z_t\}$ is a sequence of uncorrelated random variables with mean zero and variance σ^2 . Such a sequence $\{Z_t\}$ is said to be *white noise with mean zero and variance σ^2* , indicated more concisely by writing $\{Z_t\} \sim \text{WN}(0, \sigma^2)$.

Equation (14) has a unique stationary solution if, and only if, the equation $\phi(z) = 0$ has no root with $|z| = 1$, however, the possible values of $\phi_1, \dots, \phi_p$ are usually assumed to satisfy the stronger restriction,

$$\phi(z) \neq 0 \text{ for all complex } z \text{ such that } |z| \leq 1 \quad (15)$$

The unique weakly stationary solution of equation (14) is then

$$Y_t = \sum_{j=0}^{\infty} \psi_j Z_{t-j} \quad (16)$$

where ψ_j is the coefficient of z^j in the power-series expansion,

$$\frac{\theta(z)}{\phi(z)} = \sum_{j=0}^{\infty} \psi_j z^j, \quad |z| \leq 1 \quad (17)$$

Since Y_t in equation (16) is a function only of Z_s , $s \leq t$, the series $\{Y_t\}$ is said to be a *causal function* of $\{Z_t\}$, and the condition (15) is called the *causality condition* for the process (14). (Condition (15) is also frequently referred to as a *stability condition*.)

In the causal case, simple recursions are available for the numerical calculation of the sequence $\{\psi_j\}$ from the autoregressive and moving average coefficients $\phi_1, \dots, \phi_p$, and $\theta_1, \dots, \theta_q$ (see e.g. [2]). For every noncausal ARMA process, there is a causal ARMA process with the same autocovariance function, and, *under the assumption that all the joint distributions of the process are multivariate normal*, with the same joint distributions. This is one reason for restricting attention to causal models. Another practical reason is that if $\{Y_t\}$ is to be simulated sequentially from the sequence $\{Z_t\}$, the causal representation (16) of Y_t does not involve *future* values Z_{t+h} , $h > 0$.

A key feature of ARMA processes for modeling dependence in a sequence of observations is the extraordinarily large range of autocorrelation functions exhibited by ARMA processes of different orders (p, q) as the coefficients $\phi_1, \dots, \phi_p$, and $\theta_1, \dots, \theta_q$, are varied. For example, if we take any set of observations, $x_1, \dots, x_n$ and compute the sample autocorrelations, $\hat{\rho}(0)(=1), \hat{\rho}(1), \dots, \hat{\rho}(n-1)$, then for any $k < n$, it is always possible to find a causal $\text{AR}(k)$ process with autocorrelations $\rho(j)$ satisfying $\rho(j) = \hat{\rho}(j)$ for every $j \leq k$.

The mean of the ARMA process defined by equation (14) is $E(Y_t) = 0$, and the autocovariance function can be found from the equations obtained by multiplying each side of equation (14) by Y_{t-j} , $j = 0, 1, 2, \dots$, taking expectations and solving for $\gamma(j) = E(Y_t Y_{t-j})$. Details can be found in [2].

Example 3 The causal $\text{ARMA}(1, 0)$ or $\text{AR}(1)$ process is a stationary solution of the equations

$$Y_t = \phi Y_{t-1} + Z_t, \quad |\phi| < 1, \quad \{Z_t\} \sim \text{WN}(0, \sigma^2) \quad (18)$$

The autocovariance function of $\{Y_t\}$ is $\rho(h) = \sigma^2 \phi^{|h|} / (1 - \phi^2)$.

Example 4 The ARMA(0, q) or MA(q) process is the stationary series defined by

$$Y_t = \sum_{j=0}^q \theta_j Z_{t-j},$$

$$\theta_0 = 1, \{Z_t\} \sim WN(0, \sigma^2) \quad (19)$$

The autocovariance function of $\{Y_t\}$ is $\gamma(h) = \sigma^2 \sum_{j=0}^{|h|} \theta_j \theta_{j+|h|}$ if $h \leq q$ and $\gamma(h) = 0$ otherwise.

A process $\{X_t\}$ is said to be an *ARMA process with mean μ* if $\{Y_t = X_t - \mu\}$ is an ARMA process as defined by equations of the form (14).

In the following section, we consider the problem of fitting an ARMA model of the form (14), that is, $\phi(B)Y_t = \theta(B)Z_t$, to the residual series $y_1, y_2, \dots$, generated as described in the section titled “Time Series Models”. If the residuals were obtained by applying the differencing operator $(1 - B)^d$ to the original observations $x_1, x_2, \dots$, then we are effectively fitting the model

$$\phi(B)(1 - B)^d X_t = \theta(B)Z_t,$$

$$\{Z_t\} \sim WN(0, \sigma^2) \quad (20)$$

to the original data. If the order of the polynomials $\phi(B)$ and $\theta(B)$ are p and q , respectively, then the model (20) is called an *ARIMA(p, d, q) model* for $\{X_t\}$. If the residuals y_t had been generated by differencing also at some lag greater than 1 to eliminate seasonality, say, for example, $y_t = (1 - B^{12})(1 - B)^d x_t$, and the ARMA model (14) were then fitted to y_t , then the model for the original data would be a more general ARIMA model of the form,

$$\phi(B)(1 - B^{12})(1 - B)^d X_t = \theta(B)Z_t \quad (21)$$

$$\{Z_t\} \sim WN(0, \sigma^2) \quad (22)$$

Selecting and Fitting a Model to Data

In the section titled “Time Series Models”, we discussed two methods for eliminating trend and seasonality with the aim of transforming the original series to a series of residuals suitable for modeling as a zero-mean stationary series. In the section titled “ARMA Processes”, we introduced the class of ARMA models with a wide range of autocorrelation functions. By suitable choice of ARMA parameters,

it is possible to find an ARMA process whose autocorrelations match the sample autocorrelations of the residuals $y_1, y_2, \dots$, up to any specified lag. This is an intuitively appealing and natural approach to the problem of fitting a stationary time series model. However, except when the residuals are truly generated by a purely autoregressive model (i.e., an ARMA($p, 0$) model), this method turns out to give greater large-sample mean squared errors than the method of *maximum Gaussian likelihood* described in the following paragraph.

Suppose for the moment that we know the orders p and q of the ARMA process (equation (14)) that is to be fitted to the residuals $y_1, \dots, y_n$, and suppose that $\beta = (\phi_1, \dots, \phi_p, \theta_1, \dots, \theta_q, \sigma^2)$, is the vector of parameters to be estimated. The Gaussian likelihood $L(\beta; y_1, \dots, y_n)$, is the likelihood computed under the assumption that the joint distribution from which $y_1, \dots, y_n$ are drawn is multivariate normal (see **Maximum Likelihood**). Thus,

$$L(\beta; y_1, \dots, y_n) = (2\pi)^{-n/2} (\det \Gamma_n)^{-1/2} \times \exp \left(-\frac{1}{2} \mathbf{y}_n \Gamma_n^{-1} \mathbf{y}_n \right) \quad (23)$$

where $\mathbf{y}_n = (y_1, \dots, y_n)'$, Γ_n is the matrix of autocovariances $[\gamma(i - j)]_{i,j=1}^n$, and $\gamma(h)$ is the autocovariance function of the model defined by equation (14). Although direct calculation of L is a daunting task, L can be reexpressed in the *innovations form* of Schweppe [4], which is readily calculated from the minimum **mean squared error** one-step linear predictors of the observations and their mean squared errors. These in turn can be readily calculated from the *innovations algorithm* (see [5]).

At first glance, maximization of Gaussian likelihood when the observations appear to be non-Gaussian may seem strange. However, if the noise sequence $\{Z_t\}$ in the model (14) is any independent identically distributed sequence (with finite variance), the large-sample joint distribution of the estimators (assuming the true orders are p and q) is the same as in the Gaussian case (see [5, 6]). This large-sample distribution has a relatively simple Gaussian form that can be used to specify large-sample **confidence intervals** for the parameters (under the assumption that the observations are generated by the fitted model).

Maximization of L with respect to the parameter vector β is a nonlinear optimization problem, requiring the use of an efficient numerical maximization

algorithm. For this reason, a variety of simpler **estimation** methods have been developed. These generally lead to less efficient estimators, which can be used as starting points for the nonlinear optimization. Notable among these are the Hannan–Rissanen algorithm for general ARMA processes, and the Yule–Walker and Burg algorithms for purely autoregressive processes.

The previous discussion assumes that the orders p and q are *known*. However, this is rarely, if ever, the case, and they must be chosen on the basis of the observations. The choice of p and q is referred to as the problem of *order selection*. The shape of the sample autocorrelation function gives some clue as to the order of the ARMA(p, q) model that best represents the data. For example, a sample autocorrelation function that appears to be roughly of the form $\phi^{|h|}$ for some ϕ such that $|\phi| < 1$ suggests (see Example 3) that an AR(1) model might be appropriate, while a sample autocorrelation function that is small in absolute value for lags $h > q$ suggests (see Example 4) that an MA(r) model with $r \leq q$ might be appropriate.

A systematic approach to the problem was suggested by Akaike [7] when he introduced the *information criterion* known as the AIC. He proposed that p and q be chosen by minimizing $AIC(\hat{\beta}(p, q))$, where $\hat{\beta}(p, q)$ is the maximum-likelihood estimator of β for fixed p and q and

$$AIC(\beta) = -2 \ln(L(\beta)) + 2(p + q + 1) \quad (24)$$

The term $2(p + q + 1)$ can be regarded as a *penalty factor* that prevents the selection of excessive values for p and q and the accumulation of additional parameter estimation errors. If the data are truly generated by an ARMA(p, q) model, it has been shown that the AIC criterion tends to overestimate p and q and is not consistent as the sample-size approaches infinity. Consistent estimation of p and q can be obtained by using information criteria with heavier penalty factors such as the Bayesian information criterion (BIC). Since, however, there is rarely a “true” model generating the data, consistency is not necessarily an essential property of order selection methods. It has been shown in [8] that although the AIC criterion does not give consistent estimation of p and q , it is optimal in a certain sense with respect to the prediction of future values of the series. A refined small-sample version of AIC,

known as the AICC, has been developed in [9], and a comprehensive account of model selection can be found in [10].

Having arrived at a potential ARMA model for the data, the model should be checked for goodness of fit to the data. On the basis of the fitted model, the minimum mean squared error linear predictors $\hat{Y}_t$ of each Y_t in terms of Y_s , $s < t$, and the corresponding mean squared errors s_t can be computed (see the section titled “Prediction”). In fact, they are computed in the course of evaluating the Gaussian likelihood in its innovations form. If the fitted model is valid, the properties of the rescaled one-step prediction errors $(Y_t - \hat{Y}_t)/\sqrt{s_t}$ should be similar to those of the sequence Z_t/σ in the model (14), and can, therefore, be used to check the assumed white noise properties of $\{Z_t\}$ and whether the assumption of independence and/or normality is justified. A number of such tests are available (see e.g. [2, Chapter 5]).

Prediction

If $\{Y_t\}$ is a weakly stationary process with mean, $E(Y_t) = \mu$, and autocovariance function, $\text{Cov}(Y_{t+h}, Y_t) = \gamma(h)$, a fundamental property of conditional expectation tells us that the “best” (minimum mean squared error) predictor of Y_{n+h} , $h > 0$, in terms of $Y_1, \dots, Y_n$, is the conditional expectation $E(Y_{n+h} | Y_1, \dots, Y_n)$. However, this depends, in a complicated way, on the joint distributions of the random variables Y_t that are virtually impossible to estimate on the basis of a single series of observations $y_1, \dots, y_n$. However, if the sequence $\{Y_t\}$ is Gaussian, the best predictor of Y_{n+h} in terms of $Y_1, \dots, Y_n$ is a linear function, and can be calculated as described below.

The best *linear* predictor of Y_{n+h} in terms of $\{1, Y_1, \dots, Y_n\}$, that is, the linear combination $P_n Y_{n+h} = a_0 + a_1 Y_1 + \dots + a_n Y_n$, which minimizes the mean squared error, $E[(Y_{n+h} - a_0 - \dots - a_n Y_n)^2]$, is given by

$$P_n Y_{n+h} = \mu + \sum_{i=1}^n a_i (Y_{n+1-i} - \mu) \quad (25)$$

where the vector of coefficients $\mathbf{a} = (a_1, \dots, a_n)'$ satisfies the linear equation,

$$\Gamma_n \mathbf{a} = \boldsymbol{\gamma}_n(h) \quad (26)$$

with $\mathbf{y}_n(h) = (\gamma(h), \gamma(h+1), \dots, \gamma(h+n-1))$ and $\Gamma_n = [\gamma(i-j)]_{i,j=1}^n$. The mean squared error of the best linear predictor is

$$E(Y_{n+h} - P_n Y_{n+h})^2 = \gamma(0) - \mathbf{a}' \mathbf{y}_n(h) \quad (27)$$

Once a satisfactory model has been fitted to the sequence $y_1, \dots, y_n$, it is, therefore, a straightforward but possibly tedious matter to compute best linear predictors of future observations by using the mean and autocovariance function of the fitted model and solving the linear equations for the coefficients $a_0, \dots, a_n$. If n is large, then the set of linear equations for $a_0, \dots, a_n$ is large, and, so, recursive methods using the Levinson–Durbin algorithm or the innovations algorithm have been devised to express the solution for $n = k+1$ in terms of the solution for $n = k$, and, hence, to avoid the difficulty of inverting large matrices. For details see, for example [2, Chapter 2].

If an ARMA model is fitted to the data, the special linear structure of the ARMA process arising from the defining equations can be used to greatly simplify the calculation of the best linear h -step predictor $P_n Y_{n+h}$ and its mean squared error, $\sigma_n^2(h)$. If the fitted model is Gaussian, we can also compute 95% prediction bounds, $P_n Y_{n+h} \pm 1.96\sigma_n(h)$. For details see, for example [2, Chapter 5].

Example 5 In order to predict future values of the causal AR(1) process defined in Example 3, we can make use of the fact that linear prediction is a linear operation, and that $P_n Z_t = 0$ for $t > n$ to deduce that

$$\begin{aligned} P_n Y_{n+h} &= \phi P_n Y_{n+h-1} = \phi^2 P_n Y_{n+h-2} = \dots \\ &= \phi^h Y_n, \quad h \geq 1 \end{aligned} \quad (28)$$

In order to obtain forecasts and prediction bounds for the *original* series that were transformed to generate the residuals, we simply apply the inverse transformations to the forecasts and prediction bounds for the residuals.

The Frequency Viewpoint

The methods described so far are referred to as *time-domain methods* since they focus on the evolution in time of the sequence of random variables $X_1, X_2, \dots$ representing the observed data. If, however, we regard the sequence as a random function defined on the integers, then an alternative

approach is to consider the decomposition of that function into sinusoidal components, analogous to the Fourier decomposition of a deterministic function. This approach leads to the *spectral representation* of the sequence $\{X_t\}$, according to which every weakly stationary sequence has a representation

$$X_t = \int_{-\pi}^{\pi} e^{i\omega t} dZ(t) \quad (29)$$

where $\{Z(t), -\pi \leq t \leq \pi\}$ is a process with uncorrelated increments. A detailed discussion of this approach can be found, for example, in [5, 11, 12], but is outside the scope of this article. Intuitively, however, the expression (29) can be regarded as representing the random function $\{X_t, t = 0, \pm 1, \pm 2, \dots\}$ as the limit of a linear combination of sinusoidal functions with uncorrelated random coefficients. The analysis of weakly stationary processes by means of their spectral representation is referred to as *frequency-domain analysis* or *spectral analysis*. It is equivalent to time-domain analysis, but provides an alternative way of viewing the process, which, for some applications, may be more illuminating. For example, in the design of a structure subject to a randomly fluctuating load, it is important to be aware of the presence in the loading force of a large sinusoidal component with a particular frequency in order to ensure that this is not a resonant frequency of the structure.

Multivariate Time Series

Many time series arising in practice are best analyzed as components of some vector-valued (multivariate) time series $\{\mathbf{X}_t\} = (X_{t1}, \dots, X_{tm})'$ in which each of the component series $\{X_{ti}, t = 1, 2, 3, \dots\}$ is a univariate time series of the type already discussed. In multivariate time series modeling, the goal is to account for, and take advantage of, the dependence not only between observations of a single component at different times, but also between the different component series. For example, if X_{t1} is the daily percentage change in the closing value of the Dow–Jones Industrial Average in New York on trading day t , and if X_{t2} is the analog for the Australian All-Ordinaries index (see Example 2), then $\{\mathbf{X}_t\} = (X_{t1}, X_{t2})'$ is a bivariate time series in which there is very little evidence of autocorrelation in either of the two component series. However, there is strong evidence

of **correlation** between X_{t1} and $X_{(t+1)2}$, indicating that the Dow–Jones percentage change on day t is of value in predicting the All-Ordinaries percentage change on day $t + 1$. Such dependencies in the multivariate case can be measured by the **covariance** matrices,

$$\Gamma(t + h, t) = [\text{cov}(X_{(t+h),i}, X_{t,j})]_{i,j=1}^m \quad (30)$$

where the number of components, m , is two in this particular example. Weak stationarity of the multivariate series $\{\mathbf{X}_t\}$ is defined, as in the univariate case, to mean that all components have finite second **moments** and that the mean vectors $E(\mathbf{X}_t)$ and covariance matrices $\Gamma(t + h, t)$ are independent of t .

Much of the analysis of multivariate time series is analogous to that of univariate series, multivariate ARMA (or VARMA) processes being defined again by linear equations of the form (14) but with vector-valued arguments, and matrix coefficients. There are, however, some new important considerations arising in the multivariate case. One is the question of VARMA nonidentifiability. In the univariate case, an $\text{AR}(p)$ process cannot be reexpressed as a finite-order moving average process. However, this is not the case for $\text{VAR}(p)$ processes. There are simple examples of $\text{VAR}(1)$ processes that are also $\text{VMA}(1)$ processes. This lack of identifiability, together with the large number of parameters in a VARMA model and the complicated shape of the likelihood surface, introduces substantial additional difficulty into maximum Gaussian likelihood estimation for VARMA processes. Restricting attention to VAR processes eliminates the identifiability problem. Moreover, the fitting of a $\text{VAR}(p)$ process by equating covariance matrices up to lag p is asymptotically efficient and simple to implement using a multivariate version of the Levinson–Durbin algorithm (see [13, 14]).

For nonstationary univariate time series, we discussed in the section titled “Objectives” the use of differencing to transform the data to residuals suitable for modeling as zero-mean stationary series. In the multivariate case, the concept of *cointegration*, due to Granger [15], plays an important role in this connection. The m -dimensional vector time series $\{\mathbf{X}_t\}$ is said to be integrated of order d (or $I(d)$) if, when the difference operator ∇ is applied to each component $d - 1$ times, the resulting process is nonstationary, while if it is applied d times, the resulting series is stationary. The $I(d)$ process $\{\mathbf{X}_t\}$

is said to be cointegrated with cointegration vector $\boldsymbol{\alpha}$, if $\boldsymbol{\alpha}$ is an $m \times 1$ vector such that $\{\boldsymbol{\alpha}'\mathbf{X}_t\}$ is of order less than d . Cointegrated processes arise naturally in economics. Engle and Granger [16] give, as an illustrative example, the vector of tomato prices in northern and southern California (say X_{t1} and X_{t2} , respectively). These are linked by the fact that if one were to increase sufficiently relative to the other, the profitability of buying in one market and selling in the other would tend to drive the prices together, suggesting that although the two series separately may be nonstationary, the difference varies in a stationary manner. This corresponds to having a cointegration vector $\boldsymbol{\alpha} = (1, -1)'$. Statistical inference for multivariate models with cointegration is discussed in [17].

State-Space Models

State-space models and the associated Kalman recursions have had a profound impact on time series analysis. A linear state-space model for a (possibly multivariate) time series $\{\mathbf{Y}_t, t = 1, 2, \dots\}$ consists of two equations. The first, known as the *observation equation*, expresses the w -dimensional observation vector $\mathbf{Y}_t$ as a linear function of a v -dimensional state variable $\mathbf{X}_t$ plus noise. Thus,

$$\mathbf{Y}_t = \mathbf{G}_t \mathbf{X}_t + \mathbf{W}_t, \quad t = 1, 2, \dots \quad (31)$$

where $\{\mathbf{W}_t\}$ is a sequence of uncorrelated random vectors with $E(\mathbf{W}_t) = \mathbf{0}$, $\text{cov}(\mathbf{W}_t) = \mathbf{R}_t$, and $\{\mathbf{G}_t\}$ is a sequence of $w \times v$ matrices. The second equation, called the *state equation*, determines the state $\mathbf{X}_{t+1}$ at time $t + 1$ in terms of the previous state $\mathbf{X}_t$ and a noise term. The state equation is

$$\mathbf{X}_{t+1} = \mathbf{F}_t \mathbf{X}_t + \mathbf{V}_t, \quad t = 1, 2, \dots \quad (32)$$

where $\{\mathbf{F}_t\}$ is a sequence of $v \times v$ matrices, $\{\mathbf{V}_t\}$ is a sequence of uncorrelated random vectors with $E(\mathbf{V}_t) = \mathbf{0}$, $\text{cov}(\mathbf{V}_t) = \mathbf{Q}_t$, and $\{\mathbf{V}_t\}$ is uncorrelated with $\{\mathbf{W}_t\}$ (i.e., $E(\mathbf{W}_t \mathbf{V}_s') = \mathbf{0}$ for all s and t). To complete the specification, it is assumed that the initial state $\mathbf{X}_1$ is uncorrelated with all of the noise terms $\{\mathbf{V}_t\}$ and $\{\mathbf{W}_t\}$.

An extremely rich class of models for time series, including and going well beyond the ARIMA models described earlier, can be formulated within this framework (see [18]). In econometrics, the structural time series models developed in [19], in which

trend and seasonal components are allowed to evolve randomly, also fall into this framework. The power of the state-space formulation of time series models depends heavily on the Kalman recursions, which allow best linear predictors and best linear estimates of various model-related variables to be computed in a routine way. Time series with missing values are also readily handled in the state-space framework.

More general state-space models, in which the linear relationships (31) and (32) are replaced by the specification of conditional distributions, are also widely used to generate an even broader class of models, including, for example, models for time series of counts such as the numbers of reported new cases of a particular disease (see e.g. [20]).

Additional Topics

Although linear models for time series data have found broad applications in many areas of the physical, biological, and engineering sciences, there are also many areas where they have been found inadequate. Consequently, a great deal of attention has been devoted to the development of nonlinear models for such applications. These include threshold models, [21]; bilinear models, [22]; random-coefficient autoregressive models, [23]; Markov switching models, [24]; and many others. For financial time series, the ARCH (autoregressive conditionally heteroscedastic) model of Engle [25], and its generalized version, the Generalized Autoregressive Conditionally Heteroscedastic (GARCH) model of Bollerslev [26], have been particularly successful in modeling financial-returns data of the type illustrated in Figure 4.

Although the sample autocorrelation function of the series shown in Figure 4 is compatible with the hypothesis that the series is an independent and identically distributed white noise sequence, the sample autocorrelation functions of the absolute values and the squares of the series are significantly different from zero, contradicting the hypothesis of independence. ARCH and GARCH processes are white noise sequences that are nevertheless dependent. The dependence is introduced in such a way that these models exhibit many of the distinctive features (heavy-tailed marginal distributions and persistence of volatility) that are observed in financial time series.

The recent explosion of interest in the modeling of financial data, in particular, with a view to solving option-pricing and asset allocation problems, has led to a great deal of interest, not only in ARCH and GARCH models, but also to stochastic volatility models and to models evolving continuously in time. For a recent account of time series models specifically related to financial applications, see [27].

Apart from their importance in finance, continuous-time models provide a very useful framework for the modeling and analysis of discrete-time series with irregularly spaced data or missing observations. For such applications, see [28].

References

- [1] Box, G.E.P. & Luceno.
- [2] Brockwell, P.J. & Davis, R.A. (2002). *Introduction to Time Series and Forecasting*, 2nd Edition, Springer-Verlag, New York.
- [3] Box, G.E.P., Jenkins, G.M. & Reinsel, G.C. (1994). *Time Series Analysis; Forecasting and Control*, 3rd Edition, Prentice-Hall, Englewood Cliffs.
- [4] Priestley, M.B. (1981). *Spectral Analysis and Time Series*, Academic Press, New York, Vols. 1 and 2.
- [5] Brockwell, P.J. & Davis, R.A. (1991). *Time Series: Theory and Methods*, 2nd Edition, Springer-Verlag, New York.
- [6] Hamilton, J.D. (1994). *Time Series Analysis*, Princeton University Press, Princeton.
- [7] Akaike, H. (1973). Information theory and an extension of the maximum likelihood principle, in *2nd International Symposium on Information Theory*, B.N. Petrov & F. Saki, eds, Akademia Kiado, Budapest, pp. 267–281.
- [8] Schweppe, F.C. (1965). Evaluation of likelihood functions for Gaussian signals, *IEEE Transactions on Information Theory* **IT-11**, 61–70.
- [9] Harvey, A.C. (1990). *Forecasting, Structural Time Series Models and the Kalman Filter*, Cambridge University Press, Cambridge.
- [10] Burnham, K.P. & Anderson, D. (2002). *Model Selection and Multi-Model Inference*, 2nd Edition, Springer-Verlag, New York.
- [11] Bloomfield, P. (1976). *Fourier Analysis of Time Series: An Introduction*, John Wiley & Sons, New York.
- [12] Nicholls, D.F. & Quinn, B.G. (1982). *Random Coefficient Autoregressive Models: An Introduction*, Springer Lecture Notes in Statistics, p. 11.
- [13] Tsay, R.S. (2001). *Analysis of Financial Time Series*, John Wiley & Sons, New York.
- [14] Whittle, P. (1963). On the fitting of multivariate autoregressions and the approximate canonical factorization of a spectral density matrix, *Biometrika* **40**, 129–134.
- [15] Findley, D.F., Monsell, B.C., Bell, W.R., Otto, M.C. & Chen, B.-C. (1998). New capabilities and methods of

- the X-12-ARIMA seasonal adjustment program (with discussion and reply), *The Journal of Business and Economic Statistics* **16**, 127–177.
- [16] Engle, R.F. & Granger, C.W.J. (1991). *Long-Run Economic Relationships*, *Advanced Texts in Econometrics*, Oxford University Press, Oxford.
- [17] Engle, R.F. & Granger, C.W.J. (1987). Co-integration and error correction: representation, estimation and testing, *Econometrica* **55**, 251–276.
- [18] Hannan, E.J. (1973). The asymptotic theory of linear time series models, *Journal of Applied Probability* **10**, 130–145.
- [19] Hannan, E.J. & Deistler, M. (1988). *The Statistical Theory of Linear Systems*, John Wiley & Sons, New York.
- [20] Chan, K.S. & Ledolter, J. (1995). Monte Carlo EM estimation for time series models involving counts, *The Journal of the American Statistical Association* **90**, 242–252.
- [21] Subba-Rao, T. & Gabr, M.M. (1984). *An Introduction to Bispectral Analysis and Bilinear Time Series Models*, Springer Lecture Notes in Statistics, p. 24.
- [22] Shibata, R. (1980). Asymptotically efficient selection of the order for estimating parameters of a linear process, *Annals of Statistics* **8**, 147–164.
- [23] Jones, R.H. (1980). Maximum likelihood fitting of ARMA models to time series with missing observations, *Technometrics* **22**, 389–395.
- [24] Granger, C.W.J. (1981). Some properties of time series data and their use in econometric model specification, *Journal of Econometrics* **16**, 121–130.
- [25] Engle, R. (1982). Autoregressive conditional heteroscedasticity with estimates of the variance of UK inflation, *Econometrica* **50**, 987–1087.
- [26] Bollerslev, T. (1986). Generalized autoregressive conditional heteroskedasticity, *Journal of Econometrics* **31**, 307–327.
- [27] Tong, H. (1990). *Non-linear Time Series: A Dynamical Systems Approach*, Oxford University Press, Oxford.
- [28] Hurvich, C.M. & Tsai, C.L. (1989). Regression and time series model selection in small samples, *Biometrika* **76**, 297–307.

Further Reading

Wiggins, R.A. & Robinson, E.A. (1965). Recursive solution to the multichannel filtering problem, *Journal of Geophysics Research* **70**, 1885–1891.

P.J. BROCKWELL

Article originally published in Encyclopedia of Statistics in Behavioral Science (2005, John Wiley & Sons, Ltd). Minor revisions for this publication by Jeroen de Mast.

Parametric Tests

Introduction

One way to look at the definition of a parametric test is to note that a parameter is a measure taken from a population. Thus, the mean and standard deviation of data for a whole population are both parameters. Parametric tests make certain assumptions about the nature of parameters, including the appropriate probability distribution, which can be used to decide whether the result of a statistical test would be significant. In addition, they can be used to find a **confidence interval** for a given parameter.

The types of tests given here have been categorized into those that compare means, those that compare variances, those that relate to **correlation** and regression, and those that deal with frequencies and proportions.

Comparing Means

z-Tests

z-tests compare a statistic (or single score) from a sample with the expected value of that statistic in the population under the null hypothesis. The expected value of the statistic under H_0 is subtracted from the observed value of the statistic and the result is divided by its standard deviation. When the sample statistic is based on more than a single score, then the standard deviation for that statistic is its standard error (*see* **Standard Error**). Thus, this type of test can only be used when the expected value of the statistic and its standard deviation are known. The probability of a *z*-value is found from the standardized **normal distribution**, which has a mean of zero and a standard deviation of one (*see* **Probability Density Functions**).

One-Group *z*-Test for a Single Score. In this version of the test, the equation is

$$z = \frac{x - \mu}{\sigma} \quad (1)$$

where

x is the single score
 μ is the mean for the scores in the population
 σ is the standard deviation of the population.

Example 1

Single score = 10

$$\mu = 5$$

$$\sigma = 2$$

$$z = \frac{10 - 5}{2} = 2.5.$$

Critical value for z at $\alpha = 0.05$ with a two-tailed test is 1.96.

Decision: reject H_0 .

One-Group *z*-Test for a Single Mean. This version of the test is a modification of the previous one because the distribution of means is the standard error of the mean: $\sigma/\sqrt{n}$, where n is the sample size.

The equation for this *z*-test is

$$z = \frac{m - \mu}{\left(\frac{\sigma}{\sqrt{n}}\right)} \quad (2)$$

where

m is the mean of the sample
 μ is the mean of the population
 σ is the standard deviation of the population
 n is the **sample size**.

Example 2

$$m = 5.5$$

$$\mu = 5$$

$$\sigma = 2$$

$$n = 20$$

$$z = \frac{5.5 - 5}{\left(\frac{2}{\sqrt{20}}\right)} = 1.12.$$

The critical value for z with $\alpha = 0.05$ and a one-tailed test is 1.64.

Decision: fail to reject H_0 .

2 Parametric Tests

t-Tests

t-tests form a family of tests that derive their probability from Student's *t* distribution (see **Probability Density Functions**). The shape of a particular *t* distribution is a function of the **degrees of freedom**. *t*-tests differ from *z*-tests in that they estimate the population standard deviation from the standard deviation(s) of the sample(s). Different versions of the *t*-test have different ways in which the degrees of freedom are calculated.

One-Sample *t*-Test. This test is used to compare a mean from a sample with that of a population or a hypothesized value from the population, when the standard deviation for the population is not known. The null hypothesis is that the sample is from the (hypothesized) population (that is, $\mu = \mu_h$, where μ is the mean of the population from which the sample came and μ_h is the mean of the population with which it is being compared).

The equation for this version of the *t*-test is

$$t = \frac{m - \mu_h}{\left(\frac{s}{\sqrt{n}}\right)} \quad (3)$$

where

- m is the mean of the sample
- μ_h is the mean or assumed mean for the population
- s is the standard deviation of the sample
- n is the sample size.

Degrees of freedom

In this version of the *t*-test, $df = n - 1$.

Example 3

$$\begin{aligned} m &= 9, \mu_h = 7 \\ s &= 3.2, n = 10 \\ df &= 9 \\ t_{(9)} &= 1.98. \end{aligned}$$

The critical *t* with $df = 9$ at $\alpha = 0.05$ for a two-tailed test is 2.26.

Decision: fail to reject H_0 .

Two-Samples *t*-Test. This test is used to compare the means of two different samples. The null hypothesis is $\mu_1 = \mu_2$ (i.e., $\mu_1 - \mu_2 = 0$, where μ_1 and μ_2

are the means of the two populations from which the samples come).

There are two versions of the equation for this test – one when the variances of the two populations are homogeneous (see **Homogeneity of Variances**) and the variances are pooled, and one when the variances are heterogeneous and are entered separately into the equation.

Homogeneous variance

$$t = \frac{m_1 - m_2}{\sqrt{pv \times \left(\frac{1}{n_1} + \frac{1}{n_2}\right)}} \quad (4)$$

where m_1 and m_2 are the sample means of groups 1 and 2 respectively, n_1 and n_2 are the sample sizes of the two groups, and pv is the **pooled variance**

$$pv = \frac{(n_1 - 1) \times s_1^2 + (n_2 - 1) \times s_2^2}{n_1 + n_2 - 2} \quad (5)$$

where s_1^2 & s_2^2 are the variances of groups 1 and 2. When the two group sizes are the same, this simplifies to

$$t = \frac{m_1 - m_2}{\sqrt{\frac{s_1^2 + s_2^2}{n}}} \quad (6)$$

where n is the sample size of each group.

Heterogeneous variances

$$t = \frac{m_1 - m_2}{\sqrt{\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}}} \quad (7)$$

Degrees of freedom

Homogeneous variance

$$df = n_1 + n_2 - 2 \quad (8)$$

Heterogeneous variance

$$df = \frac{\left(\frac{s_1^2}{n_1} + \frac{s_2^2}{n_2}\right)^2}{\frac{\left(\frac{s_1^2}{n_1}\right)^2}{n_1 - 1} + \frac{\left(\frac{s_2^2}{n_2}\right)^2}{n_2 - 1}} \quad (9)$$

Example 4 Groups with homogeneous variance

$$m_1 = 5.3, m_2 = 4.1$$

$$\begin{aligned}
n_1 &= n_2 = 20 \\
s_1^2 &= 1.3, s_2^2 = 1.5 \\
df &= 38 \\
t_{(38)} &= 3.21.
\end{aligned}$$

The critical t for a two-tailed probability with $df = 38$, at $\alpha = 0.05$ is 2.02.

Decision: reject H_0 .

Paired t -Test. This version of the t -test is used to compare two means that have been gained from the same sample, for example, quality of products before and after a corrective action is applied to them, or from two matched samples. For each product, a difference score is found between the two scores for that product (quality after the correction minus the quality before). The null hypothesis is that the mean of the difference scores is zero ($\mu_d = 0$, where μ_d is the mean of the difference scores in the population).

The equation for this test is

$$t = \frac{m_d}{\left(\frac{s_d}{\sqrt{n}}\right)} \quad (10)$$

where

- m_d is the mean of the difference scores
- s_d is the standard deviation of the difference scores
- n is the number of difference scores.

Degrees of freedom

In this version of the t -test, the df are $n - 1$.

Example 5

$$\begin{aligned}
m_1 &= 153.2, m_2 = 145.1, m_d = 8.1 \\
s_d &= 14.6, n = 20 \\
df &= 19 \\
t_{(19)} &= 2.48.
\end{aligned}$$

The critical t for $df = 19$ at $\alpha = 0.05$ for a one-tailed probability is 1.729.

Decision: reject H_0 .

ANOVA

Analysis of variance (ANOVA) (see **Analysis of Variance**) allows the comparison of the means of

more than two different conditions to be compared at the same time in a single omnibus test. As an example, researchers might wish to study differences in productivity of shifts. The null hypothesis, which is tested, is that the mean productivity is equal for all shifts. Thus, the null hypothesis would be $\mu_1 = \mu_2 = \mu_3$, where the μ_i denote the mean daily productivity of the morning, evening, and night shifts. ANOVA partitions the overall variance in a set of data into different parts.

The statistic that is created from an ANOVA is the F -ratio. It is the ratio of an estimate of the variance between groups and an estimate of the variance, which is not explicable in terms of differences between groups, that which is due to individual differences (sometimes referred to as *error*).

$$F = \frac{\text{variance between groups}}{\text{variance due to individual differences}} \quad (11)$$

If the null hypothesis is true, then these two variance estimates will both be due to individual differences and F will be close to 1. If the values from the different groups do differ, then F will tend to be larger than 1. The probability of an F -value is found from the F distribution (see **Probability Density Functions**). The value of F , which is statistically significant, depends on the degrees of freedom. In this test, there are two different degrees of freedom that determine the shape of the F distribution: the df for the variance between groups and the df for the error.

The variance estimates are usually termed the *mean squares (MS)*. These are formed by dividing a sum of squared deviations from a mean (usually referred to as the **sum of squares**) by the appropriate degrees of freedom.

The F -ratio is formed in different ways, depending on aspects of the design. In addition, the F -value will be calculated on a different basis if the independent variables (IVs) are fixed or random. The examples given here are for fixed IVs. For variations on the calculations, see [1]. The methods of calculation shown will be ones designed to explicate what the equation is doing rather than the computationally simplest version.

One-Way ANOVA

This version of the test partitions the total variance into two components: between groups (attributed to the IV) and within groups (the error).

4 Parametric Tests

Sums of squares

The sum of squares between the groups (SS_{bg}) is formed from

$$SS_{bg} = \sum [n_i \times (m_i - m)^2] \quad (12)$$

where

n_i is the sample size in group i

m_i is the mean of group i

m is the overall mean.

The sum of squares within the groups (SS_{wg}) is formed from

$$SS_{wg} = \sum \sum (x_{ij} - m_i)^2 \quad (13)$$

where

x_{ij} is the j th data point in group i

m_i is the mean of group i .

Degrees of freedom

The df for between groups is one fewer than the number of groups:

$$df_{bg} = k - 1 \quad (14)$$

where k is the number of groups.

The degrees of freedom for within groups is the total sample size minus the number of groups:

$$df_{wg} = N - k \quad (15)$$

where

N is the total sample size

k is the number of groups

Mean squares

The MS are formed by dividing the sum of squares by the appropriate degrees of freedom:

$$MS_{bg} = \frac{SS_{bg}}{df_{bg}} \quad (16)$$

$$MS_{wg} = \frac{SS_{wg}}{df_{wg}} \quad (17)$$

F-ratio

The F -ratio is formed by

$$F = \frac{MS_{bg}}{MS_{wg}} \quad (18)$$

Example 6 Three groups each with six observations are compared (Table 1).

$$\text{Overall mean}(m) = 4.278$$

$$SS_{bg} = \sum [6 \times (m_i - 4.278)^2] = 40.444$$

$$SS_{wg} = 19.167$$

$$df_{bg} = 3 - 1 = 2$$

$$df_{wg} = 18 - 3 = 15$$

$$MS_{bg} = \frac{40.444}{2} = 20.222$$

$$MS_{wg} = \frac{19.167}{15} = 1.278$$

$$F_{(2,15)} = \frac{20.222}{1.278} = 15.826.$$

The critical F -value for $\alpha = 0.05$, with df of 2 and 15 is 3.68.

Decision: reject H_0 .

Multiway ANOVA

When there is more than one IV, the way in which these variables work together can be investigated to see whether some act as moderators for others. An example of a design with two IVs would be if researchers wanted to test whether the effects of different types of music (jazz, classical, or pop) on blood pressure might vary, depending on the age of the listeners. The moderating effects of age on the effects of music might be indicated by an interaction between age and music type. In other words, the pattern of the link between blood pressure and music type differed between the two age groups. Therefore,

Table 1 Observations and group means in a one-way design

	Group		
	A	B	C
	2	5	7
	3	7	5
	3	5	6
	3	5	4
	1	5	4
	1	4	7
Mean (m_i)	2.167	5.167	5.500

an ANOVA with two IVs will have three F -ratios, each of which is testing a different null hypothesis. The first will ignore the presence of the second IV and test the main effect of the first IV, such that if there were two conditions in the first IV, then the null hypothesis would be $\mu_1 = \mu_2$, where μ_1 and μ_2 are the means in the two populations for the first IV. The second F -ratio would test the second null hypothesis, which would refer to the **main effect** of the second IV with the existence of the first being ignored. Thus, if there were two conditions in the second IV, then the second H_0 would be $\mu_a = \mu_b$ where μ_a and μ_b are the means in the population for the second IV. The third F -ratio would address the third H_0 , which would relate to the interaction between the two IVs. When each IV has two levels, the null hypothesis would be $\mu_{a1} - \mu_{a2} = \mu_{b1} - \mu_{b2}$, where the μ_{a1} denotes the mean for the combination of the first condition of the first IV and the first condition of the second IV.

Examples are only given here of ANOVAs with two IVs. For more complex designs, there will be higher-order interactions as well. When there are k IVs, there will be 2-, 3-, 4-, ..., k -way interactions, which can be tested. For details of such designs, see [1].

This version of ANOVA partitions the overall variance into four parts: the main effect of the first IV, the main effect of the second IV, the interaction between the two IVs, and the error term, which is used in all three F -ratios.

Sums of squares

The total sum of squares (SS_{Total}) is calculated from

$$SS_{\text{Total}} = \sum \sum \sum (x_{ijp} - m)^2 \quad (19)$$

where x_{ijp} is the p th observation where IV_1 is on level i , and IV_2 is on level j . A simpler description is that it is the sum of the squared deviations of each observation from the overall mean.

The sum of squares for the first IV (SS_A) is calculated from

$$SS_A = \sum [n_i \times (m_i - m)^2] \quad (20)$$

where

n_i is the number of observations having level i for IV_1

m_i is the mean of the observations having level i for IV_1
 m is the overall mean.

The sum of squares for the second IV (SS_B) is calculated from

$$SS_B = \sum [n_j \times (m_j - m)^2] \quad (21)$$

where

n_j is the number of observations having level j for IV_2
 m_j is the mean of the observations having level j for IV_2
 m is the overall mean.

The interaction sum of squares (SS_{AB}) can be found by finding the between-cells sum of squares ($SS_{\text{b.cells}}$), where a cell refers to the combination of levels in the two IVs: for example, first level of IV_1 and first level of IV_2 . $SS_{\text{b.cells}}$ is found from

$$SS_{\text{b.cells}} = \sum \sum [n_{ij} \times (m_{ij} - m)^2] \quad (22)$$

where

n_{ij} is the number of observations for which IV_1 is on level i and IV_2 is on level j ,
 m_{ij} is the mean of the observations for which IV_1 is on level i and IV_2 is on level j ,
 m is the overall mean.

$$SS_{AB} = SS_{\text{b.cells}} - (SS_A + SS_B) \quad (23)$$

The sum of squares for the residual (SS_{res}) can be found from

$$SS_{\text{res}} = SS_{\text{Total}} - (SS_A + SS_B + SS_{AB}) \quad (24)$$

Degrees of freedom

The total degrees of freedom (df_{Total}) are found from

$$df_{\text{Total}} = N - 1 \quad (25)$$

where N is the total sample size.

The degrees of freedom for each main effect (for example, df_A) are found from

$$df_A = k - 1 \quad (26)$$

where k is the number of levels in that IV.

6 Parametric Tests

The interaction degrees of freedom (df_{AB}) are found from

$$df_{AB} = df_A \times df_B \quad (27)$$

where df_A and df_B are the degrees of freedom of the two IVs.

The degrees of freedom for the residual (df_{res}) are calculated from

$$df_{res} = df_{Total} - (df_A + df_B + df_{AB}) \quad (28)$$

Mean squares

Each mean square is found by dividing the sum of squares by the appropriate df. For example, the mean square for the interaction (MS_{AB}) is found from

$$MS_{AB} = \frac{SS_{AB}}{df_{AB}} \quad (29)$$

F-ratios

Each F -ratio is found by dividing a given mean square by the mean square for the residual. For example, the F -ratio for the interaction is found from

$$F = \frac{MS_{AB}}{MS_{res}} \quad (30)$$

Example 7 Twenty observations are collected, equally distributed over the four combinations of two factors, each having two levels (Table 2). Overall mean = 8.2

Sums of squares

$$SS_{Total} = 139.2$$

$$SS_A = 51.2$$

$$SS_B = 5.0$$

$$SS_{b,cells} = 56.4$$

$$SS_{AB} = 56.4 - (51.2 + 5.0) = 0.2$$

$$SS_{res} = 139.2 - (51.2 + 5.0 + 0.2) = 82.8$$

Degrees of freedom

$$df_{Total} = 20 - 1 = 19$$

$$df_A = 2 - 1 = 1$$

$$df_B = 2 - 1 = 1$$

$$df_{AB} = 1 \times 1 = 1$$

$$df_{res} = 19 - (1 + 1 + 1) = 16$$

Table 2 Observations in a two-way design

IV ₁ (A)	1		2	
IV ₂ (B)	1	2	1	2
	9	7	6	4
	11	10	4	4
	10	14	9	10
	10	10	5	9
	7	10	6	9
Means	9.4	10.2	6	7.2

Mean squares

$$MS_A = \frac{51.2}{1} = 51.2$$

$$MS_B = \frac{5}{1} = 5$$

$$MS_{AB} = \frac{0.2}{1} = 0.2$$

$$MS_{res} = \frac{82.8}{16} = 5.175$$

F-ratios

$$F_{A(1,16)} = \frac{51.2}{5.175} = 9.89$$

The critical value for F at $\alpha = 0.05$ with df of 1 and 16 is 4.49.

Decision: reject H_0

$$F_{B(1,16)} = \frac{5}{5.175} = 0.97$$

Decision: fail to reject H_0

$$F_{AB(1,16)} = \frac{0.2}{5.175} = 0.04$$

Decision: fail to reject H_0 .

Extensions of ANOVA. ANOVA can be extended in a number of ways, including multiple factors, random and fixed factors, crossed and nested designs, and multivariate responses. See [1] for a general description.

Comparing Variances

F-Test for Difference between Variances

Two Independent Variances. This test compares two variances from different samples to see whether they are significantly different. An example of its use could be where we want to see whether the variability

in quality of one production line differs from the variability at another line.
The equation for the F -test is

$$F = \frac{s_1^2}{s_2^2} \quad (31)$$

where the sample variance in one sample is divided by the variance in the other sample.

If the research hypothesis is that one particular group will have the larger variance, then that should be treated as group 1 in this equation. As usual, an F -ratio close to 1 would suggest no difference in the variances of the two groups. A large F -ratio would suggest that group 1 has a larger variance than group 2, but it is worth noting that a particularly small F -ratio, and therefore a probability close to 1, would suggest that group 2 has the larger variance.

Degrees of freedom

The degrees of freedom for each variance are one fewer than the sample size in that group.

Example 8

Group 1
Variance: 16
Sample size: 100

Group 2
Variance: 11
Sample size: 150
 $F = \frac{16}{11} = 1.455$

Degrees of freedom

Group 1 $df = 100 - 1 = 99$; group 2 $df = 150 - 1 = 149$

The critical value of F with df of 99 and 149 for $\alpha = 0.05$ is 1.346.

Decision: reject H_0 .

k Independent Variances. The following procedure was devised by Bartlett [2] to test differences between the variances from more than two independent groups $i = 1, \dots, k$.

The finding of the statistic B (which can be tested with the χ^2 distribution) involves a number of stages. The first stage is to find an estimate of the overall variance (S^2). This is achieved by multiplying each sample variance by its related df , which is one fewer

than the size of that sample, summing the results and dividing that sum by the sum of all the df :

$$S^2 = \frac{\sum [(n_i - 1) \times s_i^2]}{N - k} \quad (32)$$

where

n_i is the sample size for the i th group
 s_i^2 is the sample variance in group i
 N is the total sample size
 k is the number of groups.

Next, we need to calculate a statistic known as C , using

$$C = 1 + \frac{1}{3 \times (k - 1)} \times \left[\sum \left(\frac{1}{n_i - 1} \right) - \frac{1}{N - k} \right] \quad (33)$$

We are now in a position to calculate B :

$$B = \frac{2.3026}{C} \times \left[(N - k) \times \log(S^2) - \sum \{ (n_i - 1) \times \log(s_i^2) \} \right] \quad (34)$$

where $\log$ means taking logarithm to the base 10.

Degrees of freedom

$$df = k - 1, \quad \text{where } k \text{ is the number of groups} \quad (35)$$

Kanji [3] cautions against using the chi-square distribution when the sample sizes are smaller than 6 and provides a table of critical values for a statistic derived from B when this is the case.

Example 9 We wish to compare the variances of three groups: 2.62, 3.66, and 2.49, with each group having the same sample size of 10.

$$\begin{aligned} N &= 30, k = 3, N - k = 27, C = 0.994 \\ S^2 &= 2.923, \log(s_1^2) = 0.418, \log(s_2^2) = 0.563, \\ \log(s_3^2) &= 0.396 \\ \log(S^2) &= 0.466, \\ (N - k) \times \log(S^2) &= 12.579 \\ \sum [(n_i - 1) \times \log(s_i^2)] &= 12.402 \\ B &= 0.4098 \\ df &= 3 - 1 = 2 \end{aligned}$$

8 Parametric Tests

The critical value of χ^2 for $\alpha = 0.05$ for $df = 2$ is 5.99.

Decision: fail to reject H_0 .

Correlation and Regression

***t*-Test for a Single Correlation Coefficient.** This test can be used to test the statistical significance of a correlation (*see Correlation*) between two variables. It makes the assumption under the null hypothesis that there is no correlation between these variables in the population. Therefore, the null hypothesis is $\rho = 0$, where ρ is the correlation in the population.

The equation for this test is

$$t = \frac{r \times \sqrt{n-2}}{\sqrt{1-r^2}} \quad (36)$$

where

r is the correlation in the sample
 n is the sample size

Degrees of freedom

In this version of the *t*-Test, $df = n - 2$, where n is the sample size.

Example 10

$$\begin{aligned} r &= 0.4, n = 15 \\ df &= 13 \\ t_{(13)} &= 1.57. \end{aligned}$$

Critical t for a two-tailed test with $\alpha = 0.05$ and $df = 13$ is 2.16.

Decision: fail to reject H_0 .

***t*-Test for Regression Coefficient.** This tests the size of an unstandardized regression coefficient. In the case of simple regression, where there is only one predictor variable, the null hypothesis is that the regression in the population is 0; that is, that the variance in the predictor variable does not account for any of the variance in the criterion variable. In multiple linear regression, the null hypothesis is that the predictor variable does not account for any variance in the criterion variable, which is not accounted for by the other predictor variables.

The version of the *t*-test is

$$t = \frac{b}{SE} \quad (37)$$

where b is the unstandardized regression coefficient
 SE is the standard error for the regression coefficient.
 In simple regression, the standard error is found from

$$SE = \sqrt{\frac{MS_{\text{res}}}{SS_p}} \quad (38)$$

where MS_{res} is the MS of the residual for the regression

SS_p is the sum of squares for the predictor variable.

For multiple regression, the standard error takes into account the interrelationship between the predictor variable for which the SE is being calculated and the other predictor variables in the regression (see [4]).

Degrees of freedom

The degrees of freedom for this version of the *t*-test are based on p (the number of predictor variables) and n (the sample size): $df = n - p - 1$.

Example 11 In a simple regression ($p = 1$), $b = 1.3$, and the standard error of the regression coefficient is 0.5

$$\begin{aligned} n &= 30 \\ df &= 28 \\ t_{(28)} &= 2.6 \end{aligned}$$

The critical value for t with $df = 28$, for a two-tailed test with $\alpha = 0.05$ is 2.048.

Decision: reject H_0 .

In multiple regression, it can be argued that a correction to α should be made to allow for multiple testing. This could be achieved by using a Bonferroni adjustment.

***F*-Test for a Single R^2 .** This is the equivalent of a one-way ANOVA. In this case, the overall variance within the variable to be predicted (the response variable, for example, moisture percentage) is separated into two sources: that which can be explained by the relationship between the predictor variable(s) (such as various

machine settings) and the response variable, and that which cannot be explained by this relationship, the residual.

The equation for this F -test is

$$F = \frac{(N - p - 1) \times R^2}{(1 - R^2) \times p} \quad (39)$$

where N is the sample size

p is the number of predictor variables

R^2 is the squared multiple correlation coefficient (see **Coefficient of Determination (R^2)**).

Degrees of freedom

The regression degrees of freedom are p (the number of predictor variables). The residual degrees of freedom are $N - p - 1$.

Example 12 Number of predictor variables: 3

Sample size: 336

Regression df : 3

Residual df : 332

$$R^2 = 0.03343$$

$$F_{(3,332)} = \frac{(336 - 3 - 1) \times 0.03343}{(1 - 0.03343) \times 3} = 3.83$$

Critical F -value with df of 3 and 332 for $\alpha = 0.05$ is 2.63.

Decision: reject H_0 .

F -Test for Comparison of Two R^2 . This tests whether the addition of predictor variables to an existing regression model adds significantly to the amount of variance that the model explains. If only one variable is being added to an existing model, then the information about whether it adds significantly is already supplied by the t -test for the regression coefficient of the newly added variable.

The equation for this F -test is

$$F = \frac{(N - p_1 - 1) \times (R_1^2 - R_2^2)}{(p_1 - p_2) \times (1 - R_1^2)} \quad (40)$$

where N is the sample size.

The subscript 1 refers to the regression with more predictor variables and subscript 2, the regression with fewer predictor variables.

p is the number of predictor variables.

R^2 is the squared multiple correlation coefficient from a regression.

Degrees of freedom

The regression $df = p_1 - p_2$, while the residual $df = N - p_1 - 1$.

Example 13 $R_1^2 = 0.047504$

Number of predictor variables (p_1): 5

$R_2^2 = 0.03343$

Number of predictor variables (p_2): 3

Total sample size: 336

df for regression = $5 - 3 = 2$

df for residual = $336 - 5 - 1 = 330$

$F_{(2,330)} = 2.238$.

The critical value of F with df of 2 and 330 for $\alpha = 0.05$ is 3.023.

Decision: fail to reject H_0 .

Frequencies and Proportions

Chi-Square

χ^2 is a distribution against which the results of a number of statistical tests are compared. The two most frequently used tests that use this distribution are themselves called χ^2 tests and are used when the data take the form of frequencies. Both types involve the comparison of frequencies that have been found (observed) to fall into particular categories with the frequencies that could be expected if the null hypothesis were correct. The categories have to be mutually exclusive; that is, a case cannot appear in more than one category.

The first type of test involves comparing the frequencies in each category for a single variable, for example, the number of smokers and nonsmokers in a sample. It has two variants, which have a different way of viewing the null **hypothesis testing** process. One version is like the example given above, and might have the null hypothesis that the number of smokers in the population is equal to the number of nonsmokers. However, the null hypothesis does not have to be that the frequencies are equal; it could be that they divide the population into particular proportions, for example, 0.4 of the population are smokers and 0.6 are nonsmokers.

The second way in which this test can be used is to test whether a set of data is distributed

in a particular way, for example, that they form a normal distribution. Here, the expected proportions in different intervals are derived from the proportions of a normal curve, with the observed mean and standard deviation, that would lie in each interval, and H_0 is that the data are normally distributed.

The second most frequent χ^2 test is for contingency tables where two variables are involved, for example, the number of smokers and nonsmokers in two different socioeconomic groups.

All the tests described here are calculated in the following way:

$$\chi^2 = \sum \frac{(f_o - f_e)^2}{f_e} \quad (41)$$

where f_o and f_e are the observed and expected frequencies, respectively.

The degrees of freedom in these tests are based on the number of categories and not on the sample size.

One assumption of this test is over the size of the expected frequencies. When the degrees of freedom are 1, the assumption is that all the expected frequencies will be at least 5. When the df is greater than 1, the assumption is that at least 20% of the expected frequencies will be 5.

Yates [5] devised a correction for χ^2 when the degrees of freedom are 1 to allow for the fact that the χ^2 distribution is continuous and yet when $df = 1$, there are so few categories that the χ^2 values from such a test will be far from continuous; hence, the Yates test is referred to as a *correction for continuity*. However, it is considered that this variant on the χ^2 test is only appropriate when the marginal totals are fixed, that is, that they have been chosen in advance [6]. In most uses of χ^2 , this would not be true. If we were looking at gender and smoking status, it would make little sense to set, in advance, how many males and females you were going to sample, as well as how many smokers and nonsmokers.

χ^2 corrected for continuity is found from

$$\chi^2_{(1)} = \sum \frac{(|f_o - f_e| - 0.5)^2}{f_e} \quad (42)$$

where $|f_o - f_e|$ means taking the absolute value, in other words, if the result is negative, treat it as positive.

One-Group Chi-Square/Goodness of Fit

Equal Proportions. In this version of the test, the observed frequencies that occur in each category of a single variable are compared with the expected frequencies, which are such that the same proportion will occur in each category.

Degrees of freedom

The df are based on the number of categories (k); $df = k - 1$.

Example 14 A sample of 45 products are placed into three categories, with 25 in category A, 15 in B, and 5 in C.

The expected frequencies (given the null hypothesis of equal proportions) are calculated by dividing the total sample by the number of categories. Therefore, each category would be expected to have 15 people in it.

$$\begin{aligned} \chi^2 &= 13.33 \\ df &= 2. \end{aligned}$$

The critical value of the χ^2 distribution with $df = 2$ and $\alpha = 0.05$ is 5.99.

Decision: reject H_0 .

Other procedures for rates and proportions are discussed in **Statistical Methods for Counts, Rates, and Proportions**.

Test of Distribution. This test is another use of the previous test, but the example will show how it is possible to test whether a set of data is distributed according to a particular pattern. The distribution could be uniform, as in the previous example, or nonuniform.

Example 15 One hundred scores have been obtained with a mean of 1.67 and a standard deviation of 0.51. In order to test whether the distribution of the scores deviates from being normally distributed, the scores can be converted into z -scores by subtracting the mean from each and dividing the result by the standard deviation. The z -scores can be put into ranges. Given the sample size and the need to maintain at least 80% of the expected frequencies at 5 or more, the width of the ranges can be approximately half a standard deviation except for the two outer ranges, where the expected frequencies get smaller, the further they go from the mean. At the bottom of the range, as the lowest possible

score is 0, the equivalent z -score will be -3.27 . At the top end of the range, there is no limit set on the scale.

By referring to standard tables of probabilities for a normal distribution, we can find out what the expected frequency would be within a given range of z -scores. Table 3 shows the expected and observed frequencies in each range.

$$\chi^2 = 3.56$$

$$df = 8 - 1 = 7.$$

The critical value for a χ^2 distribution with $df = 7$ at $\alpha = 0.05$ is 14.07.

Decision: fail to reject H_0 .

See **Normality Tests** for other tests on normality.

Chi-Square Contingency Test

This version of the χ^2 test investigates the way in which the frequencies in the levels of one variable differ across the other variable. Once again it is for categorical data. An example could be a sample of blind people and a sample of sighted people; both groups are aged over 80 years. Each person is asked whether they go out of their house in a normal day. Therefore, we have the variable visual condition with the levels blind and sighted, and another variable whether the person goes out with the levels yes and no. The null hypothesis of this test would be that the proportions of sighted and blind people going out would be the same (which is the same as saying that the proportions staying in would be the same in each group). This can be rephrased to say that the two variables are independent of each other: the likelihood

of a person going out is not linked to that person's visual condition.

The expected frequencies are based on the marginal probabilities. In this example, that would be the number of blind people, the number of sighted people, the number of the people in the whole sample who go out, and the number of people in the whole sample who do not go out. Thus, if 25% of the entire sample went out, then the expected frequencies would be based on 25% of each group going out and, therefore, 75% of each not going out.

The degrees of freedom for this version of the test are calculated from the number of rows and columns in the contingency table: $df = (r - 1) \times (c - 1)$, where r is the number of rows and c , the number of columns in the table.

Example 16 Two variables A and B each have two levels. Level 1 of variable A has 27 people and 13 are in level 2. Level 1 of variable B has 21 people and 19 are in level 2 (Table 4).

If variables A and B are independent, then the expected frequency for the number who are in the first level of both variables will be based on the fact that 27 out of 40 (or 0.675) were in level 1 of variable A and 21 out of 40 (or 0.525) were in level 1 of variable B . Therefore, the proportion who would be expected to be in level 1 of both variables would be $0.675 \times 0.525 = 0.354375$, and the expected frequency would be $0.354375 \times 40 = 14.175$.

$$\chi^2 = 2.16$$

$$df = (2 - 1) \times (2 - 1) = 1.$$

The critical value of χ^2 with $df = 1$ and $\alpha = 0.05$ is 3.84.

Decision: fail to reject H_0 .

Table 3 The expected (under the assumption of normal distribution) and observed frequencies of a sample of 100 values

From z	To z	f_e	f_o
-3.275	-1.500	6.620	9
-1.499	-1.000	9.172	7
-0.999	-0.500	14.964	15
-0.499	0.000	19.111	22
0.001	0.500	19.106	14
0.501	1.000	14.953	15
1.001	1.500	9.161	11
1.501	5.000	6.661	7

Table 4 A contingency table showing the cell and marginal frequencies of 40 participants

		Variable A		
		1	2	Total
Variable B	1	12	9	21
	2	15	4	19
	Total	27	13	40

z-Test for Proportions

Comparison between a Sample and a Population Proportion. This test compares the proportion in a sample with that in a population (or that which might be assumed to exist in a population).

The standard error for this test is $\sqrt{\frac{\pi \times (1 - \pi)}{n}}$ where

π is the proportion in the population
 n is the sample size.

The equation for the z -test is

$$\frac{p - \pi}{\sqrt{\frac{\pi \times (1 - \pi)}{n}}} \quad (43)$$

where p is the proportion in the sample.

Example 17 In a sample of 25, the proportion to be tested is 0.7

Under the null hypothesis, the proportion in the population is 0.5

$$z = \frac{0.7 - 0.5}{\sqrt{\frac{0.5 \times (1 - 0.5)}{25}}} = 2.$$

The critical value for z with $\alpha = 0.05$ for a two-tailed test is 1.96.

Decision: reject H_0 .

Comparison of Two Independent Samples. Given two populations, the null hypothesis to be tested is that the proportion having a certain characteristic is equal in both populations. Given a sample from each population, the standard error for this test is

$$\sqrt{\frac{\pi_1 \times (1 - \pi_1)}{n_1} + \frac{\pi_2 \times (1 - \pi_2)}{n_2}} \quad (44)$$

where

π_1 and π_2 are the proportions in each population
 n_1 and n_2 are the sizes of the two samples.

When the population proportions are not known, they are estimated from the sample proportions.

The equation for the z -test is

$$z = \frac{p_1 - p_2}{\sqrt{\frac{p_1 \times (1 - p_1)}{n_1} + \frac{p_2 \times (1 - p_2)}{n_2}}} \quad (45)$$

where p_1 and p_2 are the proportions in the two samples.

Example 18

Sample 1: $n = 30$, $p = 0.7$

Sample 2: $n = 25$, $p = 0.6$

$$z = \frac{0.7 - 0.6}{\sqrt{\frac{0.7 \times (1 - 0.7)}{30} + \frac{0.6 \times (1 - 0.6)}{25}}} = 0.776.$$

Critical value for z at $\alpha = 0.05$ with a two-tailed test is 1.96.

Decision: fail to reject H_0 .

Comparison of Two Correlated Samples. The main use for this test is to judge whether there has been change across two occasions when a measure was taken from a sample. For example, researchers might be interested in whether people's attitudes to a ban on smoking in public places had changed after seeing a video on the dangers of passive smoking compared with attitudes held before seeing the video. A complication with this version of the test is over the estimate of the standard error. A number of versions exist, which produce slightly different results. However, one version will be presented here from Agresti [7]. This will be followed by a simplification, which is found in a commonly used test.

Table 5 shows that originally 40 people were in category A and 44 in category B, and on a second occasion, this had changed to 50 people

Table 5 The frequencies of 84 participants placed in two categories and noted at two different times

		After		
		A	B	Total
Before	A	15	25	40
	B	35	9	44
	Total	50	34	84

being in category A and 34 in category B. This test is only interested in the cells where change has occurred: the 25 who were in A originally but changed to B and the 35 who were in B originally but changed to A. By converting each of these to proportions of the entire sample, $25/84 = 0.297619$ and $35/84 = 0.416667$, we have the two proportions we wish to compare. The standard error for the test is

$$\sqrt{\frac{(p_1 + p_2) - (p_1 - p_2)^2}{n}}$$

where n is the total sample size.

The equation for the z -test is

$$z = \frac{p_1 - p_2}{\sqrt{\frac{(p_1 + p_2) - (p_1 - p_2)^2}{n}}} \quad (46)$$

Example 19 Using the data in Table 5,

$$p_1 = 0.297619$$

$$p_2 = 0.416667$$

$$n = 84$$

$$z = \frac{0.297619 - 0.416667}{\sqrt{\frac{(0.297619 + 0.416667) - (0.297619 - 0.416667)^2}{84}}}$$

$$= -1.304$$

The critical z with $\alpha = 0.05$ for a two-tailed test is -1.96 .

Decision: fail to reject H_0 .

When the z from this version of the test is squared, this produces the Wald test statistic.

A simplified version of the standard error allows another test to be derived from the resulting z -test: McNemar's test of change.

In this version, the equation for the z -test is

$$z = \frac{p_1 - p_2}{\sqrt{\frac{(p_1 + p_2)}{n}}} \quad (47)$$

Example 20 Once again using the same data as that in the previous example,

$$z = -1.291.$$

If this z -value is squared, then the statistic is McNemar's test of change, which is often presented in a further simplified version of the calculations, which produces the same result.

References

- [1] Montgomery, D.C. (1997). *Design and Analysis of Experiments*, 4th Edition, John Wiley & Sons.
- [2] Bartlett, M.S. (1937). Some examples of statistical methods of research in agriculture and applied biology, *Supplement to the Journal of the Royal Statistical Society* **4**, 137–170.
- [3] Kanji, G.K. (1993). *100 Statistical Tests*, Sage Publications, London.
- [4] Draper, N.R. & Smith, H. (1998). *Applied Regression Analysis*, 3rd Edition, John Wiley & Sons.
- [5] Yates, F. (1934). Contingency tables involving small numbers and the χ^2 test, *Supplement to the Journal of the Royal Statistical Society* **1**, 217–235.
- [6] Neave, H.R. & Worthington, P.L. (1988). *Distribution-Free Tests*, Routledge, London.
- [7] Agresti, A. (2002). *Categorical Data Analysis*, 2nd Edition, John Wiley & Sons, Hoboken.

DAVID CLARK-CARTER

Article originally published in *Encyclopedia of Statistics in Behavioral Science* (2005, John Wiley & Sons, Ltd). Minor revisions for this publication by Jeroen de Mast.

Box and Cox Transformation

The usual assumptions when one analyzes data are the standard assumptions of the linear model (*see Analysis of Variance*), i.e., the additivity of effects, the constancy of variance, the normality of the observations, and the independence of observations. If these assumptions are not met by the data, Tukey [1] suggested two alternatives: either a new analysis must be devised to meet the assumptions, or the data must be transformed to meet the usual assumptions. If a satisfactory transformation is found, it is almost always easier to use than to develop a new method of analysis.

Tukey suggested a family of transformations with an unknown power parameter and Box and Cox [2] modified it to

$$y^{(\lambda)} = \begin{cases} (y^{(\lambda)} - 1)/\lambda & \text{for } \lambda \neq 0 \\ \log y, & \text{for } \lambda = 0 \end{cases} \quad (1)$$

where y is an original observation, $y^{(\lambda)}$ the “new” observation, and λ a real unknown parameter.

Box and Cox assumed that for some λ the n transformed observations

$$\mathbf{y}^{(\lambda)} = \begin{bmatrix} y_1^{(\lambda)} \\ y_2^{(\lambda)} \\ \vdots \\ y_n^{(\lambda)} \end{bmatrix} \quad (2)$$

are independent and normally distributed with constant variance σ^2 and mean vector

$$E\mathbf{y}^{(\lambda)} = A\boldsymbol{\theta} \quad (3)$$

where A is a known $n \times p$ matrix and $\boldsymbol{\theta}$ is a vector of parameters associated with the transformed observations.

Box and Cox [2] estimate the parameters by *maximum likelihood* (*see Maximum Likelihood*) as follows. First, given λ ,

$$\hat{\boldsymbol{\theta}}(\lambda) = (A'A)^{-1}A'\mathbf{y}^{(\lambda)} \quad (4)$$

and

$$\hat{\sigma}^2(\lambda) = \mathbf{y}^{(\lambda)'}[I - A(A'A)^{-1}A']\mathbf{y}^{(\lambda)} \quad (5)$$

are the estimators of $\boldsymbol{\theta}$ and σ^2 , respectively. Second, λ is estimated by maximizing the log likelihood function

$$l(\lambda) = -n \log \hat{\sigma}^2(\lambda)/2 + \log J(\lambda : y) \quad (6)$$

where

$$J(\lambda : y) = \prod_{i=1}^n y_i^{\lambda-1} \quad (7)$$

The maximum-likelihood estimator of λ , say $\hat{\lambda}$, is then substituted into equation (5), which determines the estimators of the other parameters.

Since the original Box–Cox article, there have been many related papers. For example, Draper and Cox [3] derive the precision of the maximum-likelihood estimator of λ for a simple random sample, i.e., $Ey_i^{(\lambda)} = \mu$ for $i = 1, 2, \dots, n$. They found that the approximate variance is

$$V(\hat{\lambda}) = \frac{2}{3n\Delta^2} \left(1 - \frac{1}{3}\gamma_1^2 + \frac{7}{18}\gamma_2 \right)^{-1} \quad (8)$$

where

$$\gamma_1 = \mu_3/\sigma^3 \quad (9)$$

$$\gamma_2 = \mu_4/\sigma^4 - 3 \quad (10)$$

$$\Delta = \lambda\sigma/(1 + \lambda\mu) \quad (11)$$

σ^2 is the variance, and μ_i is the i th central moment of the original observations.

Hinkley [4] generalized the Box–Cox transformation to include power transformations to symmetric distributions and showed his “quick” estimate of λ to be consistent and have a limiting **normal distribution**.

Literature. There are many references related to the Box and Cox transformation. Poirier [5] has extended the Box–Cox work to truncated normal distributions. Some recent results related to large-sample behavior are given by Hernández and Johnson [6]. The Box and Cox [2], Draper and Cox [3], Hinkley [4], and Poirier [5] articles are quite informative but highly technical.

At the textbook level, Chapter 10 of Box and Tiao [7] and Chapter VI of Zellner [8] are excellent introductions to the Box–Cox theory of transformations.

References

- [1] Tukey, J.W. (1957). On the comparative anatomy of transformations, *Annals of Mathematical Statistics* **28**, 602–632.
- [2] Box, G.E.P. & Cox, D.R. (1964). An analysis of transformations (with Discussion), *Journal of the Royal Statistical Society. Series B* **26**, 211–252.
- [3] Draper, N.R. & Cox, D.R. (1969). On distributions and their transformation to normality, *Journal of the Royal Statistical Society. Series B* **31**, 472–476.
- [4] Hinkley, D.V. (1975). On power transformations to symmetry, *Biometrika* **62**, 101–111.
- [5] Poirier, D.J. (1978). The use of the Box-Cox transformation in limited dependent variable models, *Journal of the American Statistical Association* **73**, 284–287.
- [6] Harnández, F. & Johnson, R.A. (1979). Technical Report No. 545, Department of Statistics, University of Wisconsin, Madison.
- [7] Box, G.E.P. & Tiao, G.C. (1973). *Bayesian Inference in Statistical Analysis*, Addison-Wesley, Reading.
- [8] Zellner, A. (1971). *An Introduction to Bayesian Inference in Econometrics*, John Wiley & Sons, New York.

LYLE D. BROEMELING

Article originally published in Encyclopedia of Statistical Sciences, 2nd Edition (2005, John Wiley & Sons, Inc.). Minor revisions for this publication by Jeroen de Mast.

Cumulative Distribution Function (CDF)

The probability that a random variable X does not exceed a value x , regarded as a function of x , is called the *cumulative distribution function of x* . A common notation is

$$F_X(x) = \Pr[X \leq x] \quad (1)$$

See also **Probability Theory**. From the definition, it follows that $0 \leq F_X(x) \leq 1$, and $F_X(x)$ is a nondecreasing function of x .

The notation $F_X(x)$ is intended to indicate that the (mathematical) function of x represents properties of the random variable X . For continuous distributions, the derivative of F_X is the *probability density function* (see **Probability Density Function (PDF)**).

If $F_X(x) \rightarrow 0$ as $x \rightarrow -\infty$ and $F_X(x) \rightarrow 1$ as $x \rightarrow \infty$, the distribution is called *proper*; otherwise, it is called *improper*.

Data Collection

A distinction needs to be made between numbers, adjectives, and other forms of description and data. Numbers and adjectives can be available simply as entities in themselves or as data. The existence of numbers and adjectives does not imply the existence of a set of data, but the existence of a set of data implies the existence of some descriptive forms such as numbers, adjectives, phrases, pictures, graphs, etc. For example, the set of numbers {3, 1, 0, 4, 9, 6} and the set of adjectives {small, pretty, personable, harsh, susceptible, resistant} do not in themselves convey information about any phenomenon. When numbers and adjectives convey information about some entity such as 0, 1, 3, and 4 worms in four particular apples and small, pretty, heat-resistant spores of a bacterium, these facts are denoted as *data*. A *datum* is defined to be a fact (numerical or otherwise) from which a conclusion may be drawn such as, none of the four apples have the same number of worms. A datum contains information whereas a number, adjective, or other form of description may not.

The information in a set of data may be contaminated (mixed up or confounded) with other kinds of information. For example, the particular four apples may have been the only ones in a basket of apples that contained worms and the apples were from a part of the orchard that had a very light infestation of the codling moth. Thus we do not know if 4 out of 120 wormy apples per basket represents the proportion of wormy apples in the entire population of apples or only for this variety in lightly infested orchards. The various factors affecting variation and response are denoted as *sources of variation* in the data.

When the various effects in a set of data are inseparable, they are said to be completely confounded. Partial confounding of effects occurs when the effects are partly mixed together, but it is possible to obtain estimates of each of the various sources of variation. No confounding of effects (orthogonality) implies maximum separability of sources of variation. For example, suppose that there is technician-to-technician variation, day-to-day variation, and method-to-method variation. If technician I conducts method 1 on day 1, technician II method 2 on day 2, technician III method 3 on day 3, etc., one cannot separate the technician effect from the method

effect or from the day effect. If, on the other hand, all technicians used all methods on each of several days, the different effects would be completely separable. The preceding plan would be an orthogonal (completely unconfounded) one, whereas the former would be completely confounded.

Thus the plan or the design of the investigation has a large influence on the confounding aspects in data and hence on the amount of information derivable from the data. With adequate planning and foresight, data can be obtained that have little or no confounding of sources of variation of interest to the investigator. On the other hand, unplanned and/or haphazard collection of data, no matter how voluminous, most frequently results in highly confounded data sets which may be of little or no use in studying effects of various sources of variation. There is a tendency for people to believe that largeness of a data set removes the effects of confounding. This is not necessarily true. The closer one gets to obtaining responses for 100% of the population, the nearer one is to the values of the population parameters (i.e., *see Standard Error* of estimates will be smaller). In this sense, largeness does remove the difficulties of haphazard or selective sampling. However, the complete enumeration of a subpopulation provides no information on the other subpopulation parameters. In this sense, largeness of a set of data does nothing to remove the biases.

Aspects of Data Collection

Whenever data are to be collected, several aspects require consideration. The first one to consider is *why* these data should be collected and what uses these data will serve. If there is no purpose or no use for these data, why collect them? The ready accessibility of computers, together with recording equipment, makes it a simple matter to record voluminous sets of data with relatively little effort. As an illustration, a number of instruments for recording temperature, pulse rate, blood pressure, changes in blood pressure, etc., are attached to a dog undergoing surgery. A minicomputer is used to determine the frequency with which temperature, pulse rate, blood pressure measurements, etc., are recorded. As a second example, suppose that temperature in 100ths of a degree Celsius and humidity in 100ths of a percent are to be recorded every second over a 20-year period in 1000 locations of a field; a

2 Data Collection

minicomputer is used to determine the frequency with which observations are recorded and the observations may be recorded on tape and/or paper. This results in $60 \times 60 \times 24 \times 365.25 \times 1000 \times 20 = 631\,150\,000\,000$ measurements on temperature and the same number on humidity. The data set would consist of more than 1.25 trillion observations. Before embarking on such a large **data collection** venture, one should seriously consider taking fewer measurements (e.g., one every hour rather than 3600 per hour), the uses for data of this nature, and which individuals would actually use these data. Unfortunately, large data sets have been and are being collected, merely on the premise that they might be useful to someone at some time. A computer is available, and data are collected to use the capabilities of the computer.

A second aspect of data collection is *what* data will be collected. All pertinent data for studying a phenomenon should be collected. Thus it is essential to determine what data to collect in light of why the data are collected. Sufficient data should be collected to achieve the goals and purposes of the investigation. Provision should also be made to obtain pertinent data that surface during the course of an investigation. The investigator must be aware of evidence that comes to his or her attention, and obtain the necessary data to explain the phenomenon. An example was given above illustrating how a large data set might be obtained. In some investigations, such as sending missiles to the moon, one can obtain only a few observations on selected entities. In both these cases, it is essential to determine what data to collect.

A third aspect of data collection is *how* and *where* the data are to be collected. The statistician can be very helpful in designing and planning the investigation and in determining how and where the data are to be collected. The how and where of data collection are intimately entwined with the plan and type of the investigation. Also involved are how to measure and quantify information on the various phenomena in an investigation. Measuring methods and gauges must be carefully considered, and it should be ascertained if a measuring instrument is measuring what it is supposed to, and with adequate precision.

Two other aspects to be considered are – *who* is to collect the data and *when* are they to be obtained. Unless these aspects are firmly regulated, the data

may not be collected or only a part of the data may be obtained. In many investigations it is imperative to have unbiased and highly trained technicians to carry out the investigation. The timing of an investigation may be important to its success, and/or it may be necessary to carry out the investigation in a specified time period. If these aspects are ignored, the data may be so unreliable as to make them useless.

In addition to the *why, what, how, where, who, and when* aspects of data collection, it is imperative to have a complete, written description of all data obtained. If this information is recorded only in the investigator's mind, it may soon be lost and be unavailable to other investigators desiring to use the data. Plans need to be made for storage and/or disposal of all data collected. Discarding valuable data that are needed in future research or the storage of useless data result in economic losses that can be avoided with proper planning. Data must be adequately organized and stored on storage devices.

Types of Data Collection

The main types of data collection are historical data, *observational studies* (see **Observational Studies**, including **Internal and External Quality Measures**), and experimental investigations. Historical data (or *happenstance data*, as Box, Hunter, and Hunter call such data sets) are data that have been collected for another purpose, and the investigator has no control over the way in which they are collected. In experiments, the investigator manipulates the system he or she samples data from, thus controlling the sources of variation in the system (varying the factors of interest according to a plan, while keeping other factors as constant as possible). Observational studies are in between in terms of control: the investigator can tailor the data collection design to the purpose he or she has in mind (controlling what data to collect where and when and by whom). The sources of variation are not controlled, however, but the investigator has to work with the values they happen to have.

These different levels of control have their impact on the analysis of the data. The usefulness of historical data is often quite limited. Although the data set can be quite large, unless the observations were guided by an appropriate sampling scheme, there is no reason why they should be representative. Further, strong confounding among sources of variation, or

the omission of factors that have a substantial effect may make it difficult to draw valid conclusions about the research questions. When working with historical data, the investigator must always critically assess whether the data set is usable for his purpose, or that it may be better to discard it.

For observational studies and sample surveys, the investigator draws a sample from the process or population, aiming for this sample to be representative and thus allow an extrapolation of the conclusions to the population (in the case of a *census* one aims for a sample that covers 100% of the population). Some types of sampling methods are

1. Simple random sample

The population does not contain subpopulations and each sampling unit in the population has an equal and independent chance of being selected. An equivalent definition is that every possible sample of size n has an equal chance of being selected.

2. Stratified simple random sample

The population is composed of subpopulations and a simple random sample of sampling units is selected within *each* of the subpopulations (strata).

3. Cluster simple random sample

The population is composed of subpopulations (clusters). A simple random sample of subpopulations (clusters) is made; then a simple random sample is made within each of the selected clusters. When the clusters are areas, the sample survey design is known as an *area-simple random sample*.

4. Every k th sampling unit with a random start

In many cases, a list of sampling units or a serial ordering of sampling units is available. The list is partitioned in subgroups of k items each. A number between 1 and k , say t , is randomly selected. Then the t th item, the $(t + k)$ th item, the $(t + 2k)$ th item, the $(t + 3k)$ th item, etc., form the sample. An example is a sample consisting of products taken from the production line at

9, 9.30, 10, 10.30 a.m., and so on. This survey design is called *systematic* or *cyclic*.

5. More complex designs

There are many types of sample survey designs involving more complexity than the above. Several of the types described above may be included in these designs.

Experimental design is built on the principles laid down by Fisher: **randomization**, **blocking**, and **factorial** designs (*see Assessment of Experimental Designs*). Especially for making inferences about causal relations (*see Causality*), experimental data are generally considered the most powerful. Because of the large extent of control, experimental data are less subject to complications such as confounding.

Further Reading

- Campbell, S.K. (1974). *Flaws and Fallacies in Statistical Thinking*, Prentice-Hall, Englewood Cliffs, (A humorous and enlightening account of many misconceptions about and of statistics).
- Cochran, W.G. (1977). *Sampling Techniques*, 3rd Edition, John Wiley & Sons, New York, (A lucid and clear account of sample survey techniques, theory, and application is presented. The text is an introductory book in sampling with prerequisites of a course in statistical methods and some mathematical statistics).
- Deming, W.E. (1950). *Some Theory of Sampling*, John Wiley & Sons, New York.
- Fisher, R.A. (1966). *Design of Experiments*, 8th Edition, Hafner (Macmillan), (Fisher's classical and authoritative book on the principles of experimental design).
- Haack, D.G. (1979). *Statistical Literacy. A Guide to Interpretation*, Duxbury Press, North Scituate, (An elementary and introductory book that deals with statistical doublespeak, which is defined as the "inflated, involved, and often deliberately ambiguous use of numbers" in many areas of our life).

WALTER T. FEDERER

Article originally published in Encyclopedia of Statistical Sciences, 2nd Edition (2005, John Wiley & Sons, Inc.). Minor revisions for this publication by Jeroen de Mast.

Degrees of Freedom

The term *degrees of freedom* is used in a number of different fields of science and has a variety of meanings. In statistics, it is most commonly thought of as a parameter (ν) in the χ^2 distribution (denoted χ^2_ν) and in distributions related thereto, such as central and noncentral t and F (see **Probability Density Functions**).

The term *degrees of freedom* for this parameter and the rationale for it go back to R. A. Fisher and his correspondence in 1912 with W. S. Gosset ("Student"). Fisher realized that based on n independent normally distributed sample values $x_1, \dots, x_n$, with sample mean $\bar{x}$, the denominator of the t statistic involves dividing $\sum_{i=1}^n (x_i - \bar{x})^2$ by $(n - 1)$, not by n as Gosset had done and as one would if the true mean μ were known and the sum of squared deviations would be $\sum_{i=1}^n (x_i - \mu)^2$ (see the discussion in [1, pp. 72–73]). In the 1800s Bienaymé and Helmert had proved that $\sum (x_i - \mu)^2 / \sigma^2$ and $\sum (x_i - \bar{x})^2 / \sigma^2$ have χ^2_n and χ^2_{n-1} distributions, respectively, σ^2 being the true variance. However, Fisher took a geometric approach, which is presented in a general way by Stuart and Ord [2, Section 11.2, Exercise 11.6, Section 16.2]; this provides a rationale for applying the term *degrees of freedom* to the parameter of the χ^2 distribution.

Briefly, the joint density of $(x_1, \dots, x_n)$ in Euclidean n space (the sample space of $x_1, \dots, x_n$) when $\mu = 0$ and $\sigma = 1$ is constant over a hypersphere $u = x_1^2 + \dots + x_n^2$ of dimension n , centered at the origin. If one now imposes p functionally independent linear restrictions of the form

$$a_1 x_1 + \dots + a_n x_n = 0 \quad (1)$$

and then restricts the x 's to lie on p hyperplanes through the origin. "The result ... will be to constrain the variables to a hypersphere of p fewer dimensions", i.e., of dimension $n - p$, defining a χ^2_{n-p} distribution. In this context, the term *degrees of freedom* means the dimension of the space in which the relevant hypersphere is constructed.

Degrees of freedom also played a part in the widening rift between Fisher and Karl Pearson [1, pp. 84–87]. Pearson had used χ^2_3 tables to assess the goodness of fit of data in 2×2 contingency tables, but with row and column totals fixed, there

are three linear constraints on four cell frequencies, not one, implying that one should use χ^2 having $(4 - 3 = 1)$ degree of freedom.

In another connotation, the degrees of freedom of a model for expected values of random variables is the excess of the number of variables over the number of parameters in the model. In the case of the *analysis of variance* (see **Analysis of Variance**) or general linear model with independent, homoscedastic **residuals**, the expected value of the **sum of squares** of differences between observed values and corresponding expected values, fitted by *least squares* (see **Least-Squares Estimation**), is

$$\begin{aligned} &(\text{number of degrees of freedom}) \\ &\times (\text{residual variance}) \end{aligned}$$

Further insight into degrees of freedom can be gained from large-sample properties of likelihood-ratio tests. Under certain regularity conditions, suppose that the number of parameters free to vary under a null hypothesis H_0 (possibly composite) and under a more general alternative H_1 are r_0 and r_1 , respectively. Let λ_n be the likelihood ratio, with n being the **sample size**. Then under H_0 [3, Section 13.8] as n increases,

$$-2 \log \lambda_n \rightarrow \chi^2_\nu \quad (2)$$

where $\nu = r_1 - r_0$. Here the degrees of freedom parameter is the difference between the numbers of parameters free to vary under H_1 and H_0 .

References

- [1] Box, J.F. (1978). *R. A. Fisher, The Life of a Scientist*, John Wiley & Sons, New York.
- [2] Stuart, A. & Ord, J.K. (1987). *Kendall's Advanced Theory of Statistics*, 5th Edition, Oxford University Press, New York, Vol. 1.
- [3] Wilks, S.S. (1962). *Mathematical Statistics*, John Wiley & Sons, New York.

Further Reading

Siegel, A.F. (1979). The noncentral chi-squared distribution with zero degrees of freedom and testing for uniformity, *Biometrika* **66**, 381–386.

JEROEN DE MAST

Article originally published in *Encyclopedia of Statistical Sciences*, 2nd Edition (2005, John Wiley & Sons, Inc.). Minor revisions for this publication by Jeroen de Mast.

Graphical Representation of Data

Graphs, charts, and diagrams offer effective display and enable easy comprehension of complex, multifaceted relationships. Gnanadesikan and Wilk [1] point out that “man is a geometrical animal and seems to need and want pictures for parsimony and to stimulate insight.” Various forms of statistical graphs have been in use for over 200 years.

Beniger and Robyn [2] cite the following as first or near-first uses of statistical graphs: Playfair’s [3] use of the bar chart to display Scotland’s 1781 imports and exports for 17 countries; Fourier’s [4] cumulative distribution of population age in Paris in 1817; Lalanne’s [5] contour plot of temperature by hour and month; and Perozzo’s [6] stereogram display of Sweden’s population for the period 1750–1875 by age groupings. Fienberg [7] notes that Lorenz [8] made the first use of the probability–probability (P–P) plot in 1905. Each of these plots is discussed subsequently. Beniger and Robyn [2], Cox [9], and Fienberg [7] provide additional historical accounts and insights on the evolution of graphics.

Today, graphical methods play an important role in all aspects of a statistical investigation – from the initial exploratory plots, through various stages of analysis, to the final communication and display of results. Many persons consider graphical displays as the single most effective, robust statistical tool.

Not only are graphical procedures helpful but in many cases they are essential. Tukey [10] claims that “the greatest value of a picture is when it forces us to notice what we never expected to see.” This is no better exemplified than by Anscombe’s data sets [11], where plots of four equal-size data sets (Figure 1) reveal large differences among the sets even though all sets produce the same linear regression summaries. Mahon [12] maintains that statisticians’ responsibilities include communication of their findings to decision makers, who frequently are statistically naive, and the best way to accomplish this is through the power of the picture.

Good graphs should be simple, self-explanatory, and not deceiving. Cox [9] offers the following guidelines:

1. The axes should be clearly labeled with the names of the variables and the units of measurement.
2. Scale breaks should be used for false origins.
3. Comparison of related diagrams should be easy, for example, by using identical scales of measurement and placing diagrams side by side.
4. Scales should be arranged so that systematic and approximately linear relations are plotted at roughly 45° to the x axis.
5. Legends should make diagrams as nearly self-explanatory (i.e., independent of the text) as is feasible.
6. Interpretation should not be prejudiced by the technique of presentation.

Most of the graphs discussed here, which involve spatial relationships, are implicitly or explicitly on Cartesian or rectangular coordinate grids, with axes that meet at right angles. The horizontal axis is the *abscissa* or x axis and the vertical axis is the *ordinate* or y axis. Each point on the grid is uniquely specified by an x and a y value, denoted by the ordered pair (x, y) . Ordinary graph paper utilizes linear scales for both axes. Other scales commonly used are logarithmic (see Figure 8) and inverse distribution functions (see Figure 6). Craver [13] includes a discussion of plotting techniques with over 200 graph papers that may be copied without permission of the publisher.

This discussion includes old and new graphical forms that have broad application or are specialized but commonly used. Taxonomies based on the uses of graphs have been addressed by several authors, including Tukey [14], Fienberg [7], and Schmid and Schmid [15]. This discussion includes references to more than 50 different graphical displays (Table 1) and is organized according to the principal functions of graphical techniques:

- Exploration
- Analysis
- Communication and display of results
- Graphical aids

Some graphical displays are used in a variety of ways; however, each display here is discussed only in the context of its widest use.

2 Graphical Representation of Data

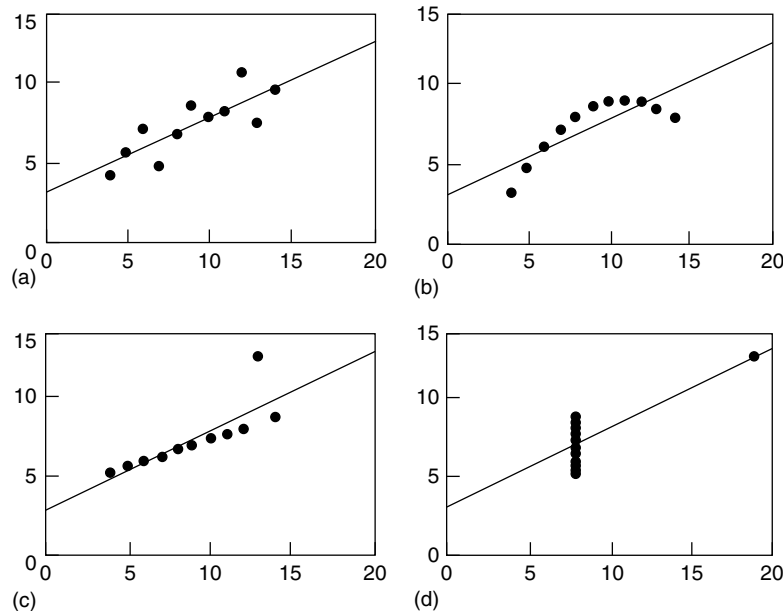


Figure 1 Anscombe's [11] plots of four equal-size data sets, all of which yield the same regression summaries [Reprinted with permission from American Statistical Association. ©1973.]

Exploratory Graphs

Exploratory graphs are used to help diagnose characteristics of the data and to suggest appropriate statistical analyses and models. They usually do not require assumptions about the behavior of the data or the system or mechanism that generated the data.

Data Condensation

A listing or tabulation of data can be very difficult to comprehend, even for relatively small data sets. Data condensation techniques, discussed in most elementary statistics texts, include several types of *frequency distributions* (see, e.g., Freund [16] and Johnson and Leone [17]). These distributions associate the frequency of occurrence with each distinct value or distinct group of values in a data set. Ordinarily, data from a continuous variable will first be grouped into intervals, preferably of equal length, which completely cover the range of the data without overlap. The number or length of these intervals is usually best determined from the size of the data set, with larger sets effectively able to support more intervals. Table 2 presents four commonly

used forms: frequency, relative frequency, cumulative frequency, and cumulative relative frequency. Here carbon monoxide emissions (grams per mile) of 794 cars are grouped into intervals of length 24, where the upper limit is included (denoted by the square upper interval bracket) and the lower limit is not (denoted by the lower open parenthesis). Columns 4 and 6 depict relative frequencies, which are scaled versions (divided by 794) of columns 3 and 5, respectively.

The four distributions tabulated in Table 2 are useful data summaries; however, plots of them can help the data analyst develop an even better understanding of the data. A *histogram* is a bar graph associating frequencies or relative frequencies with data intervals. The histogram for carbon monoxide data shown in Figure 2 clearly shows a positively skew (see **Skewness**), unimodal distribution with modal interval (72–96). Other forms of histograms use symbols such as dots (*dot-array diagram*) or asterisks in place of bars, with each symbol representing a designated number of counts.

A *frequency polygon* is similar to a histogram. Points are plotted at coordinates representing interval midpoints and the associated frequency; consecutive

Table 1 Graphical displays used in the analysis and interpretation of data

Exploratory plots		
Data condensation	Relationship among variables	
	Two variables	Three or more variables
Histogram	Scatter plot	Labeled scatter plot
Dot-array diagram	Sequence plot	Glyphs and metroglyphs
Stem and leaf diagram	Autocorrelation plot	Weather-vane plot
Frequency polygon	Cross- correlation plot	Biplot
Ogive		Face plots
Box and whisker plot		Fourier plot
		Cluster trees
		Similarity and preference maps
		Multidimensional scaling displays
Graphs used in the analysis of data		
Distribution assessment	Model adequacy and assumption verification	Decision making
Probability plot	Average <i>versus</i> standard deviation	Control chart
Q-Q plot	Residual plots	CUSUM chart
P-P plot	Partial residual plot	Youden plot
Hanging histogram	Component-plus-residual plot	Half-normal plot
Rootogram		C_p plot
Poissonness plot		Ridge trace
Communication and display of results		
Quantitative graphics	Summary of statistical analyses	Graphical aids
Bar chart	Means plots	Power curves
Pictogram	Sliding reference distribution	Sample-size curves
Pie chart	Notched-box plot	Confidence limits
Contour plot	Factor space/response	Nomographs
Stereogram	Interaction plot	Graph paper
Color map	Contour plot	Trilinear coordinates
	Predicted response plot	
	Confidence region plot	

points are connected with straight lines (e.g., in Table 2, plot column 3 *vs* column 2). The form of this graph is analogous to that of a **probability density function**.

A disadvantage of a grouped-data histogram is that individual data points cannot be identified since all the data falling in a given interval are indistinguishable. A display that circumvents this difficulty is the *stem and leaf diagram*, a modified histogram with “stems” corresponding to interval groups and “leaves” corresponding to bars. Tukey [10] gives a thorough discussion of stem and leaf and its variations.

An *ogive* is a graph of the cumulative frequencies (or cumulative relative frequencies) against the upper limits of the intervals (e.g., from Table 2, plot column 5 *vs* the upper limit of each interval in column 1) where straight lines connect consecutive points. An ogive is a grouped-data analog of a graph of the empirical cumulative distribution function and is especially useful in graphically estimating percentiles (see **Quantiles**), which are data values associated with specified cumulative percents. Figure 3 shows the ogive for the carbon monoxide data and how it is used to obtain the 25th percentile (i.e., the lower quartile).

4 Graphical Representation of Data

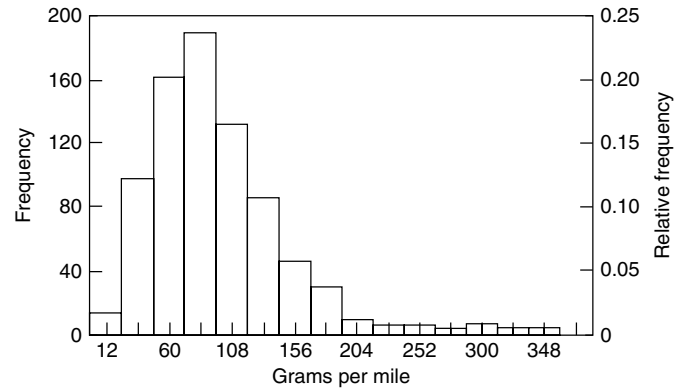


Figure 2 Histogram of Environmental Protection Agency surveillance data (1957–1967) on carbon monoxide emissions from 794 cars

Table 2 Frequency distributions of carbon monoxide data

(1)	(2)	(3)	(4)	(5)	(6)
Interval	Interval midpoint	Frequency	Relative frequency	Cumulative Frequency	Cumulative relative frequency
1. (0–24]	12	13	0.016	13	0.016
2. (24–48] ^(a)	36	98	0.123	111	0.140
3. (48–72]	60	161	0.203	272	0.343
4. (72–96]	84	189	0.238	461	0.581
5. (96–120]	108	148	0.186	609	0.767
6. (120–144]	132	85	0.107	694	0.874
7. (144–168]	156	45	0.057	739	0.931
8. (168–192]	180	30	0.038	769	0.969
9. (192–216]	204	10	0.013	779	0.981
10. (216–240]	228	5	0.006	784	0.987
11. (240–264]	252	5	0.006	789	0.994
12. (264–288]	276	1	0.001	790	0.995
13. (288–312]	300	2	0.003	792	0.997
14. (312–336]	324	1	0.001	793	0.999
15. (336–360]	348	1	0.001	794	1.000

^(a) Notation designates inclusion of all values greater than 24 and less than or equal to 48

Another display, which highlights five important characteristics of a data set, is a *box and whisker* or box plot. The box, usually aligned vertically, encloses the interquartile range (see **Descriptive Statistics**), with the lower line identifying the 25th percentile (lower *quartile*, see **Quartiles**) and the upper one the 75th (upper quartile). A line sectioning the box displays the 50th percentile (median) and its relative position within the interquartile range. The whiskers at either end may extend to the extreme values, or, for large data sets, to the 10th/90th

or 5th/95th percentiles. These plots are especially convenient for comparing two or more data sets, as shown in Figure 4 for winter snowfalls of Buffalo and Rochester, New York. (See Tukey [10] for further discussion, and McGill *et al.* [18] for some variations.)

Relationships between Two Variables

Often of interest is the relationship, if any, between x and y , or the development of a model to predict y

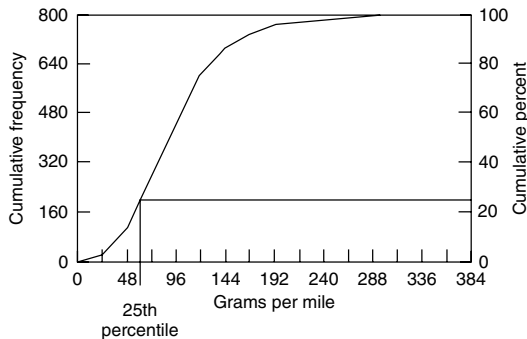


Figure 3 Ogive of automobile carbon monoxide emissions data shown in Figure 2

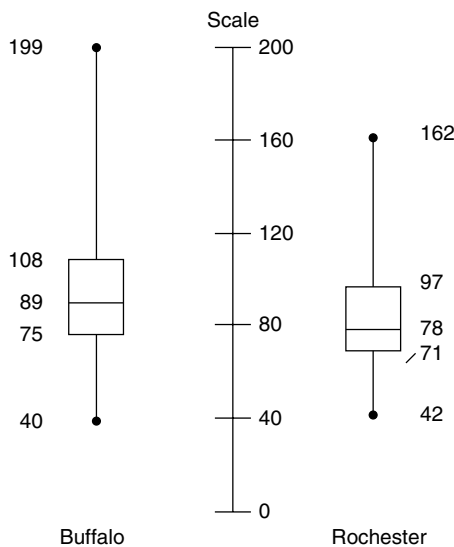


Figure 4 Box and whisker plots comparing winter snowfalls of Buffalo and Rochester, New York, (1939–1940 to 1977–1978) and demonstrating little distributional difference (contrary to popular belief). (Based on local climatological data gathered by the National Oceanic and Atmospheric Administration, National Climatic Center, Asheville, NC)

given the value of x . Ordinarily, an (x, y) measurement pair is obtained from the same experimental unit, as shown in Table 3.

A usual first step is to construct a *scatter plot*, a collection of plotted points representing the measurement pairs (x_i, y_i) , $i = 1, \dots, n$. The importance of scatter plots was seen in Figure 1. These four very different sets of data yield the same regression line

Table 3

Unit	x	y
Object	Diameter	Weight
Person	Age	Height
Product	Raw material purity	Quality

(drawn on the plots) and associated statistics [11]. Consequently, the numerical results of an analysis, without the benefit of a look at plots of the data, could result in invalid conclusions.

The objective of regression analyses is to develop a mathematical relationship between a measured response or dependent variable, y , and two or more predictor or independent variables, $x_1, x_2, \dots, x_p$. A usual initial step is to plot y versus each of the x variables individually and to plot each x_i versus each of the other x 's. This results in a total of $p + p(p - 1)/2$ plots. Plots of y versus x_i enable one to identify the x_i variables that appear to have large effects, to assess the form of a relationship between the y and x_i variables, and to determine whether any unusual data points are present. Plots of x_i versus x_j , $i \neq j$, help to identify strong correlations that may exist among the predictor variables. It is important to recognize such correlations because least-squares regression techniques work best when these correlations are small [19] (see **Collinearity; Data Collection**).

In many instances data are collected sequentially in time and a plot of the data versus the sequence of collection can help identify sources of important effects. Figure 5 shows a *sequence plot* of gasoline mileage of an automobile versus the sequence of gasoline fill-ups. The large seasonal effect (summer mileage is higher than winter mileage) and the increase in gasoline mileage due to a major tune-up are clearly evident.

Observations collected sequentially in time, such as the gasoline mileage data plotted in Figure 5, form a time series (see **Time Series Analysis**). Statistical modeling of a time series is largely accomplished by studying the correlation between observations separated by $1, 2, \dots, n - 1$ units in time. For example, the lag-1 autocorrelation is the correlation coefficient between observations collected at time i and time $i + 1$, $i = 1, 2, \dots, n - 1$; it measures the linear relationship among all pairs of consecutive observations (see **Autocorrelation Function**). An *autocorrelation*

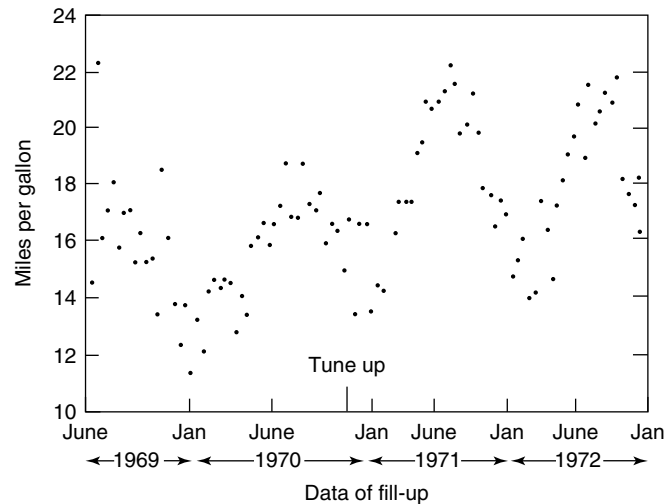


Figure 5 Sequence plot of gasoline mileage data. Note the seasonal variation and the increased average value and decreased variation in mileage after tune-up

plot may be used to study the “correlation structure” in the data where the lag- j autocorrelation coefficient computed between observations i and $i + j$ is plotted *versus* lag j . A *cross-correlation plot* between two time series is developed similarly. The lag- j correlation coefficient, computed between observations i in one series and observations $i + j$ in the other series, is plotted *versus* the lag j . Box and Jenkins [20] discuss the construction and interpretation of these plots, in addition to presenting examples on the use of sequence plots and the modeling of time series.

Relationships among More than Two Variables

Scatter plots directly display relationships between two variables. Values of a third variable can be incorporated in a *labeled scatter plot*, in which each plotted point (whose location designates the values of two variables) is labeled by a symbol designating a level of the third variable. Anderson [21] extended these to “pictorialized” scatter plots, called *glyphs* and *metroglyphs*, where each coordinate point is plotted as a circle and has two or more rays emanating from it; the length of each ray is indicative of the value of the variable associated with that ray. Bruntz *et al.* [22] developed a variation of the glyph for four variables, called the *weathervane plot*, where values of two of the variables are again indicated by the plotted coordinates and the other two by using

variable-sized plotting symbols and variable-length arrows attached to the symbols.

Tukey [10], Gabriel’s *biplot* [23], and Mandel [24] provide innovative methods for displaying two-way tables of data. All three approaches involve fitting a model to the table of data and then constructing various plots of the coefficients in the fitted model to study the relationships among the variables.

Chernoff [25] used facial characteristics to display values of up to 18 variables through *face plots*. Each face represents a multivariate datum point, and each variable is represented by a different facial characteristic, such as size of eyes or shape of mouth. Experience has shown that the interpretation of these plots can be affected by how the variables are assigned to facial characteristics.

Exploratory plots of raw data are usually less effective with measurements on four or more variables (e.g., height, weight, age, sex, race, etc., of a subject). The usual approach then is to reduce the dimensionality of the data by grouping variables with common properties or identifying and eliminating unimportant variables. The variables plotted may be functions of the original variables as specified by a statistical model. Exploratory analysis of the data is then conducted with plots of the “reduced” data. Andrews’ *Fourier plots* [26], *data clustering* [27], *similarity and preference mapping* [28], and geometrical representations of multidimensional scaling

analyses [29] are examples of such procedures. One often attempts to determine whether there are two or more groups (i.e., clusters) of observations within the data set. When several different groups are identified, the next step is usually to determine why the groups are different. These methods are sophisticated and require computer programs for implementation on a routine basis. (See Gnanadesikan [30], Everitt [31].)

Graphs Used in the Analysis of Data

The graphical methods discussed next generally depend on assumptions of the analysis. Decisions made from these displays may be either subjective in nature, such as a visual assessment of an underlying distribution, or objective, such as an out-of-control signal from a control chart.

Distribution Assessment and Probability Plots

The *probability plot* is a widely used graphical procedure for data analysis. Since other graphical techniques discussed in this article require a basic understanding of it, a brief discussion follows (see **Probability Plots**).

A probability plot on linear rectangular coordinates is a collection of two-dimensional points specifying corresponding quantiles from two distributions. Typically, one distribution is empirical and the other is a hypothesized theoretical one. The primary purpose is to determine visually if the data could have arisen from the given theoretical distribution. If the empirical distribution is similar to the theoretical one, the expected shape of the plot is approximately a straight line; conversely, large departures from linearity suggest different distributions and may indicate how the distributions differ.

Imagine a sample of size n in which the data $y_1, \dots, y_n$ are independent observations on a random variable Y having some continuous distribution function (df) (see **Probability Density Function (PDF)**). The ordered data $y_{(i)}, y_{(1)} \leq \dots \leq y_{(n)}, i = 1, \dots, n$, represent sample **Quantiles** and are plotted against theoretical quantiles, $x_i = F^{-1}(p_i)$, where F^{-1} denotes the inverse of F , the hypothesized df of Y . Moreover, $F(y)$ may involve unknown location (v) and scale (δ) parameters (not necessarily the mean and the standard deviation) as long as $F((y - v)/\delta)$ is completely specified. If Y has a df F ,

then $p_i = F(x_i) = F((y_{(i)} - v)/\delta)$; x_i is called the *reduced $y_{(i)}$ variate* and is a function of the unknown parameters.

Selection of p_i for use in plotting has been much discussed. Suggested choices have been $i/(n+1)$, $(i-1/2)/n$, $(2i-1)/2n$, and $(i-3/8)/(n+1/4)$. Kimball [32] discusses some choices of p_i in the context of probability plots.

Now $F^{-1}(p_i)$ is not expressible in closed form for most commonly encountered distributions and thus provides an obstacle to easy evaluations of x_i . An equivalent procedure often employed to avoid this difficulty is to plot $y_{(i)}$ against p_i on *probability paper*, which is rectangular graph paper with an F^{-1} scale for the p axis. The p_i entries on the p axis are commonly called *plotting positions*. Naturally, a different type of probability paper is needed for each family of distributions, F .

Normal probability paper, with F^{-1} based on the normal distribution, is most common, although many others have been developed (see, e.g., King [33]). Normal probability paper is available in two versions: arithmetic probability paper, where the data axis has a linear scale, and logarithmic probability paper, where the data axis has a natural logarithmic scale. The latter version is used to check for a lognormal distribution (see **Probability Density Functions**).

To illustrate the procedure, two data sets of size $n = 15$ have been plotted on normal (arithmetic) probability paper shown in Figure 6, using $p_i =$

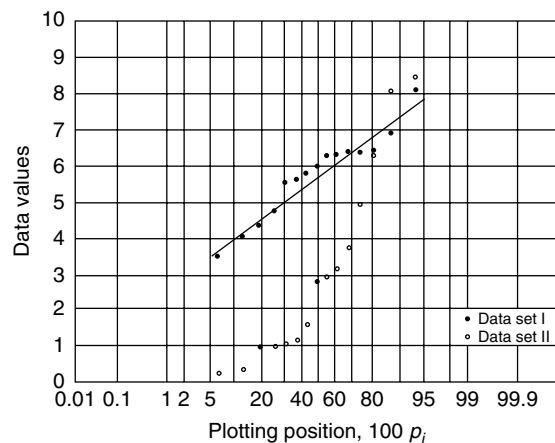


Figure 6 Comparison of two data sets plotted on normal probability paper. Set I can be adequately approximated by a **normal distribution**, whereas set II cannot

8 Graphical Representation of Data

$i/(n+1)$. Here the horizontal axis is labeled in percentage; $100p_i$ has been plotted against the i th smallest observation in each set. The plotted points of set I appear to cluster around the straight line drawn through them, visually supportive evidence that these data come from a population that can be adequately approximated by a normal distribution. The points for set II, however, bend upward at the higher percents, suggesting that the data come from a distribution with a longer upper tail (i.e., larger upper quantiles) than the normal distribution.

If set I data are viewed as sufficiently normal, graphical estimates of the mean (μ) and the standard deviation (σ) are easily obtained by noting that the 50th percentile of the normal distribution corresponds to the mean and that the difference between the 84th and 50th percentiles corresponds to one standard deviation. The respective graphical estimates from the line fitted by eye for μ and σ are 5.7 and $7.0 - 5.7 = 1.3$, respectively.

The conclusions from Figure 6 were expected because the data from set I are, in fact, random normal deviates with $\mu = 6$ and $\sigma = 1$. The data in set II are random lognormal deviates with $\mu = 0.6$ and $\sigma = 1$. A plot of set II data on logarithmic probability paper produces a more nearly linear collection of points.

Visual inference, such as determining here whether the collection of points forms a straight line, is fairly easy, but should be used with theoretical understanding to enhance its reliability. A curvilinear pattern of points based on a large sample offers more evidence against the hypothesized distribution than does the same pattern based on a smaller sample. For example, the plot of set I data exhibits some asymmetric irregularities, which are due to random fluctuations; however, a similar pattern of irregularity based on a much larger sample would be much more unlikely from a normal distribution. Daniel [34] and Daniel and Wood [35] give excellent discussions on the behavior of normal probability plots.

The probability plot is a special case of the quantile–quantile ($Q-Q$) plot [36] (see **Probability Plots**), which is a quantile–quantile comparison of two distributions, either or both of which may be empirical or theoretical; whereas a probability plot is typically a display of sample data on probability paper (i.e., empirical vs theoretical). $Q-Q$ plots are particularly useful because a straight line will result when comparing the distributions of X and

Y , whenever one variable can be expressed as a linear function of the other. $Q-Q$ plots are relatively more discriminating in low-density or low-frequency regions (usually the tails) of a distribution than near high-density regions, since in low-density regions quantiles are rapidly changing functions of p . In the plot, this translates into comparatively larger distances between consecutive quantiles in low-density areas. The quantiles in Figure 6 illustrate this, especially the larger empirical ones of set II.

A related plot considered by Wilk and Gnanadesikan [36] is the $P-P$ plot. Here, for varying x_i , $p_{i1} = F_1(x_i)$ is plotted against $p_{i2} = F_2(x_i)$, where F_j , $j = 1, 2$, denotes the df (empirical or theoretical). If $F_1 = F_2$ for all x_i , the resulting plot is a straight line with unit slope through the origin. This plot is especially discriminating near high-density regions, since here the probabilities are more rapidly changing functions of x_i than in low-density regions. The $P-P$ plot is not as widely used as the $Q-Q$ plot since it does not remain linear if either variable is transformed linearly (e.g., by a location or scale change).

For large data sets, an obvious approach for comparison of data with a probability model is a graph of a fitted theoretical density (parameters estimated from data), with the appropriate scale adjustment, superimposed on a histogram. Gross differences between the ordinates of the distributions are easily detected. A translation of the differences to a reference line (instead of a reference curve) to facilitate visual discrimination is easily accomplished by hanging the bars of the histogram from the density curve [37].

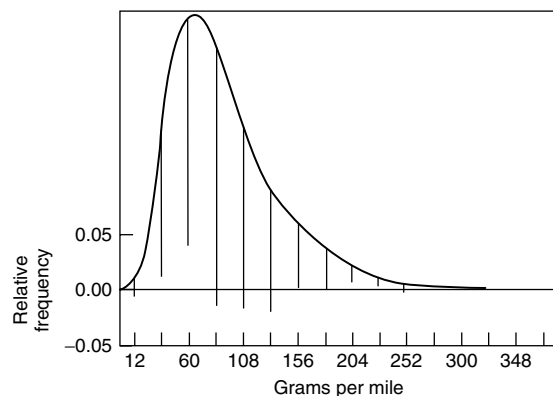


Figure 7 Hanging carbon monoxide histogram from fitted lognormal distribution

Figure 7 illustrates the *hanging histogram*, where the histogram for the carbon monoxide data (Figure 2) is hung from a lognormal distribution. Slight, systematic variation about the reference line suggests that the data are slightly more skewed to the right in the high-density area than in the lognormal distribution.

Further improvements in detecting systematic variation may be achieved by rootograms [37]. A *hanging rootogram* is analogous to a hanging histogram except that the square roots of the ordinate values are graphed. The *suspended rootogram* is an upside-down graph of the **residuals** about the baseline of the hanging rootogram.

Graphical assessments for discrete distributions can also be made by comparing the histogram of the data to the fitted probability density, $p(x)$. However, as with continuous distributions, curvilinear discrimination may be difficult and linearizing procedures are helpful. A general approach is to determine a function of $p(x)$, say, $r(x) = r(p(x))$, which is linearly related to a function of x , say, $s(x)$. Then using sample data one calculates relative frequencies to estimate $p(x)$, evaluates $r(x)$, and plots $r(x)$ against $s(x)$. The absence of systematic departures from linearity offers some evidence that the data could arise from density $p(x)$. The slope and intercept will be functions of the parameters and can be used to estimate the parameters graphically. A suitable $r(x)$ may be obtained by simply transforming $p(x)$; for example, taking logarithms of the density of the discrete **Pareto** distribution, where $p(x) \propto x^\lambda$, gives $r(x) = \log p(x)$ and $s(x) = \log x$. In other cases, ratios of consecutive probabilities (e.g., $p(x+1)/p(x)$) are linear functions of $s(x)$ [38]. Table 4 summarizes these ratios for three commonly encountered discrete distributions.

Ord [39] expands on the foregoing ideas by defining a class of discrete distributions where $r(x) = xp(x)/(p(x-1))$ is a linear function of x (i.e., $s(x) = x$), thereby keeping the same abscissa scale. Distributions in this class are binomial, negative binomial, Poisson, logarithmic, and uniform.

These graphical tests for discrete distributions may be difficult to interpret because the sample relative frequencies have nonhomogeneous variances. This difficulty may be compounded when using ratios of the relative frequencies as functions of $s(x)$. These procedures are therefore recommended more as exploratory than confirmatory.

Another graphical technique for the Poisson distribution (see **Probability Density Functions**) is the *Poissonness plot* [40], similar in spirit to probability plotting. It can also be applied to truncated Poisson data or any one-parameter exponential family of discrete distributions such as the binomial.

For further insights, see Parzen [41].

Model Adequacy and Assumption Verification

Any statistical analysis is based on certain assumptions. Those usually associated with least-squares regression analysis are that experimental errors are independent and have a homogeneous variance and a normal (Gaussian) distribution. It is standard practice to check these assumptions as part of the analysis. These checks, most often done graphically, have the desirable by-product of forcing the analyst to look at the data critically. This can be effectively accomplished by graphical analysis of both the raw data and residuals from the fitted model. In addition to assumption verification, this evaluation frequently results in the discovery of unusual observations or unsuspected relationships. Most of the plots discussed below are applications of graphical forms previously discussed.

If repeat observations have been obtained for each of k groups representing different situations or conditions being studied, a scatter plot of the group standard deviation, s_i , versus the group mean, $\bar{y}_i$, $i = 1, \dots, k$, will appear random and show little correlation when the homogeneous variance assumption is satisfied. Box *et al.* [42] point out that if these assumptions are not satisfied, this plot can be used to determine a transformed measurement scale on which the assumptions will be more nearly satisfied (Figure 8).

The normal distribution assumption can be checked in this situation from a histogram of the **Residuals**, $r_{ij} = y_{ij} - \bar{y}_i$, between the observations in each group (y_{ij}) and the average of the group ($\bar{y}_i$). This histogram will tend to be bell-shaped if the normal distribution and homogeneous variance assumptions are satisfied. Alternatively, especially for small data sets, the r_{ij} 's may be plotted on normal probability paper. The expected shape is a straight line if these assumptions are appropriate.

Replicate observations often are not available. However, the analysis assumptions and adequacy of the form of the model can still be checked by

Table 4

Distribution	$p(x)$	$r(x)$	=	Intercept	+	Slope	$\times$	$s(x)$
Binomial	$\binom{n}{x}\pi^x(1-\pi)^{n-x}, x = 0, \dots, n$	$\frac{p(x+1)}{p(x)}$		$-\frac{\pi}{1-\pi}$		$\frac{(n+1)\pi}{1-\pi}$		$\frac{1}{x+1}$
Poisson	$\frac{e^{-\lambda}\lambda^x}{x!}, x = 0, 1, 2, \dots$	$\frac{p(x)}{p(x+1)}$		$\frac{1}{\lambda}$		$\frac{1}{\lambda}$		x
Pascal	$\binom{x-1}{k-1}\pi^k(1-\pi)^{x-k}, x = k, k+1, \dots$	$\frac{p(x)}{p(x+1)}$		$\frac{1}{1-\pi}$		$\frac{k-1}{1-\pi}$		$\frac{1}{x}$

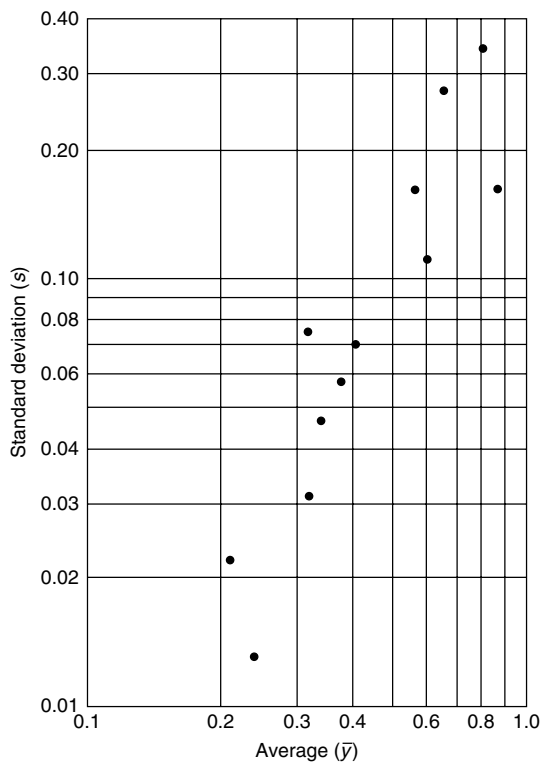


Figure 8 Toxic agent data [42, Table 7.11]. Linear log–log relationship suggests that a power transformation will produce homogeneous variances. The slope of the line indicates the necessary power [©WileyInterscience.]

constructing plots of the residuals or standardized residuals [43], from the fitted model. The residual associated with observation y_i is $r_i = y_i - \hat{y}_i$, where $\hat{y}_i$ is the value of y_i predicted by the model fitted to the data. Four types of *residual plots* are routinely constructed [44]: plots on normal probability paper, residuals (r_i) versus predicted values ($\hat{y}_i$), sequence plot of residuals, and residuals (r_i) versus predictor variables (x_j). Note that although the residuals are not mutually independent, the effect of

the correlation structure on the utility of these plots is negligible [45].

The *plot of residuals on normal probability paper* provides a check on the normal distribution assumption. Substantive deviations from linearity may be due to a nonnormal distribution of experimental errors, the presence of atypical (i.e., outlying) data points, or an inadequate model (Figure 9).

The *residuals (r_i) versus the predicted values ($\hat{y}_i$) plot* will show a random distribution of points (no trends, shifts, or peculiar points) if all assumptions are satisfied. Any curvilinear relationship indicates that the model is inadequate (Figure 10). Nonhomogeneous variance is indicated if the spread in r_i changes with $\hat{y}_i$. When the spread increases linearly with $\hat{y}_i$, a log transformation (i.e., replace y in the analysis

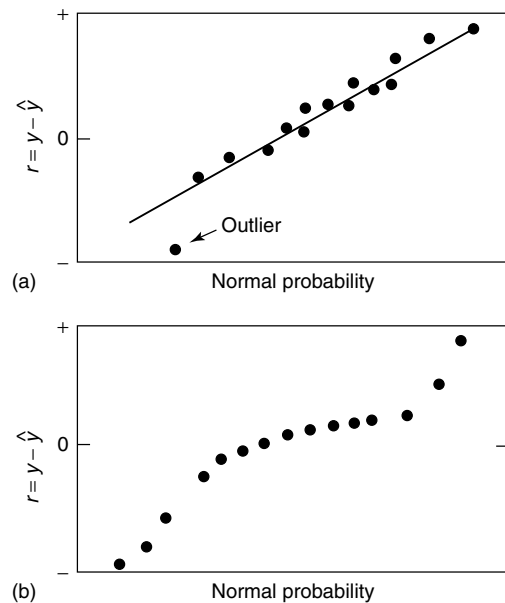


Figure 9 Normal probability plot of residuals. Plot (a) shows an outlying data point and plot (b) shows a set of residuals not normally distributed

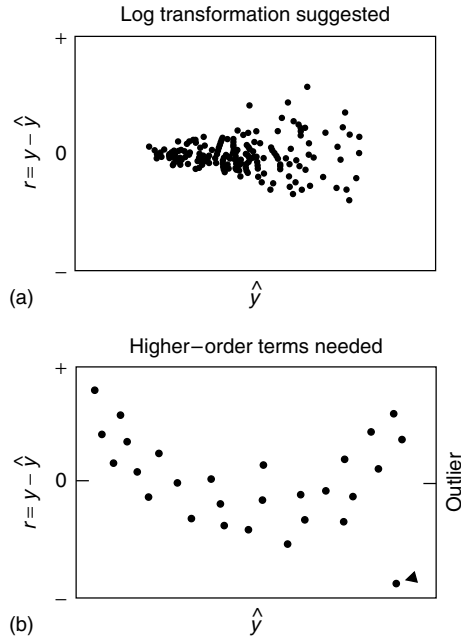


Figure 10 Residuals *versus* fitted values. Plot (a) shows increased residual variability with increasing fitted values and plot (b) shows a curvilinear relationship

by $y' = \log y$) will often produce a response scale on which the homogeneous variance assumption will be satisfied (Figure 10). Plots of the residuals (r_i) *versus* raw observations (y_i) are of little value because they will always show a linear correlation with a value of $(1 - R^2)^{1/2}$, where R^2 is the coefficient of determination (see **Coefficient of Determination (R^2)**) of the fitted model [46].

If all analysis assumptions are satisfied, the *sequence plot of the residuals* (r_i) can be expected to show a random distribution of points and contain no trends, shifts, or atypical points. Any trends or shifts here suggest that one or more variables not included in the model may have changed during the collection of the data (Figure 11). This plot may show cycles in the residuals and other dependencies, indicating that the assumption of independence of experimental errors is not appropriate. This assumption can also be checked by constructing an *autocorrelation plot* of the residuals (see earlier discussion).

The *residuals* (r_i) *versus* the *predictor variables* (x_i) *plot* should also show a random distribution of points. Any smooth patterns or trends

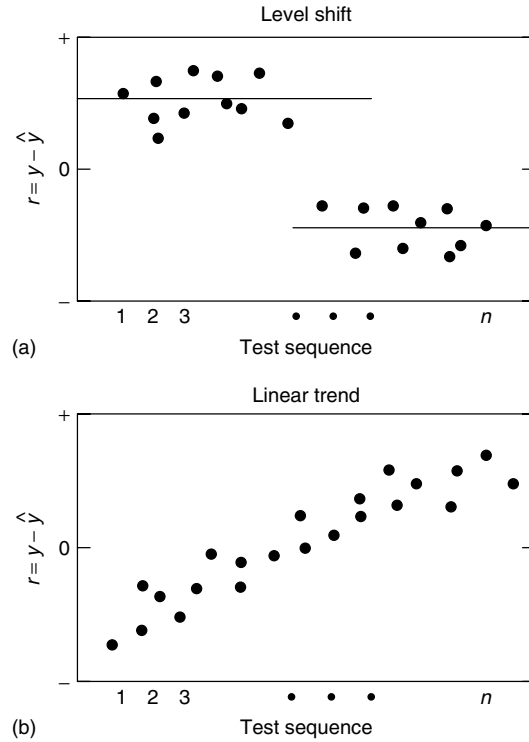


Figure 11 Sequence plot of residuals. Plot (a) shows an abrupt change and plot (b) shows a gradual change due to factors not accounted for by the model

suggests that the form of the model may not be appropriate (Figure 12).

The scatter plots of residuals discussed above are not independent of each other. Peculiarities and trends observed in one plot usually show up in one or more of the others. Collectively, these plots provide a good check on data behavior and the model construction process.

With a large number of predictor variables it is sometimes hard to see relationships between y and x_i in scatter plots. Larsen and McCleary [47] developed the *partial residual plot* to overcome this problem. Wood [48] and Daniel and Wood [35] refer to these as *component-plus-residual plots*. A regression model must be fitted to the data before the plot can be constructed. In effect, the relationships of all the other variables are removed and the plot of the component-plus-residual *versus* x_i displays only the computed relationship between y and x_i and the residual variation in the data.

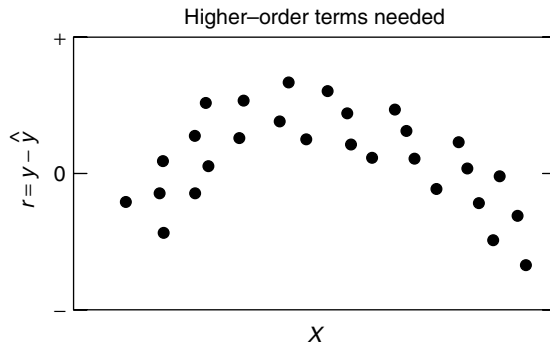


Figure 12 Residuals *versus* a predictor variable

Plots for Decision Making

At various points in a statistical analysis, decisions concerning the effects of variables and differences among groups of data are made. Statisticians and other scientists have developed a variety of statistical techniques that use graphical displays to make these decisions. In some instances (e.g., control charts, Youden plots) these contain both the data in raw or reduced form and a measure of their uncertainty. The user, in effect, makes decisions from the plot rather than by calculating a test statistic. In other situations (e.g., half-normal plot, ridge trace) a measure of uncertainty is not available but the analyst makes decisions concerning the magnitude of an effect or the appropriateness of a model by assessing deviations from expected or desired appearance.

The *control chart* (see **Control Charts, Overview**) is widely used to control industrial production and analytical measurement processes [49]. It is a *sequence plot* of a measurement or statistic (average, range, etc.) *versus* time sequence together with limits to reflect the expected random variation in the plotted points. For example, on a plot of sample averages, limits of ± 3 standard deviations are typically shown about the process average to reflect the uncertainty in the averages. The process is considered to be out of control if a plotted average falls outside the limits. This suggests that a process shift has occurred and a search for an **assignable cause** should be made. The *cumulative sum control chart* (see **Cumulative Sum (CUSUM) Chart**) is another popular **process control** technique, particularly useful in detecting small process shifts [50, 51].

Ott [53] used the control chart concept to develop his analysis-of-means plotting procedure for the interpretation of data that would ordinarily be analyzed by analysis-of-variance techniques. Schilling [54, 55] systematized Ott's procedure and extended it past the cross-classification designs to **incomplete block** experiments and studies involving random effects. The analysis-of-means procedure enables those familiar with control chart concepts and technology to quickly develop an ability to analyze the results from experimental designs.

The *Youden plot* [56] was developed to study the ability of laboratories to perform a test procedure. Samples of similar materials, A and B, are sent to a number of laboratories participating in a collaborative test. Each laboratory runs a predetermined number of replicate tests on each sample for a number of different characteristics. A Youden plot is constructed for each measured characteristic. Each point represents a different laboratory, where the average of the replicate results on material A is plotted *versus* the average results on material B (Figure 13). Differences along a 45° line reflect between-lab variation. Differences in the direction perpendicular to the 45° line reflect within-lab and lab-by-material interaction variation. An uncertainty ellipse can be used to identify problem laboratories. Any point outside this ellipse is an indication that the associated laboratory's results are significantly different from those of the laboratories within the ellipse. Mandel and Lashof [52] generalized and extended

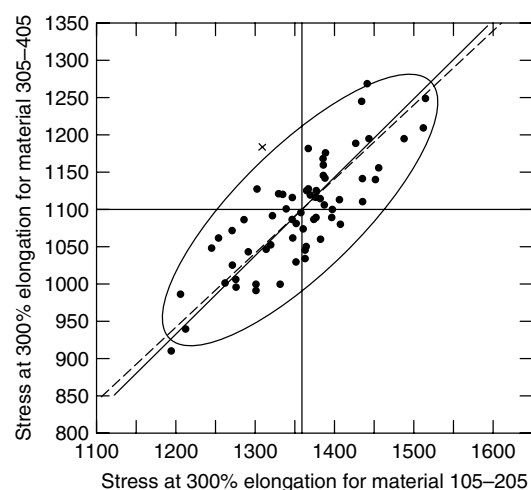


Figure 13 Youden plot [52] [©American Society of Quality, 1974.]

the construction and interpretation of Youden's plot. Ott [57] showed how this concept can be used to study paired measurements. For example, "before" and "after" measurements are frequently collected to evaluate a process change or a manufacturing stage of an industrial process.

Daniel [34] developed the **Half-Normal Plot** to interpret two-level factorial and fractional **Factorial Experiments**. It displays the absolute value of the $n - 1$ contrasts (i.e., main effects and interactions) from an n -run experiment, *versus* the probability scale of the half-normal distribution (Figure 14). Identification of large effects (positive or negative) is enhanced by plotting the absolute value of the contrast. Alternatively, to preserve the sign of the contrast, the contrast value may be plotted on normal probability paper [42]. With no significant effects, either plot will appear as a straight line. Significant effects are indicated by the associated contrast falling off the line. Although this plot is usually assessed visually, uncertainty limits and decision guides have been developed by Zahn [58, 59]. The half-normal plot is not restricted to two-level experiments and can be used in the interpretation of any experiment for which the effects can be described by independent 1-degree-of-freedom contrasts.

The half-normal plot was significant in marking the beginning of extensive research on probability plotting methods in the early 1960s. During this time the statistical community became convinced of the usefulness and effectiveness of graphical techniques; many of these developments were discussed earlier.

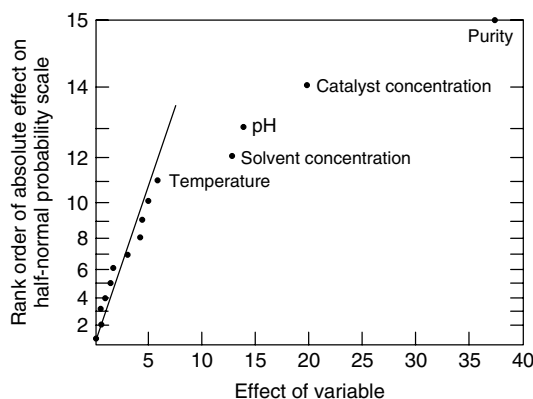


Figure 14 Half-normal plot from a 2^{5-1} factorial experiment showing four important effects on a color response

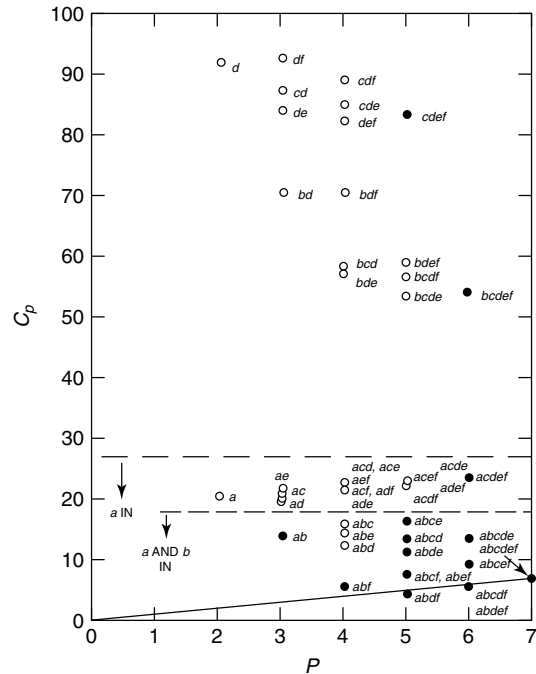


Figure 15 C_p plot [60]. The six variables in the equation are denoted by a , b , c , d , e , and f . The number of terms in the equation is denoted by P [Reprinted with permission from American Statistical Association. ©1966.]

Research in the 1960s and 1970s also focused on regression analysis and the fitting of equations to data when the predictor variables (x 's) are correlated. Two displays, the C_p plot and the *ridge trace*, were developed as graphical aids in these studies.

The C_p plot, suggested by C. L. Mallows and popularized by Gorman and Toman [60] and Daniel and Wood [35], is used to determine which variables should be included in the regression equation. It (Figure 15) is constructed by plotting, for each equation considered, C_p versus p , where $C_p = \text{RSS}_p / s^2 - (n - 2p)$, p is the number of terms in the equation, RSS_p is the residual **sum of squares** for the p -term equation of interest, n is the total number of observations, and s^2 is the residual mean square obtained when all the variables are included in the equation. If all the important terms are in the equation, $C_p = p$. The line $C_p = p$ is included on the plot and one looks for the equations that have points falling near this line. Points above the line indicate that significant terms have not been included. The objective is to find the equation with the smallest

number of terms for which $C_p \doteq p$.

The *ridge trace*, developed by Hoerl and Kennard [19] and discussed by Marquardt and Snee [61], identifies which regression coefficients are poorly estimated because of correlations among the predictor variables (x 's). This is accomplished by plotting the regression coefficients *versus* the **bias** parameter, k , which is used to calculate the ridge regression coefficients, $\hat{\beta} = (\mathbf{X}'\mathbf{X} + k\mathbf{I})^{-1}\mathbf{X}'\mathbf{Y}$ (Figure 16). Coefficients whose sign and magnitude are affected by correlations among the predictor variables (x 's) will change rapidly as k increases. Hoerl and Kennard recommend that an appropriate value for k is the smallest for which the coefficients “stabilize” or change very little as k increases. If the coefficients remain nearly constant for $k > 0$, then $k = 0$ is suggested; this indicates that the least-squares coefficients should be used.

Communication and Display of Results

Quantitative Graphics

Quantitative graphics encompass general methods of presenting and summarizing numerical information,

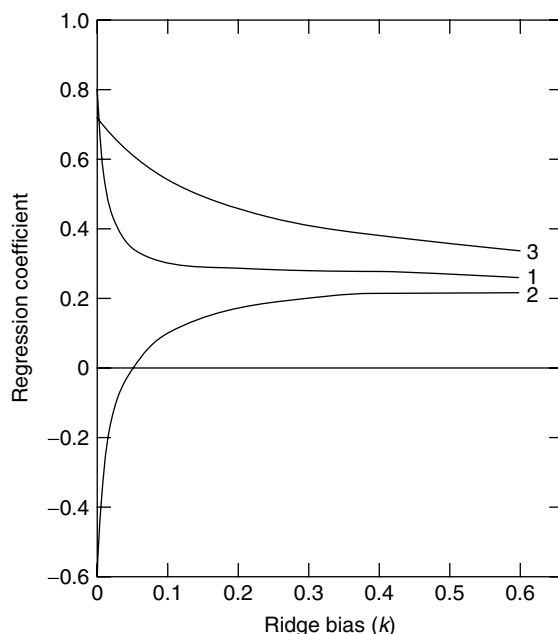


Figure 16 Ridge trace showing instability of regression coefficients 1 and 2

usually for easy comprehension by the layman. There can be no complete catalog of these graphics since their form is limited only by one's ingenuity. Beniger and Robyn [2] trace the historical development since the seventeenth century of quantitative graphics and offer several pictorial examples illustrating the improved sophistication of graphs. Schmid and Schmid [15], in a comprehensive handbook, discuss and illustrate many quantitative graphical forms. A few of the more common forms follow briefly.

A *bar chart* is similar in appearance to a histogram (Figure 2) but more general in application. It is frequently used to make quantitative comparisons of qualitative variables, as a company might do to compare revenues among departments. Comparisons within and between companies are easily made by superimposing a similar graph of another comparable company and identifying the bars.

A *pictogram* is similar to the bar chart except that “bars” consist of objects related to the response tallied. For example, figures of people might be used in a population comparison where each figure represents a specified number of people [62]. When partial objects are used (e.g., the bottom half of a person), it should be stated whether the height or volume of the figure is proportional to frequency.

A *pie chart* is a useful display for comparing attributes on a relative basis, usually a percentage. The angle formed by each slice of the pie is an indication of that attribute's relative worth. Governments use this device to illustrate the number of pennies of a taxpayer's dollar going into each major budget category.

A *contour plot* may effectively depict a relationship among three variables by one or more contours. Each contour is a locus of values of two variables associated with a constant value of the third. A relief map that shows latitudinal and longitudinal locations of constant altitude by contours or isolines is a familiar example. Contours displaying a **response surface**, $y = f(x_1, x_2)$, are discussed subsequently (Figure 21). A *stereogram* is another form for displaying spatial relationship of three variables in two dimensions, similar to a draftsman's three-dimensional perspective drawing in which the viewing angle is not perpendicular to any of the object's surfaces (planes).

The field of *computer graphics* has rapidly expanded, offering many new types of graphics, see Tufte [63].

Summary of Statistical Analysis

Many of the displays previously discussed are useful in summarizing and communicating the results of a study. For example, histograms or dot-array diagrams are used to display the distribution of data. Box plots are useful (see Figure 4) for large data sets or when several groups of data are compared.

The results of many analyses may be displayed by a *means plot* of the group means ($\bar{y}$) together with a measure of their uncertainty, such as **standard error** ($\bar{y} \pm \text{SE}$), confidence limits ($\bar{y} \pm t\text{SE}$), or least significant interval limits ($LSI = \bar{y} \pm \text{LSD}/2$). The use of the LSI [65] is particularly advantageous because the intervals provide a useful decision tool (Figure 17). Any two averages are significantly different at the assigned probability level if and only if their LSIs do not overlap. Intervals based on the standard error and confidence limits for the mean do not have this straightforward interpretation. Similar intervals can be developed for other multiple comparison procedures, such as those developed by Tukey (honest significant difference), Dunnett, or Scheffé.

Box *et al.* [42] use the *sliding reference distribution* to display and compare means graphically. The distribution width is determined by the **Standard Error** of the means. Any two means not within the bounds of the distribution are judged to

be significantly different. They also use this display in the interpretation of factor effects in two-level factorial experiments.

The *notched-box plot* of McGill *et al.* [18] is the nonparametric analog of the LSI-type plots discussed above. The median of the sample and confidence limits for the difference between two medians are displayed in this plot; any two medians whose intervals do not overlap are significantly different. McGill *et al.* [18] discuss other variations and applications of the box plot.

Factorial experimental designs (see **Factorial Experiments**) are widely used in science and engineering. In many instances the effects of two and three variables are displayed on *factor space/response plots* with squares and cubes similar to those shown in Figure 18. The numbers in these figures are the mean responses obtained at designated values of the independent variables studied. These figures show the region, or factor space, over which the experiments were conducted, and aid the investigator in determining how the response changes as one moves around in the experimental region.

An *interaction plot* is used to study the nature and magnitude of the interaction of two variables from a designed experiment by plotting the response (averaged over all replicates) *versus* the level of one of the independent variables while holding the

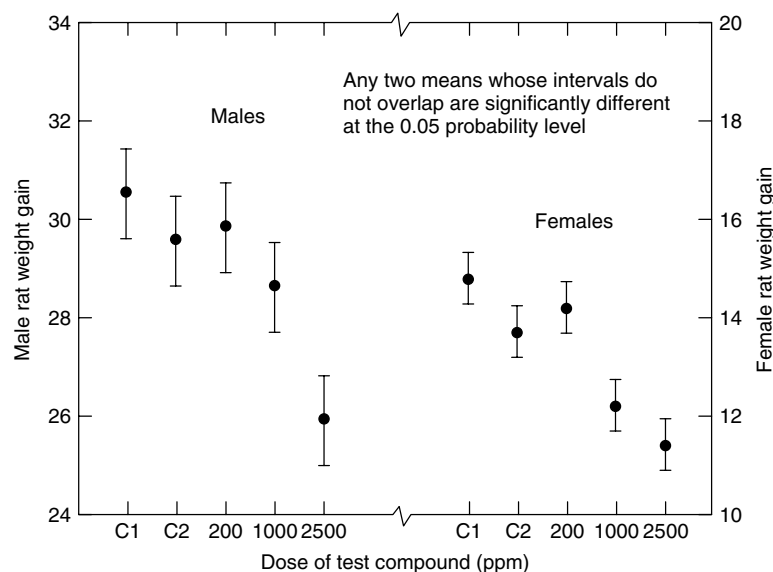


Figure 17 Least significant interval plot [64] [©International Biometrics Society, 1979.]

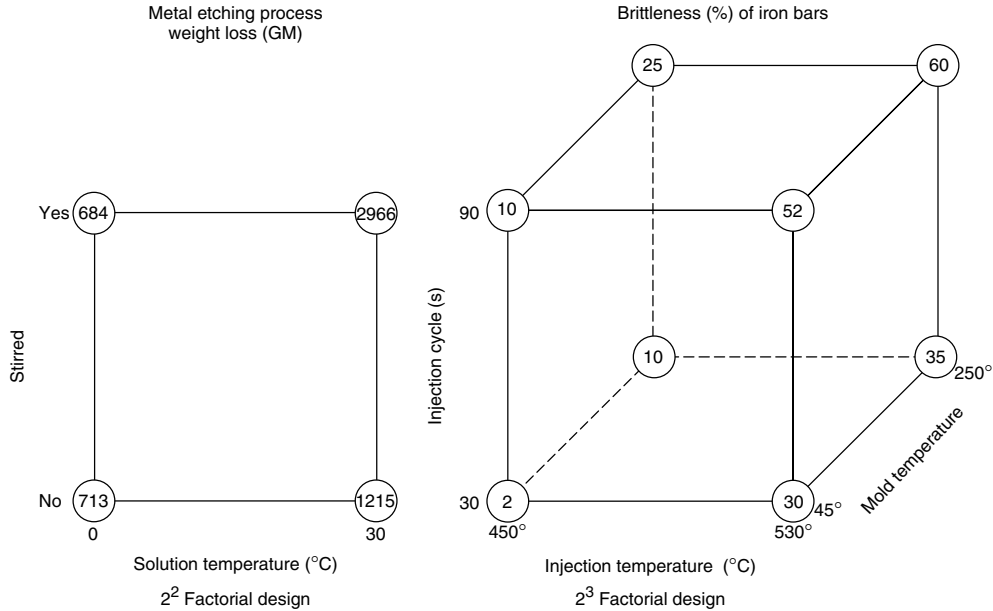


Figure 18 Factor space/response plots showing results of 2^2 and 2^3 experiments. The average responses are shown at the corners of the figures [66]

other independent variable(s) fixed at a given level. In Figure 19 it is seen by the nonparallel lines that the size of the effect of solution temperature on weight loss depends on whether stirring occurred. Solution temperature and stirring are said to *interact*; another way of saying that their effects are not additive. Monlezun [67] discusses ways of plotting and interpreting three-factor interactions.

The objective of many experiments is to develop a prediction model for the response (y) of the system as a function of experimental (predictor) variables. A typical two-predictor variable model, based on a second-order Taylor series, is

$$y = b_0 + b_1X_1 + b_2X_2 + b_{12}X_1X_2 + b_{11}X_1^2 + b_{22}X_2^2 \quad (1)$$

where the b 's are coefficients estimated by regression analysis techniques. One of the best ways to understand all the effects described by this equation is by a *contour plot*, which gives the loci of X_1 and X_2 values associated with a fixed value of the response. By constructing contours on a rectangular $X_1 - X_2$ grid for a series of fixed values of y , one obtains a global picture of the response surface (Figure 20a).

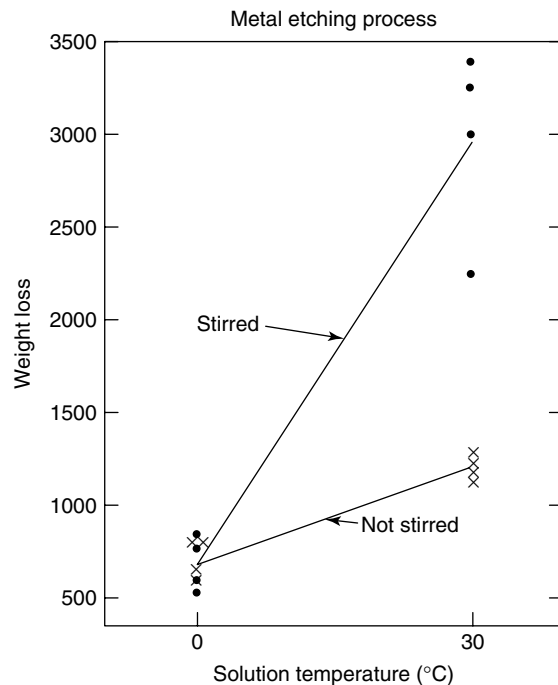


Figure 19 Nonparallelism of lines suggests an interaction between stirring and solution temperature. Stirring has a larger effect at 30 °C than at 0 °C

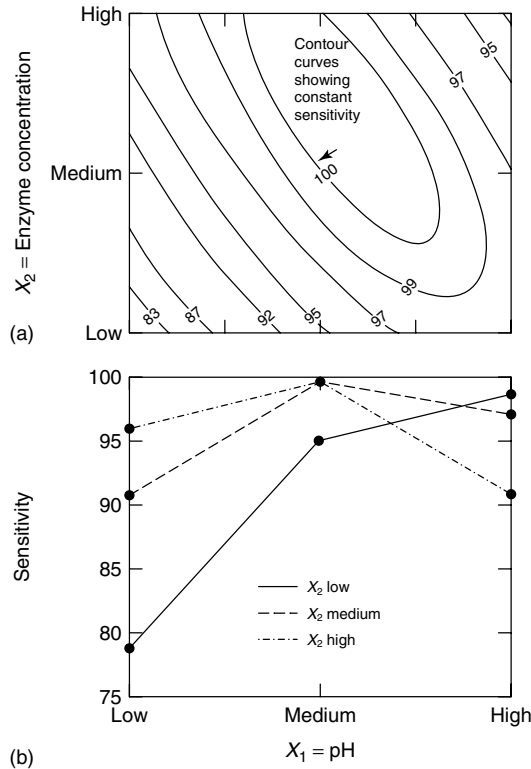


Figure 20 Contour plot of (a) quadratic response surface and (b) X_1 – X_2 interaction plot [70] [©American Association for Clinical Chemistry, 1979.]

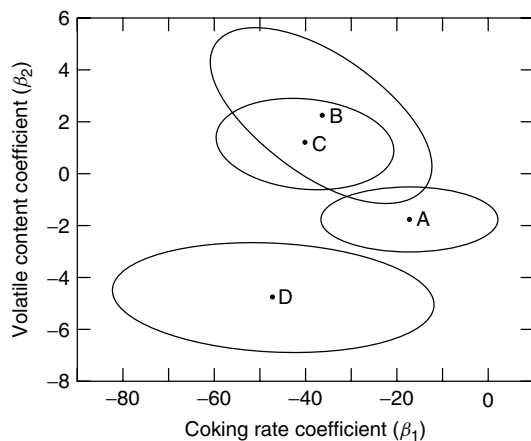


Figure 21 Joint 95% β_1 – β_2 confidence regions for coal classes A, B, C, and D [74] [©American Society for Quality, 1973.]

The response surface of a mixture system such as a gasoline or paint can be displayed by a contour plot on a triangular grid [68, 69]. Another use of trilinear coordinates is shown in Figure 22.

Predicted response plots help in interpreting interaction terms in regression equations [71]. The interaction between X_i and X_j is studied by fixing the other variables in the equation at some level (e.g., set $X_k = \bar{X}_k, k \neq i, j$) and constructing an interaction plot of the values predicted by the equation for different values of X_i and X_j (Figure 20b). Similar plots are also useful in interpreting response surface models for mixture systems [72] and in the analysis of the results of mixture screening experiments [73].

A contour plot is also the basis for a *confidence region plot*, used to display a plausible region of simultaneous hypothesized values of two or more parameters. For example, Figure 21 [74] shows a contour representing the joint confidence region (see **Confidence Intervals**) for regression parameters β_1 and β_2 associated with each of four coal classes. The figure clearly shows that these four groups do not all have the same values of regression parameters.

Confidence region plots on trilinear coordinates [75] can also be used to display and identify variable dependence in two-way contingency tables where one variable is partitioned into three categories (see also Draper *et al.* [76]). A two-dimensional plot results (Figure 22) from the constraint that the sum of the

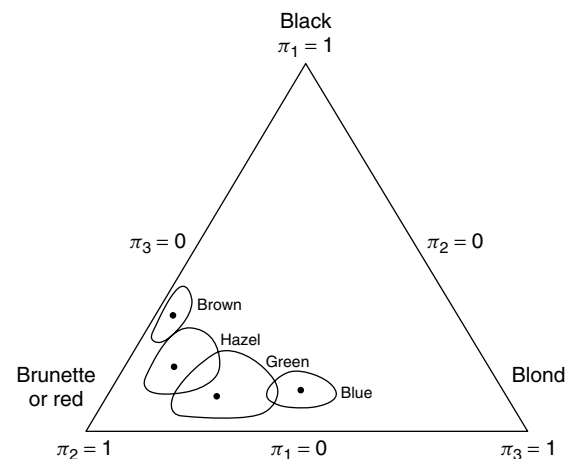


Figure 22 Joint 95% confidence region plots for hair color probabilities for brown, hazel, green, and blue eye colors [75] [Reprinted with permission from American Statistical Association. ©1974.]

three probabilities, π_1, π_2 , and π_3 , is unity. It is also possible to display three-dimensional response surfaces or confidence regions *via* a series of two-dimensional slices through three-dimensional space.

Graphical Aids

Statisticians, like chemists, physicists, and other scientists, rely on graphical devices to help them do their job more effectively. The following graphs are used as “tools of the trade” in offering a parsimonious representation for complex functional relationships among two or more variables.

Graphs of *power functions* (see **Power**) or *operating characteristics* (e.g., Natrella [77]) are used for the evaluation of error probabilities of *hypothesis tests* (see **Hypothesis Testing**) when expressed as a function of the unknown parameter(s). Test procedures are easily compared by superimposing two power curves on the same set of axes. *Sample-size curves* are constructed from a family of operating characteristic curves, each associated with a different **sample size**. These curves are useful in planning the size of an experiment to control the chances of wrong decisions. Natrella [77] offers a number of these for some common testing situations.

Similar contour graphics using families of curves indexed by sample size are also useful for determining confidence limits (see **Confidence Intervals**). These are especially convenient when the endpoints are not expressible in closed form, such as those for the correlation coefficient or the success probability in a Bernoulli sample [78].

Nomographs are graphical representations of mathematical relationships, frequently involving more than three variables. Unlike most graphics, they do not offer a picture of the relationship, but only enable the determination of the value of (usually) any one variable from the specification of the others. Levens [79] discusses techniques for straightline and curved scale nomograph construction. Other statistical nomographs have appeared in the *Journal of Quality Technology*.

References

- [1] Gnanadesikan, R. & Wilk, M.B. (1969). in *Multivariate Analysis*, P.R. Krishnaiah, ed, Academic Press, New York, Vol. 2, pp. 593–637.
- [2] Beniger, J.R. & Robyn, D.L. (1978). Quantitative graphics in **statistics**: a brief history, *American Statistician* **32**, 1–11.
- [3] Playfair, W. (1786). *The Commercial and Political Atlas*, London.
- [4] Fourier, J.B.J. (1821). *Recherches Statistiques sur la Ville de Paris et le Department de la Seine*, Vol. 1, pp. 1–70.
- [5] Lalanne, L. (1845). *Appendix to Cours Complet de météorologie de L.F. Kaemtzt*, translated and annotated by C. Martins, Paris.
- [6] Perozzo, L. (1880). Della rappresentazione grafica di una collettività di individui nella successione del tempo, *Annals of Statistics* **12**, 1–16.
- [7] Fienberg, S.E. (1979). Graphical methods in statistics, *American Statistician* **33**, 165–178.
- [8] Lorenz, M.O. (1905). Methods for measuring concentration of wealth, *Journal of the American Statistical Association* **9**, 209–219.
- [9] Cox, D.R. (1978). Some remarks on the role of **statistics** in graphical methods, *Applied Statistics* **27**, 4–9.
- [10] Tukey, J.W. (1977). *Exploratory Data Analysis*, Addison-Wesley, Reading, Mass.
- [11] Anscombe, F.J. (1973). Graphs in statistical analysis, *American Statistician* **27**, 17–21.
- [12] Mahon, B.H. (1977). *Journal of the Royal Statistical Society. A* **140**, 298–307.
- [13] Craver, J.S. (1980). *Graph Paper from Your Copier*, H.P. Books, Tucson, Ariz.
- [14] Tukey, J.W. (1972). in *Statistical Papers in Honor of George W. Snedecor*, T.A. Bancroft, ed, Iowa State University Press, Ames, Iowa.
- [15] Schmid, C.F. & Schmid, S.E. (1979). *Handbook of Graphic Presentation*, 2nd Edition, John Wiley & Sons, New York.
- [16] Freund, J.E. (1976). *Statistics: A First Course*, 2nd Edition, Prentice-Hall, Englewood Cliffs.
- [17] Johnson, N.L. & Leone, F.C. (1977). *Statistics and Experimental Design in Engineering and the Physical Sciences*, 2nd Edition, John Wiley & Sons, New York, Vol. 1.
- [18] McGill, R., Tukey, J.W. & Larsen, W.A. (1978). Variations of box plots, *American Statistician* **32**, 12–16.
- [19] Hoerl, A.E. & Kennard, R.W. (1970). Ridge regression: biased estimation of nonorthogonal problems, *Technometrics* **12**, 55–70.
- [20] Box, G.E.P. & Jenkins, G.M. (1970). *Time Series Analysis: Forecasting and Control*, Holden-Day, San Francisco.
- [21] Anderson, E. (1960). A semigraphical method for the analysis of complex problems, *Technometrics* **2**, 387–391.
- [22] Bruntz, S.M., Cleveland, W.S., Kleiner, B. & Warner, J.L. (1974). *Proceedings of the Symposium on Atmospheric Diffusion and Air Pollution*, American Meteorological Society, 125–128.

-
- [23] Gabriel, K.R. (1971). The biplot graphic display of matrices with application to principal component analysis, *Biometrika* **58**, 453–467.
- [24] Mandel, J. (1971). A new analysis of variance model for non-additive data, *Technometrics* **13**, 1–18.
- [25] Chernoff, H. (1973). Using faces to represent points in k-dimensional space graphically, *Journal of the American Statistical Association* **68**, 361–368.
- [26] Andrews, D.F. (1972). Plots of high-dimensional data, *Biometrics* **28**, 125–136.
- [27] Hartigan, J.A. (1975). *Clustering Algorithms*, Wiley-Interscience, New York.
- [28] Green, P.E. & Carmone, F.J. (1970). *Multidimensional Scaling and Related Techniques in Marketing Analysis*, Allyn & Bacon, Boston.
- [29] Lingoes, J.C., Roskam, E.E. & Borg, I. (1979). *Geometrical Representations of Relational Data – Readings in Multidimensional Scaling*, Mathesis Press, Ann Arbor.
- [30] Gnanadesikan, R. (1977). *Methods for Statistical Data Analysis of Multivariate Observations*, John Wiley & Sons, New York.
- [31] Everitt, B.S. (1978). *Graphical Techniques for Multivariate Data*, North-Holland, Amsterdam.
- [32] Kimball, B.F. (1960). On the choice of plotting positions on probability paper, *Journal of the American Statistical Association* **55**, 546–560.
- [33] King, J.R. (1971). *Probability Charts for Decision Making*, Industrial Press, New York.
- [34] Daniel, C. (1959). Use of half-normal plots in interpreting factorial two-level experiments, *Technometrics* **1**, 311–341.
- [35] Daniel, C. & Wood, F.S. (1980). *Fitting Equations to Data*, 2nd Edition, John & Wiley Sons, New York.
- [36] Wilk, M.B. & Gnanadesikan, R. (1968). Probability plotting methods for the analysis of data, *Biometrika* **55**, 1–17.
- [37] Wainer, H. (1974). The suspended rootogram and other visual displays: an empirical validation, *American Statistician* **28**, 143–145.
- [38] Dubey, S.D. (1966). Graphical tests for discrete distributions, *American Statistician* **20**, 23–24.
- [39] Ord, J.K. (1967). Graphical methods for a class of discrete distributions, *Journal of the Royal Statistical Society. A* **13**, 232–238.
- [40] Hoaglin, D.C. (1980). A Poissonness plot, *American Statistician* **34**, 146–149.
- [41] Parzen, E. (1979). A very good discussion of the properties of quantile functions, density quantile, *Journal of the American Statistical Association* **74**, 105–121.
- [42] Box, G.E.P., Hunter, W.G. & Hunter, J.S. (1978). *Statistics for Experimenters*, Wiley-Interscience, New York.
- [43] Draper, N.R. & Behnken, D.W. (1972). Residuals and their variance patterns, *Technometrics* **14**, 101–111.
- [44] Draper, N.R. & Smith, H. (1981). *Applied Regression Analysis*, 2nd Edition, John Wiley & Sons, New York.
- [45] Anscombe, F.J. & Tukey, J.W. (1963). The examination and analysis of residuals, *Technometrics* **5**, 141–160.
- [46] Jackson, J.E. & Lawton, W.H. (1967). *Technometrics* **9**, 339–341.
- [47] Larsen, W. & McCleary, S. (1972). The use of partial residual plots in regression anal, *Technometrics* **14**, 781–790.
- [48] Wood, F.S. (1973). The use of individual effects and residuals in fitting equations to data, *Technometrics* **15**, 677–695.
- [49] Grant, E.L. & Leavenworth, R.S. (1980). *Statistical Quality Control*, 5th Edition, McGraw-Hill, New York.
- [50] Barnard, G.A. (1959). Control charts and stochastic processes, *Journal of the Royal Statistical Society. B* **21**, 239–271.
- [51] Lucas, J.M. (1976). The design and use of V-mask control schemes, *Journal of Quality Technology* **8**, 1–12.
- [52] Mandel, J. & Lashof, T.W. (1974). Interpretation and generalization of. Youden's two-sample diagram, *Journal of Quality Technology* **6**, 22–36.
- [53] Ott, E.R. (1967). *Industrial Quality Control* **24**, 101–109.
- [54] Schilling, E.G. (1973). A systemic approach to the analysis of means, *Journal of Quality Technology Part I*, **5**, 93–108.
- [55] Schilling, E.G. (1973). A systemic approach to the analysis of means, *Journal of Quality Technology Parts 2 and 3* **5**, 147–159.
- [56] Youden, W.J. (1959). Graphical diagnosis of interlaboratory test results, *Industrial Quality Control* **15**, 133–137.
- [57] Ott, E.R. (1957). *Industrial Quality Control* **13**, 1–4.
- [58] Zahn, D.A. (1975). Modifications of and revised critical values for the half-normal plot, *Technometrics* **17**, 189–200.
- [59] Zahn, D.A. (1975). An empirical study of the half-normal plot, *Technometrics* **17**, 201–212.
- [60] Gorman, J.W. & Toman, R.J. (1966). Selection of variables for fitting equations to data, *Technometrics* **8**, 27–51.
- [61] Marquardt, D.W. & Snee, R.D. (1975). Ridge regression in practice, *American Statistician* **29**, 3–20.
- [62] Joiner, B.L. (1975). *International Statistical Review* **43**, 339–340.
- [63] Tufte, E.R. (1997). *Visual Explanation; Images and Quantities, Evidence and Narrative*, Graphics Press, Cheshire.
- [64] Snee, R.D., Acuff, S.K. & Gibson, J.R. (1979). A useful method for the analysis of growth studies, *Biometrics* **35**, 835–848.
- [65] Andrews, H.P., Snee, R.D. & Sarnier, M.H. (1980). Graphical display of means, *American Statistician* **34**, 195–199.
- [66] Bennett, C.A. & Franklin, N.L. (1954). *Statistical Analysis in Chemistry and the Chemical Industries*, John Wiley & Sons, New York.
- [67] Monlezun, C.J. (1979). Two-dimensional plots for interpreting interactions in the three-factor analysis of variance model, *American Statistician* **33**, 63.
- [68] Cornell, J.A. (1981). *Experiments with Mixtures*, Wiley-Interscience, New York.

- [69] Snee, R.D. (1979). Experimenting with mixtures, *Chemtech* **9**, 702–710.
- [70] Rautela, G.S., Snee, R.D. & Miller, W.K. (1979). Response-surface co-optimization of reaction conditions in **clinical** chemical methods, *Clinical Chemistry* **25**, 1954–1964.
- [71] Snee, R.D. (1973). Some Aspects of nonorthogonal data analysis. Part I. Developing prediction equations, *Journal of Quality Technology* **5**, 67–79.
- [72] Snee, R.D. (1975). *Technometrics* **17**, 425–430.
- [73] Snee, R.D. & Marquardt, D.W. (1976). Screening concepts and designs for experiments with mixtures, *Technometrics* **18**, 19–29.
- [74] Snee, R.D. (1973). *Journal of Quality Technology* **5**, 109–122.
- [75] Snee, R.D. (1974). Graphical display of two-way contingency tables, *American Statistician* **28**, 9–12.
- [76] Draper, N.R., Hunter, W.G. & Tierney, D.E. (1969). Which product is better? *Technometrics* **11**, 309–320.
- [77] Natrella, M.G. (1963). *Experimental Statistics, National Bureau of Standards Handbook 91*, U.S. Government Printing Office, Washington, DC.
- [78] Pearson, E.S. & Hartley, H.O. (1970). *Biometrika Tables for Statisticians*, Cambridge University Press, Cambridge, Vol. 1.
- [79] Levens, A.S. (1959). *Nomography*, John Wiley & Sons, New York.

RONALD D. SNEE AND CHARLES G. PFEIFER

Article originally published in Encyclopedia of Statistical Sciences, 2nd Edition (2005, John Wiley & Sons, Inc.). Minor revisions for this publication by Jeroen de Mast.

Hypothesis Testing

Introduction

Many problems lead to repetitions of an experiment with just two possible outcomes, e.g., (yes, no), (dead, alive), (pass, fail), which reminds us of coin tossing, where we get either heads or tails. Our interest is usually the proportion p of responses. We may estimate p by the observed proportion in our experiments. Often, a theory or hypothesis says that p should have some specific value, and our purpose is to test this hypothesis. To fix our notation, let us speak of coin tossing, where our hypothesis is that the coin is “fair” (i.e., that $p = 1/2$). This is a *statistical hypothesis* since it governs the pattern of outcomes of a series of independent but identical experiments, Bernoulli trials, in fact.

Common sense suggests that if x/n , the observed proportion of heads in n tosses or trials, is near $1/2$, this hypothesis gains support from our data, whereas if x/n is far from $1/2$, we will begin to doubt that $p = 1/2$. Of course, random fluctuations can lead to values of x/n far from $1/2$ even if $p = 1/2$. As n becomes larger, these become less likely, so we should then have a more powerful test. **Probability Theory** enables us to sharpen these ideas. The number of heads, x , in n trials has a *Binomial* distribution (see **Probability Density Functions**) which may be approximated by a *normal* or *Gaussian* distribution with mean np and variance npq , which equal $n/2$ and $n/4$, respectively, for the “null hypothesis” $p = 1/2$. Thus, if the coin is fair, the chance that x will differ from $n/2$ by more than $1.96\sqrt{n/4}$ is fairly accurately $1/20$. For example, if $n = 100$, the probability that x is outside the interval $(50 - 1.96 \times 5, 50 + 1.96 \times 5)$ or roughly $(40, 60)$ when the coin is fair is about 5% or $1/20$. It has become commonplace in science to say that deviations of less than 2σ (2 is an approximation to 1.96 and here $\sigma = \sqrt{n/4} = 5$ is the standard deviation of the approximately normally distributed “statistic” x) are to be expected and so ignored, but that greater deviations are “significant” (i.e., worth looking into; see **Significance Level**). There is nothing sacred about 5%. Scientists use this as a guide – not for making decisions, which depend upon many other factors.

If we have to reach a yes–no decision, where should we draw the line? Clearly, if we knew that

p could never be less than $1/2$, we would only doubt the “ $p = 1/2$ ” hypothesis if x/n were suspiciously larger than $1/2$. A small x/n can only be due to random fluctuations. Thus our test must depend on the *alternative* hypothesis, as well as on the null hypothesis. Again, the (e.g., financial) consequences of saying erroneously that $p = 1/2$ when $p \neq 1/2$ might be much less than of saying erroneously that $p \neq 1/2$ when in fact $p = 1/2$. Then we would want to be very confident before we reject the null hypothesis. Hence the consequences of the decision must play a role. Finally, experiments cost time and money, so we want to make as few as are necessary to reach a “reliable” decision. With “reliable” clearly defined, how many observations are needed? It is clear that we can never be sure that our decision is right, but we can control the probability that it is right. It is also evident that we must be careful if we first look at the data collected before deciding on n – one might be tempted to cheat.

The development of the ideas in the last paragraph comprises the subject matter of *hypothesis testing*, which is intimately connected with a large part of the theory of mathematical statistics, especially **Estimation** (see also **Overview of Statistics; Probability Theory**). Most experiments have more than two outcomes; e.g., they may yield single or multiple measurements whose postulated probability distributions usually depend on several parameters, not just one as we had above. Null and alternative hypotheses may now be more complicated. It is then easy to devise various procedures for testing the same hypothesis. Hence we will need rules for selecting tests.

If one takes a Bayesian view of the general inference problem, hypothesis tests do not arise – all knowledge is summarized in (prior and posterior) probability distributions of the parameters. Classical hypothesis testing seeks only the best control of the probabilities of the two kinds of erroneous decisions, illustrated above, when the null hypothesis and **sample size** are specified beforehand. In *decision theory*, explicit account is taken of the numerical consequences of decisions. *Sequential analysis* studies particularly efficient methods for terminating a sequence of experiments.

Nonparametric methods involve tests that are valid no matter what the distributions of the observations may be. *Multiple testing* theory began with ways of testing hypotheses suggested by the data rather

2 Hypothesis Testing

than being given *a priori*. *Robustness* is concerned with tests that remain sensitive to departures from the null hypothesis when the distribution of the observations falls in some broad class, and so is a sharpened form of nonparametrics that tries to assume nothing.

Early Ideas

As with so much statistics, the early examples are astronomical. An essay in 1734 by Daniel Bernoulli won a prize at the French Academy, which set the following problem: Is the near coincidence of the orbital planes of the planets an accident or a consequence of their creation? Bernoulli invented a test statistic, i.e., a quantity that would be small if the planes were oriented randomly and independently of one another, and large if they were near coincidence. A modern improvement would use the length of the sum of the unit vectors that are perpendicular to each of the planes. Needless to say, in both cases the quantity is far larger than would arise often by chance if the orbits were in independent randomly oriented planes.

Although many tests were invented in the interim, the χ^2 test [1], Student's t -test [2, Paper 2], and Fisher's derivation of the distribution of the **correlation** coefficient opened the modern era of statistics. Many important fields of application for statistics were found in the experimental sciences as well as the world of affairs, since in more and more areas, data were being consciously gathered to investigate specific hypotheses. Many of these experiments were small. Mathematical techniques became available to find the exact probability distributions of many quantities (i.e., statistics) computed from samples. It was then realized that there were many possible tests and estimators in every problem and so a problem of choice. To bring order out of chaos, it was necessary to seek general principles that would tell the statistician which method *should* be used in any given situation.

While Karl Pearson [3] in his book on the scientific method, (*The Grammar of Science* (1892)), had supported the use of Bayes' formula (the principle of inverse probability) to join previous and current data, he rarely used it. His χ^2 statistic,

$$\chi^2 = \sum_{i=1}^k \frac{(x_i - np_i)^2}{(np_i)} \quad (1)$$

tests whether the numbers $x_1, \dots, x_k$ of k types in a sample size n have been drawn from a population whose members are of k types in proportions $p_1, \dots, p_k$. The expected value of x_i is np_i , so χ^2 will be large when at least some of the x_i 's are far from their expectations when we might doubt the hypothesis. If the calculated value is χ_0^2 , he suggested that one should look at P_0 , the probability that samples of n from the multinomial $(p_1, \dots, p_k)$ population would give values of χ^2 greater than χ_0^2 . If we doubt the hypothesis of the given data yielding χ_0^2 , we should also doubt it for data yielding any larger value of χ^2 . He showed how to compute P_0 , the size of which is a measure of the agreement between the data and the assumed distribution $(p_1, \dots, p_k)$. If P_0 is small, we have two explanations – a rare event has happened, or the assumed distribution is wrong. This is the essence of the *significance test* argument. Not to reject the null hypothesis (that the proportions are $p_1, \dots, p_k$) means only that it is accepted for the moment on a provisional basis. "Student" took the same approach for his t -statistic and for r , the correlation statistic [2, Paper 3] with two important asides. First, he remarks that he would prefer to use inverse probability – but does not know how to set his priors. Second, he recognizes that actual samples may not come from normal distributions which he has had to assume, and gives a reason why his t -test might yet be reliable for samples from some nonnormal populations.

By 1925 we find the subject, almost as one finds it today, in Fisher's highly influential book, *Statistical Methods for Research Workers* [4]. The hypothesis being tested, often the *status quo* (so to speak), is the *null hypothesis*. For the fair-coin example, this is $p = 1/2$. If Pearson's χ^2 were applied to testing that a six-sided die was fair, the null hypothesis would be that $p_1 = p_2 = \dots = p_6 = 1/6$. **Significance Level** of 5% and 1% are suggested. Thus, if the P_0 for χ^2 is less than 0.05, it is said to be "significant". Fisher prepared tables to make these and many other significance tests easier to apply by computing the value of the statistic that would need to be exceeded in order that P_0 be less than 0.05, 0.01, and so on. Such numbers are *significance points* of the relevant statistic. In his book, however, **P Values** have not entirely disappeared. It is said that Fisher tabulated significance points rather than P -values only because Karl Pearson held the copyrights to the P -value tables!

By 1925, Fisher [4] had established his theory of point estimation, including the concepts of sufficiency, consistency, efficiency, likelihood, and **Maximum Likelihood**. But he never had a theory of significance test construction, although he invented or developed almost all the common tests, including some basic **Nonparametric Tests** (e.g., use of normal scores, **randomization**, and permutation tests). All his life, Fisher objected in the strongest terms to the use of inverse probability, but clearly felt the need for something like it since he tried to develop his own substitute, fiducial probability. Further, he always thought and wrote as a mathematician applying his skills in areas of basic natural science, but never, for example, in industry, so he was not attracted to decision theory.

Theory, to assess and derive optimal significance tests, was provided in the next 10 years by Neyman and Egon Pearson. The story of their collaboration has been given by Pearson [5]. The key papers are conveniently collected in Neyman [6], Neyman and Pearson [7], and Pearson [8]. Pearson, like “Student”, was interested in the use of statistical procedures to guide routine work in industry where rules of thumb were needed. Neyman brought a more philosophical and formal approach to the problem.

If, for example, the test for the fairness of a coin mentioned above is routinely applied (recall that the **significance level** chosen was 5%), and if “significant” means that we reject the null hypothesis $p = 1/2$, then two kinds of errors may be made. Five percent of the time when p actually equals $1/2$, we will say it does not (i.e., reject the null hypothesis that $p = 1/2$). However, when p is not equal to $1/2$ (e.g., $p = 3/4$), we will sometimes get a number of heads x in n tosses which fall in the range

$$\begin{aligned} & \left((n/2) - 1.96\sqrt{n/4} \right) \\ & (n/2) + 1.96\sqrt{n/4} \end{aligned}$$

If this happens, we will not reject the hypothesis $p = 1/2$ even though it is false. These are called, respectively, *errors of the first and second kinds*. Here, since we are setting up a method for routine application, we imagine that we will always make a definite assertion, or decision, every time. Hence we can talk about long-run rates of erroneous decisions. In nonroutine science, hypotheses do not stand or fall on the results of a statistical test, and this formulation

is less appropriate. The only distinction between significance tests and hypothesis tests seems to be Fisher’s insistence that there is one. Neither gives the weight of evidence for a hypothesis which is perhaps what one might like to have. Significance tests are usually derived or evaluated using the theory of hypothesis testing. It is somewhat ironic that, although Fisher was adamantly opposed to the latter, his tabulations of significant points and his later derivation of nonnull distributions set everything up for hypothesis-testing theory.

In their first paper (1928), Neyman and Pearson defined these errors and investigated an intuitive principle for test construction that yields likelihood-ratio (LR) tests [7]. The latter had appealed to them because Fisher had shown how to use likelihood to obtain “good” estimators and because it arises if one uses the Bayesian approach – which they were reluctant to do. The LR principle (to be explained later) was extraordinarily fruitful; it enabled them to derive all the tests known, including Pearson’s χ^2 , from one rule.

But the nagging question – what is the *best* test? – remained. In 1928, they had implicitly formulated their criterion: keep the rate of errors of the first kind constant, minimize the rate of errors of the second kind. By 1933, all the key definitions had been made to allow an orderly presentation whose centerpiece is the Neyman–Pearson lemma. The subject called *hypothesis testing* has been developed on this basis ever since; the classic text is that of Lehmann [9]. Various other subdisciplines split off as the years went by. The concern that tests should not be dependent on, or insensitive to, distributional assumptions respectively led to distribution-free or nonparametric and robust methods (*see Nonparametric Tests*). Efforts to minimize the cost or sample size in routine sampling led to sequential analysis. The vagueness of the significance level can be resolved if one takes into account the costs of making errors; this led to decision theory. These last two developments are largely due to A. Wald [10] and arose during and after World War II.

The theory of hypothesis testing gradually came to be interpreted rather more rigidly than its originators had in mind. This gave rise to many paradoxes. Practitioners tend to avoid these by less informal usage and *ad hoc* procedures. Some of the paradoxes are eliminated by recognizing that some problems are neither estimation nor testing but an incompletely

4 Hypothesis Testing

formulated activity to do with the discovery of statistical regularities in data. Other paradoxes do not appear in the Bayesian formulation and so have led to its revival and development (*see Overview of Statistics*). However, the key ideas of hypothesis testing will always remain central to the theory of statistics and its applications. We will now define the main concepts needed in the theory in terms of examples since orderly developments are easy to find at a variety of levels of abstraction (e.g., Lehmann [9] and Bickel and Doksum [11]). The paradoxical aspects may be seen in Cox and Hinkley [12]. Barnett [13] gives a general discussion of statistical inference.

The Neyman–Pearson Lemma

In the binomial and multinomial examples above, the data is a sample from a distribution fully specified on the null hypothesis; i.e., on the null hypothesis the probability of the sample is uniquely determined. However, there are many alternative hypotheses (e.g., $p \neq 1/2$ for the binomial case). We say then that the null hypothesis is *simple* and the alternative is *composite*.

Had the data been a sample from a Gaussian distribution with mean μ and variance σ^2 , a common null hypothesis is $\mu = 0$, with σ^2 unspecified. This null hypothesis is composite. To test $\mu = 0$, one instinctively thinks of rejecting the null hypothesis if $\bar{x}$ differs too much from zero. But use of the median makes sense, too. Which is better? Also, here we can speak of the probability density rather than the probability of the data; let “probability function” or “likelihood” denote either.

The problem may be stated fairly generally now. Suppose that the probability function of the data $\mathbf{x}$ is $p(\mathbf{x}; \theta)$, where $\mathbf{x}$ is a point in the *sample space* $\mathcal{X}$ and θ is a point in the parameter space Θ (*see Probability Theory*). Define the null hypothesis by $\theta \in \Theta_0 \subset \Theta$ and the alternative hypothesis by $\theta \in \Theta_1 \subset \Theta$, ($\Theta_0 \cap \Theta_1 = \emptyset$). In the binomial case $\mathcal{X} = \{0, 1, \dots, n\}$ and $\theta = p$, $\Theta = \{p | 0 \leq p \leq 1\}$, $\Theta_0 = \{1/2\}$, $\Theta_1 = \{p | p \neq 1/2\}$. We will sometimes use H_0 and H_1 for the null and alternative hypothesis. In the examples, we rejected the null hypothesis when $\mathbf{x}$ fell in a set such that some function, the test statistic, say $t(\mathbf{x})$, of the data took too extreme a value. This means that any test divides $\mathcal{X}$ into two mutually exclusive

regions; R , the rejection or critical region and its complement, the acceptance region. Choosing a test is choosing a test statistic $t(\mathbf{x})$ or equivalently a region R . For the test to have **Significance Level** α (e.g., 0.05), we require R such that

$$\Pr(\mathbf{x} \in R | \theta \in \Theta_0) \leq \alpha \quad (2)$$

This ensures that, although the left-hand side of equation (2) may vary as θ takes different values in Θ_0 , the probability of type 1 errors, or false alarms, never exceeds α . Now there will usually be many such regions R , so we want to find one that will maximize

$$\Pr(\mathbf{x} \in R | \theta \in \Theta_1) \quad (3)$$

This probability, for a fixed R , is a function of θ , the **Power** function of the test R . It is the probability of rejecting the null hypothesis, which is what we want to do when $\theta \in \Theta_1$, the alternative hypothesis. Hence maximizing equation (3) is the same as minimizing the probability of errors of the second kind.

To progress with this general formulation, we consider the simplest case when $\Theta_0 = \{\theta_0\}$, $\Theta_1 = \{\theta_1\}$ (i.e., both the null and alternative hypotheses are simple). Then the problem posed by equations (2) and (3) is as follows: find R such that

$$\int_R p(\mathbf{x}, \theta_0) d\mu(\mathbf{x}) \leq \alpha \quad (4)$$

$$\int_R p(\mathbf{x}, \theta_1) d\mu(\mathbf{x}) = \max \quad (5)$$

It is clearly advantageous to have equality in equation (4). This is usually possible except for discrete distributions (see below). Assume that there are many regions R satisfying equation (4) with equality. To achieve the maximum required in equation (5) we would want to choose points for R where $p(\mathbf{x}, \theta_1)$ is relatively large but where $p(\mathbf{x}, \theta_0)$ is relatively small. This suggests that the optimal R might be R , the set of points $\mathbf{x}$ where the ratio of $p(\mathbf{x}, \theta_1)$ to $p(\mathbf{x}, \theta_0)$ is greatest, i.e.,

$$\frac{p(\mathbf{x}, \theta_1)}{p(\mathbf{x}, \theta_0)} \geq k \quad (6)$$

where k is chosen so that equation (4) is satisfied. Then R and R both satisfy equation (5), but

$$\int_{R^*} p(\mathbf{x}, \theta_1) d\mu(\mathbf{x}) \geq \int_R p(\mathbf{x}, \theta_1) d\mu(\mathbf{x}) \quad (7)$$

since in the part of R not intersecting R , $p(\mathbf{x}, \theta_1) \geq kp(\mathbf{x}, \theta_0)$, while in the part of R not intersecting R , $p(\mathbf{x}, \theta_1) \leq kp(\mathbf{x}, \theta_0)$. This is the Neyman–Pearson lemma – the *most powerful test* of θ_0 versus θ_1 is based on the LR region (6).

Example 1 Let x be binomial (n, p) , the null hypothesis $p = p_0$, and the alternative hypothesis $p = p_1 > p_0$. The LR is

$$\frac{\binom{n}{x} p_1^x (1-p_1)^{n-x}}{\binom{n}{x} p_0^x (1-p_0)^{n-x}} = K \left(\frac{p_1/(1-p_1)}{p_0/(1-p_0)} \right)^x \quad (8)$$

Because $p_1 > p_0$, this ratio increases with x so that the critical region is $\{c, c+1, \dots, n\}$ when we choose c to make

$$\sum_c^n \binom{n}{x} p_0^x (1-p_0)^{n-x} \quad (9)$$

as near to α as possible. Had p_1 been less than p_0 , the critical region would have been in the lower tail of the distribution. Both tests are as intuition would suggest; $p_1 > p_0$ is a one-sided alternative hypothesis and we have a one-tailed test.

To get α exactly here, we must resort to *randomization*. For example, if $n = 10$, $p_0 = 1/2$, $\alpha = 0.05$, and $p_1 > 1/2$, observe that $\Pr(6 \text{ heads}) = 28/256$, $\Pr(7 \text{ heads}) = 8/256$, and $\Pr(8 \text{ heads}) = 1/256$. Consider the following rule:

If we get 6 heads, reject with probability δ ;
if we get 7 or 8 heads, reject with probability 1.

Then

$$\Pr(\text{reject} | p = \frac{1}{2}) = \frac{28\delta + 8 + 1}{256} = \frac{1}{20} \quad (10)$$

if $\delta = 19/140$. For this reason, the general theory does not use a critical region but a function $\phi(\mathbf{x})$, the probability of rejection given sample $\mathbf{x}$. In the simplest case, $\phi = 1$ or 0 and is the indicator function of the critical region.

The example shows another special feature. When the alternative to $p = p_0$ is $p_1 > p_0$, the *same* test is also optimal for all $p_1 > p_0$, and is said to be *uniformly most powerful* (UMP), clearly a most desirable property. However, had the alternative been $p \neq p_0$, a two-sided alternative, this argument fails. Clearly, we need to have two-tailed tests (i.e., some of each tail in the critical region), but we need a theory to tell us how much of each.

Finally, had the null hypothesis been $p \leq p_0$ (composite) and the alternative $p \geq p_1 > p_0$, the same test may be shown to be UMP.

Example 2 Let $X_1, \dots, X_n$ be independent and Gaussian (μ, σ^2) with *known* σ^2 , and the null and alternative hypotheses be $\mu = \mu_0$, $\mu = \mu_1$, respectively. The LR is

$$\frac{(\sqrt{2\pi}\sigma)^{-n} \exp\{-\sum (x_i - \mu_1)^2 / 2\sigma^2\}}{(\sqrt{2\pi}\sigma)^{-n} \exp\{-\sum (x_i - \mu_0)^2 / 2\sigma^2\}} \quad (11)$$

which reduces to

$$\exp\{(n/2\sigma^2)[(\mu_1 - \mu_0)\bar{x} - \mu_1^2 + \mu_0^2]\} \quad (12)$$

Thus, if $\mu_1 > \mu_0$, the ratio is large for large positive $\bar{x}$, so the critical region is $\bar{x} > c$. To calculate c we observe that on the null hypothesis $\mu = \mu_0$, $(\bar{x} - \mu_0)/(\sigma/\sqrt{n})$ is standard Gaussian, so the $\alpha = 0.05$ test is to reject if $\bar{x} > \mu_0 + 1.64\sigma/\sqrt{n}$. Once again this is a UMP test for all $\mu_1 > \mu_0$. The reader may observe that $\bar{x}$ is here a *sufficient statistic* for μ , so we would expect the best test to use it.

The power function of this test is, by definition,

$$\begin{aligned} \Pr(\bar{x} > \mu_0 + 1.64\sigma/\sqrt{n} | \mu) \\ = \Pr\left(z = \frac{\bar{x} - \mu}{\sigma/\sqrt{n}} > \frac{\mu_0 - \mu}{\sigma/\sqrt{n}} + 1.64\right) \end{aligned} \quad (13)$$

where $z = (\bar{x} - \mu)\sqrt{n}/\sigma$ now has a standard Gaussian distribution. As a function of μ , the power goes from zero at $\mu = -\infty$, to 0.05 at $\mu = \mu_0$, and to unity as μ increases to infinity.

A statistician is often asked: *How big a sample should be taken?* In the example above, if one can guess σ from previous experience, has settled on a 5% test, and is willing to pay what it costs to get a power of, say 0.90 to detect some particular μ , then equation (13) and the normal tables will provide the value of n that is needed (*see also Sample-Size Determination*). This value will increase as μ gets closer to μ_0 or as desired power is increased.

The power function (equation 13) for fixed n is a monotonically increasing function of μ . The power functions of other plausible tests (e.g., one that uses the sample median instead of the mean) will also increase. But at no value of μ can any of these curves ever rise above the curve given by equation (13)! It is the envelope of all possible power curves

6 Hypothesis Testing

for this problem. This illustrates graphically how our test is optimal. If there were no UMP test, no one power function would always be best (i.e., greatest for every μ).

Example 3 Let $X_1, \dots, X_n$ be independent and identically (i.i.d.) with the Cauchy density $\pi^{-1}(1 + (x - \theta)^2)^{-1}$; let the null hypothesis be θ_0 and the alternative be $\theta_1 > \theta_0$. The LR no longer simplifies – there is no single sufficient statistic for θ . The value of the L.R. for any sample depends on the values of both θ_0 and θ_1 . Finding the value of k to get a test of size α would probably involve applying the central limit theorem to the logarithm of the LR.

There is no UMP test – θ_1 always occurs in the test statistic.

A strategy often invoked in such cases is to seek the limit of the foregoing test as $\theta_1 \downarrow \theta_0$ – the *locally most powerful* (LMP) one-sided test. Using the general notation, the logarithm of the LR is

$$\log p(\mathbf{x}, \theta_1) - \log p(\mathbf{x}, \theta_0) \quad (14)$$

When θ is real, and θ_1 is near θ_0 , this is approximately proportional to

$$\partial \log p(\mathbf{x}, \theta_0) / \partial \theta_0 \quad (15)$$

and large values of this will be the LMP test of θ_0 *versus* values of θ just greater than θ_0 . Alternatively, if one thinks of all the possible power functions for tests of θ_0 *versus* $\theta > \theta_0$ and chooses the test whose slope is greatest at the origin, one again derives the test statistic equation (15). When the power of a test based on $\mathbf{R}$ is

$$\gamma(\theta) = \int_R p(\mathbf{x}, \theta) d\mu(\mathbf{x}) \quad (16)$$

the slope at θ is given by

$$\gamma'(\theta_0) = \int_R \frac{\partial p}{\partial \theta_0}(\mathbf{x}, \theta_0) d\mu(\mathbf{x}) \quad (17)$$

A similar argument to that leading to the Neyman–Pearson lemma tells us here that $\gamma'(\theta_0)$ will be a maximum, subject to $\gamma(\theta_0) = \alpha$ if R is based on large values of

$$\frac{1}{p(\mathbf{x}, \theta_0)} \frac{\partial p(\mathbf{x}, \theta_0)}{\partial \theta_0} = \frac{\partial}{\partial \theta_0} \log p(\mathbf{x}, \theta_0) \quad (18)$$

as we have just seen.

In the Cauchy case above, the LMP test is based on large values of

$$t(\mathbf{x}) = \sum_{i=1}^n \frac{x_i - \theta_0}{1 + (x_i - \theta_0)^2} \quad (19)$$

The asymptotic distribution of $t(\mathbf{x})$ is easily obtained from the central limit theorem.

Thinking again in terms of the graph of the power function, suppose that we have a two-sided alternative $\theta \neq \theta_0$ and attach equal importance to small deviations above and below θ_0 . It is then natural to seek a critical region R such that

$$\gamma(\theta_0) = \alpha, \gamma'(\theta_0) = 0, \gamma''(\theta_0) = \text{maximum} \quad (20)$$

A Neyman–Pearson type of argument then tells us that the optimal critical region is the set of points $\mathbf{x}$ such that

$$\frac{\partial^2 p(\mathbf{x}, \theta_0)}{\partial \theta_0^2} \geq k_1 \frac{\partial p(\mathbf{x}, \theta_0)}{\partial \theta_0} + k_2 p(\mathbf{x}, \theta_0) \quad (21)$$

Such a test is said to be an LMP *unbiased* test. If applied to Example 2, it leads to the 5% level test; reject if

$$\frac{|\sqrt{n}(\bar{x} - \theta_0)|}{\gamma} > 1.96 \quad (22)$$

the equal-tailed version of the UMP (and hence LMP) test for one-sided alternatives. If applied to Example 3, it does not lead to the two-sided version of the LMP test given above.

When only one parameter is involved, much more theory is available, but most real problems involve several parameters.

Several-Parameter Problems

Difficulties arise when dealing with problems involving several parameters. The null hypothesis is usually composite. To have a fixed significance level, we need a rejection region that is *similar* to the whole sample space in that its probability content is constant on H_0 . Parameters left unspecified by H_0 are *nuisance parameters*. Even if the null hypothesis is simple, the power function of a test will vary with all the parameters. One wants it to increase rapidly in all directions away from its value at θ_0 . But sacrificing slope in one direction

may increase it in another! This dilemma can be resolved only by weighing the importance of each direction, according to its consequences. These difficulties were recognized very early. As a practical matter, one may use an extension of the LR, or asymptotically equivalent procedures involving maximum-likelihood estimators. There is, however, a great deal of elegant theory of optimal tests when they exist.

Example 4 Let $X_1, \dots, X_n$ be independent and Gaussian (μ, σ^2) and let the null hypothesis be $\mu = \mu_0$ and the alternative $\mu \neq 0$. Here σ^2 is unspecified and so is a nuisance parameter. Ignoring constants,

$$\log \text{likelihood} = -n \log \sigma - \frac{1}{2\sigma^2} \sum (x_i - \mu)^2 \quad (23)$$

Maximizing over all values of μ we find $\hat{\mu} = \bar{x}$, and over all values of $\sigma^2 > 0$ it yields

$$\begin{aligned} \max_{H_1}(\log \text{likelihood}) \\ = -\frac{n}{2} \log \frac{\sum (x_i - \bar{x})^2}{n} - \frac{n}{2} \end{aligned} \quad (24)$$

If we set $\mu = \mu_0$ in equation (23) and maximize, we find that

$$\begin{aligned} \max_{H_0}(\log \text{likelihood}) \\ = -\frac{n}{2} \log \frac{\sum (x_i - \mu_0)^2}{n} - \frac{n}{2} \end{aligned} \quad (25)$$

Forming the log-likelihood ratio of Neyman and Pearson but after these maximizations (i.e., dividing equation 24 by equation 25), we get

$$-\frac{n}{2} \log \left(\frac{\sum (x_i - \bar{x})^2}{\sum (x_i - \mu_0)^2} \right) \quad (26)$$

Hence

$$\begin{aligned} + 2 \log \text{L.R.} \\ = -n \log \left\{ \frac{1}{1 + n(\bar{x} - \mu_0)^2 / \sum (x_i - \bar{x})^2} \right\} \end{aligned} \quad (27)$$

will be large if

$$t^2 = \frac{n(\bar{x} - \mu_0)^2}{\sum (x_i - \bar{x})^2 / (n-1)} \quad (28)$$

the square of Student's t , is large. Note that for large n , equations (27) and (28) are essentially equal. The null distribution of t does not contain σ^2 , so the t -test has similar rejection regions.

Had we done the same in the case of known σ^2 , we would have obtained

$$2 \log \text{L.R.} = n(\bar{x} - \mu_0)^2 / \sigma^2 \quad (29)$$

which on the null hypothesis has a χ^2 distribution with 1 degree of freedom. The same is asymptotically true of t^2 and so of equation (27). This illustrates *Wilks theorem*, which says roughly that as $n \rightarrow \infty$, $2 \log \text{LR}$ has a χ^2 distribution whose **degrees of freedom** are equal to the number of parameters specified by the null hypothesis.

The Pearson χ^2 test may be shown to be an approximation of an LR test. It also produces an answer when the null hypothesis specifies several parameters, unlike Examples 3 and 4.

Example 5 Reconsider Example 4. There are two unknown parameters μ and σ^2 and $\bar{x}$ and $s^2 = \sum (x_i - \bar{x})^2 / (n-1)$ are sufficient statistics for them. It follows from the theory of sufficiency that any good test will only use these two summaries of the data. Further, any test should be unchanged if we measure the x 's from a new origin and on a new scale (e.g., if they are temperatures, switch from Fahrenheit to Celsius). $\bar{x} - \mu$ and s^2 are independent of origin and $(\bar{x} - \mu)/s$ is also independent of scale. Thus, by demanding that our test should depend upon sufficient statistics and be *unchanged* under the natural group of transformations (here scale and location) which leave the problem unchanged, we arrive at the t -test. Among such invariant tests, we may now show that it has UMP properties, as in Example 2. Essentially, by invoking these general principles, we are able to reduce this problem so that the Neyman-Pearson lemma can be applied.

For a long time **nonparametric tests** remained outside the Neyman-Pearson theory – note that all the examples above are parametric. However, by requiring invariance under the group of symmetries of the problem, they were united with optimal **parametric tests**.

Example 6 Test that the cumulative distributions F and G are identical (i.e., $F(x) = G(x)$), given samples $(X_1, \dots, X_m)$ and $(Y_1, \dots, Y_n)$ from F and G .

When nothing is known of the common distribution, our test should be invariant under all monotonic transformations since they preserve the identity of F and G . Hence permissible tests can only depend on the ranks of the X and Y observations in the list of $m + n$ observations. Within this class of tests we can select members by Neyman–Pearson methods if we define specific alternatives to $F = G$; for example, that $G(x) = F(y - \Delta)$, $\Delta \geq 0$. If F is the logistic distribution, this line of argument may be used to prove that the Wilcoxon test is LMP.

One of the most surprising discoveries in the whole theory is that one need *not* pay a great price in power in order to buy a reliable significance level. To prove such assertions, one needs to introduce a definition of test efficiencies, e.g., asymptotic relative efficiency.

Example 7 Suppose that X_i are independently normal with means μ_i and variances σ^2 ($i = 1, \dots, n$) and that we have an independent estimate s^2 of σ^2 , distributed as $\sigma^2 \chi_f^2 / f$. To test a prespecified hypothesis $\sum_1^n C_i \mu_i = 0$ with $\sum_1^n C_i = 0$, we would refer $\sum C_i X_i / \left(s \sqrt{\sum C_i^2} \right)$ to the t -distribution with f degrees of freedom. If, however, we began inventing such hypotheses *after* we had looked at the data, this method would be improper (e.g., we might seek the largest difference $X_i - X_j$, which certainly is not normal with mean zero and variance two).

We may allow for all such possible comparisons of the μ_i 's by using the Cauchy inequality

$$\left\{ \sum_1^n C_i X_i = \sum C_i (X_i - \bar{X}) \right\}^2 \leq \sum_1^n C_i^2 \sum_1^n (X_i - \bar{X})^2 \quad (30)$$

so that

$$\frac{(\sum C_i X_i)^2}{\sum C_i^2 (s^2)} \leq \frac{\sum_1^n (X_i - \bar{X})^2}{s^2} \quad (31)$$

The left-hand side (lhs) is the square of our original statistic, and the rhs is $(n - 1)$ times a variable distributed as F with $n - 1$ and f degrees of freedom, which may thus be used to make a proper test. If, at the outset, we agree to look only

at differences of pairs, a different and more stringent procedure is available.

In the inspection of data one often wants to make a rough test to see if some aspect of it is exceptional, and similar care is always required, although it is rarely possible to formulate the problem as precisely as above. Tukey has called this the *multiplicity problem*.

Miscellaneous Remarks

The account above gives only the simplest cases of the basic ideas and is biased toward application. Through the references and other entries, especially **Parametric Tests; Nonparametric Tests; Normality Tests; Power; P Values; Significance Level; Overview of Statistics and Probability Theory**, the reader may find descriptions of those further aspects he or she seeks. The reader is fortunate that almost all the key papers have been collected, and most still make fascinating reading. Any such study will take the reader deeply into estimation theory and the mathematical structure of probability distributions.

On logical issues, not even hinted at here, Cox and Hinkley [12] is recommended. The role of conditionality is particularly puzzling. Neyman–Pearson theory is centered on maximizing power and uses all conceivable samples. But there are often aspects of the sample that do not depend on the parameter being tested, and a good case may be made for considering only all samples with the same such aspects as our given sample. But these two principles may be conflicting. One is often tempted to use conditional tests to eliminate nuisance parameter difficulties. One practical topic closely related to hypothesis testing but not even mentioned above is interval estimation (*see Confidence Intervals*). This arises when tests are examined to find the set of parameter values which cannot be rejected as null hypotheses. Lack of space prevents any discussion of robust tests.

Some of these logical conundrums are eliminated by adopting a Bayesian approach to inference, which, however, brings its own difficulties. Cox and Hinkley [12] give a brief discussion (*see also Overview of Statistics*). Box and Tiao [14] illustrate the Bayesian handling of many common statistical models. Box [15] attempts a partial synthesis of Bayesian and classical inference in which

significance tests are used for model checking (model to data) but Bayesian methods are used for estimation (data to model).

References

- [1] Pearson, K. (1900). On the criterion that a given system of deviations from the probable in the case of a correlated system of variables is such that it can be reasonably supposed to have arisen from random sampling, *Philosophical Magazine* **50**, 157–175. (The paper that introduced the chi-square goodness-of-fit statistic.)
- [2] Student. (1958). In *Student's Collected Papers*, E.S. Pearson & J. Wishart, eds, Cambridge University Press, Cambridge. (Papers 2 and 3 are references 1908a, 1908b; many of his papers have a modern ring.)
- [3] Pearson, K. (1892). *The Grammar of Science*, Adam and Charles Black, London
- [4] Fisher, R.A. (1925). *Statistical Methods for Research Workers*, Oliver & Boyd, Edinburgh. (New editions of this classic still appear.)
- [5] Pearson, E.S. (1966). In *Research Papers in Statistics*, F.N. David, ed, John Wiley & Sons, New York
- [6] Neyman, J. (1967). *A Selection of the Early Statistical Papers of J. Neyman*, University of California Press, Berkeley
- [7] Neyman, J. & Pearson, E.S. (1966). *Joint Statistical Papers of J. Neyman and E. S. Pearson*, University of California Press, Berkeley. (All 10 papers deal with the development of the basic notions of testing.)
- [8] Pearson, E.S. (1966). *The Selected Papers of E. S. Pearson*, Cambridge University Press, Cambridge. (A selection of papers illustrating Pearson's approach to statistical problems, including tests.)
- [9] Lehmann, E. (1959). *Testing Statistical Hypotheses*, John Wiley & Sons, New York. (This classic text has no modern sequel and is still the basis of all graduate courses on testing.)
- [10] Wald, A. (1955). *Selected Papers in Probability and Statistics*, McGraw-Hill, New York. (Wald introduced decision theory and sequential analysis and some of the modern mathematical styles of statistics, so one can see here the important first steps.)
- [11] Bickel, P.J. & Doksum, K.A. (1977). *Mathematical Statistics*, Holden-Day, San Francisco. (Introductory graduate text on classical statistic inference.)
- [12] Cox, D.R. & Hinkley, D.V. (1974). *Theoretical Statistics*, John Wiley & Sons, New York. (The most extensive discussion of examples and counterexamples of inference procedures.)
- [13] Barnett, V. (1973). *Comparative Statistical Inference*, John Wiley & Sons, New York. (A nonpartisan and largely nontechnical account of all methods of statistical inference putting hypothesis testing in context.)
- [14] Box, G.E.P. & Tiao, G.C. (1973). *Bayesian Inference in Statistical Analysis*, Addison-Wesley, Reading. (A very readable text leading to practical procedures. For philosophical and deeper issues the reader must go to their references.)
- [15] Box, G.E.P. (1980). Sampling and Bayes' inference in scientific modelling and robustness, *Journal of the Royal Statistical Society. Series A* **143**, 383–430

Further Reading

- Fisher, R.A. (1971). In *Collected Papers of R. A. Fisher*, J.H. Bennett, ed, University of Adelaide Press, Adelaide, Vols. 1–4
- Ferguson, T. (1967). *Mathematical Statistics*, Academic Press, New York. (A decision theoretic approach.)

G. S. WATSON

Article originally published in Encyclopedia of Statistical Sciences, 2nd Edition (2005, John Wiley & Sons, Inc.). Minor revisions for this publication by Jeroen de Mast.

Maximum Likelihood

Introduction

The term *maximum likelihood* refers to a general method of **Estimation** with important historical and practical significance. Consider a sample $y_1, y_2, \dots, y_n$ drawn independently from a distribution with density or probability function $f(y; \theta)$ with an unknown vector parameter θ . The likelihood function is defined as

$$L(\theta) = \prod_{i=1}^n f(y_i; \theta) \quad (1)$$

The maximum-likelihood estimate (MLE) is the value $\theta = \hat{\theta}$ which maximizes $L(\theta)$ over the set of all possible values for θ . In practice, it is usually more convenient to maximize the logarithm of the likelihood function, $l(\theta) = \ln L(\theta)$. From calculus, it follows that $\hat{\theta}$ satisfies the score equation $\partial l / \partial \theta = 0$.

Early references to the method of maximum likelihood are attributed to Gauss, Laplace, and Edgeworth [1]. However, the prominent English statistician R. A. Fisher is unquestionably responsible for popularizing the technique and for identifying many of its statistical properties [2].

The method has a strong heuristic appeal: choose as your parameter estimate the one which makes the observed data seem most likely. As it turns out, it is often the best estimate possible, particularly in large samples. Inference is made easy by the fact that MLEs are consistent and asymptotically normal under broad regularity assumptions, with a variance that can be estimated from the observed or expected information matrix. The MLE is also invariant under one-to-one transformations of the parameters, so to obtain the MLE of a transformation of the original parameters, one need only apply the transformation to the original MLE.

Optimal Properties

The following is a list of the most important properties of the MLE in large samples (i.e., $n \rightarrow \infty$):

1. Consistency

The MLE converges in probability to the true value of the parameter (i.e., for any $\epsilon > 0$, $\lim_{n \rightarrow \infty} P(|\hat{\theta}_n - \theta| > \epsilon) = 0$).

2. Asymptotic normality

The distribution of $n^{1/2}(\hat{\theta} - \theta)$ converges to a **normal distribution** with zero mean and a **covariance** matrix which is the inverse of the information matrix $\mathbf{I}$. From a practical point of view it is more convenient to say that $\hat{\theta}$ is approximately distributed as $N(\theta, \mathbf{I}^{-1}/n)$.

3. Asymptotic efficiency

The MLE is the best asymptotically normal estimate in terms of its variance in large samples. Put more precisely, if $\tilde{\theta}$ is another estimate such that $n^{1/2}(\tilde{\theta} - \theta) \xrightarrow{\mathcal{L}} N(\theta, \mathbf{C})$, where $\mathbf{C}$ is a fixed matrix, then $\mathbf{C} \geq \mathbf{I}^{-1}/n$ in the sense that $\mathbf{C} - \mathbf{I}^{-1}/n$ is a positive semidefinite matrix.

Regularity Conditions

It is important to be aware of the general regularity conditions for the optimal properties of the MLE. Most advanced statistics texts have a discussion of these conditions [3, 4], with the classic reference being Cramér [5]. The following three conditions are commonly given:

1. The observed data points $y_1, y_2, \dots, y_n$ are independently and identically distributed (i.i.d.) according to a density or probability function $f(y; \theta)$, where θ has finite dimension m . This condition is less restrictive than it may seem, given that y_i may be a vector. In fact, the method of maximum likelihood is popular and appropriate for many regression problems where the distributions of the data points are not identical, and many texts give less restrictive assumptions.
2. The underlying density or probability function is identifiable, that is, $f(y; \theta_1) = f(y; \theta_2)$ for all y implies that $\theta_1 = \theta_2$.
3. The density is “smooth” in the sense that f has derivatives up to the third order with finite expectation, and the information matrix

$$\mathbf{I} = E_{\theta} \left[\frac{\partial^2 \ln f(y; \theta)}{\partial \theta_j \partial \theta_k} \right], \quad j, k = 1, \dots, m \quad (2)$$

exists and is nonsingular. The latter conditions ensure that the asymptotic covariance matrix exists.

These conditions are satisfied for many models of interest, such as the binomial, normal, and Poisson distributions (see **Probability Density Functions**).

Maximum-Likelihood Calculations

Maximum-likelihood **estimation** involves finding a global maximum of a function of one or several parameters. For certain models, the solution can be expressed as a simple function of the data. However, more often the solution must be posed as a nonlinear optimization problem. Typically, it involves solving a system of nonlinear equations.

The main methods in use today are the Newton–Raphson (NR) or quasi-Newton (QN) methods, and the Fisher scoring (FS) algorithm [6]. The NR algorithm is based on an approximation of the log likelihood by a quadratic function through a Taylor series expansion of the score functions. To implement the NR algorithm, one has to provide second-order derivatives for the log-likelihood function – the so-called Hessian matrix. Initial values for the parameters are updated and the process is repeated until convergence is obtained. If the initial value for parameters is close enough to the maximum, the NR algorithm usually converges quickly. However, if the initial value is poorly chosen it may fail. In particular, the Hessian matrix can become a nonpositive-definite. Upon convergence, at the final iteration, the inverse of the Hessian matrix provides an approximation to the asymptotic covariance of the MLE. In the QN algorithm only the first derivatives are used, and the second derivatives are estimated on the basis of the results from previous iterations. The QN algorithm involves line search methods, that is, maximization of the log likelihood along a given ray in the parameter space.

The difference between the NR and FS algorithms is that the latter uses the expectation of the Hessian matrix rather than the Hessian matrix itself, as in the NR algorithm. There are two versions of the FS algorithm. In the first version, the expectation of the Hessian matrix is approximated as the sum of cross-products of first derivatives. In the second version, the exact calculated information matrix is used. Thus, to use this version of the FS algorithm one has to have a formula for the information matrix as a function of the parameters calculated prior to the maximization procedure.

Other methods are sometimes used for maximum-likelihood estimation. One that deserves special mention is the EM (expectation–maximization) algorithm. Certain likelihoods may be thought as involving missing data, with the most notable example being random effects models. The EM method works in this setting by maximizing the expectation of the log-likelihood iteratively for the complete data. This may aid in difficult maximization problems by taking advantage of simple closed-form solutions available for the “M” stage. Other advantages of the **EM algorithm** include its natural statistical interpretation and its property of producing an increasing sequence of log likelihood values in the specified parameter space. The principal drawback is that it may be relatively slow.

Example: Estimation of a Proportion

We observe the occurrence of a certain event for n individuals, where y_i is 1 if the event occurs and 0 otherwise. It can be assumed that events are independent among individuals and have the same probability of occurrence θ . The likelihood can be written as

$$L(\theta) = \prod_{i=1}^n \theta^{y_i} (1 - \theta)^{1-y_i} = \theta^m (1 - \theta)^{n-m} \quad (3)$$

where m is the number of events observed among the n individuals. The log likelihood is

$$l(\theta) = m \ln(\theta) + (n - m) \ln(1 - \theta) \quad (4)$$

To find the maximum, we consider the following score equation:

$$\frac{dl}{d\theta} = \frac{m}{\theta} - \frac{n - m}{1 - \theta} = 0 \quad (5)$$

which has the unique solution $\hat{\theta} = m/n$.

Discussion and Extensions

In theory, the method of maximum likelihood is very simple: determine an appropriate sampling distribution for the data, write down the likelihood as a function of the unknown parameters, and solve for the estimate. Of course, in practice it is not always so easy. Solutions to the likelihood score equations

usually must be arrived at by numerical methods, and may be computationally intensive or inaccurate. For some models and data sets it may be difficult to demonstrate that the likelihood has a unique global maximum. Many small-sample MLEs can be shown to be biased (*see* **Bias of an Estimator**). The presence of nuisance parameters can exacerbate the computational difficulties. Often the MLE is quite nonrobust to **Outliers** and, unlike competing general methods such as **Least-Squares Estimation** and the method of **moments** (*see* **Estimation**), computation of the MLE requires that the distribution of the data be completely parameterized. Uniformly minimum variance unbiased (UMVU) theory and Bayesian methods compete with maximum-likelihood estimation with regard to optimal properties, but also require complete specification of the distribution, and may be even harder to implement.

To overcome some of the difficulties of maximum-likelihood estimation, modified methods have been proposed such as restricted maximum likelihood (REML), conditional likelihood, pseudolikelihood, quasi-likelihood, partial likelihood, and M estimation. These methods were all inspired by the powerful

heuristic appeal and conceptual simplicity of the original formulation of the method of maximum likelihood.

References

- [1] Stigler, S.M. (1986). *The History of Statistics: The Measurement of Uncertainty Before 1900*, Harvard University Press, Cambridge.
- [2] Rao, C.R. (1992). R.A. Fisher: The founder of modern statistics, *Statistical Science* 7, 34–48.
- [3] Cox, D.R. & Hinkley, D.V. (1974). *Theoretical Statistics*, Chapman & Hall, London.
- [4] Stuart, H. & Ord, J.K. (1991). *Kendall's Advanced Theory of Statistics*, Oxford University Press, New York, Vol. 2.
- [5] Cramér, H. (1945). *Mathematical Methods of Statistics*, Princeton University Press, Princeton.
- [6] Thisted, R.A. (1988). *Elements of Statistical Computing*, Chapman & Hall, New York.

TOR D. TOSTESON AND EUGENE DEMIDENKO

Article originally published in Encyclopedia of Biostatistics, 2nd Edition (2005, John Wiley & Sons, Inc.). Minor revisions for this publication by Jeroen de Mast.

Missing Data and Imputation

Introduction

Missing data is a comprehensive term that comprises various aspects of incomplete information on the studied variables. An example of incomplete information is a longitudinal study of consumer satisfaction, with several surveys performed over a period. The data on consumer satisfaction is obtained by interviewing a sample from the customers' population.

1. Some of the customers from the intended sample may refuse to participate in the study. At the analysis stage, the information from those individuals will be missing on all the items and in all the subsequent surveys. This type of missing data is a *unit nonresponse* for all the surveys in the study.
2. Some customers may participate in initial studies but drop out from the study afterwards. In this attrition case with dropouts, the missing data is a still considered a *unit nonresponse* but only for the late surveys in which they do not participate.
3. In each particular survey, respondents may decide not to respond to specific question, e.g. when questions are sensitive, or the respondent ends the interview before its completion. The resulting data thus contains *item nonresponses*. Item nonresponse may also occur when respondents have no opinion or knowledge on issues, and the questions do not contain categories to cover such eventualities. Furthermore, if the respondents provide information that does not pass the logical tests (e.g. outside allowable range), the data will be again coded as *item nonresponse*.
4. Incomplete information may also result from *partial recording*. Examples of such partial recording are values known to be outside prespecified limits, or data with *interval censoring*, which are known to be between specific values.
5. In the **quality control** settings, *partial recording* may occur when the data are *batched* and the available information pertains to the entire batch rather than to the individual observations. Such partial recordings are missing data by design. A

different type of partial recording by design can occur when a subset of the data is subjected to more expensive measurements, while the values for the rest of the data are obtained by a simpler and more error-prone method of measurements.

In sample surveys, unit and item nonresponses can hardly be avoided and are an integral part of any study (see **Design and Testing of Questionnaires; Data Management in Survey Sampling**). In the quality control settings, efforts to avoid missing data can be more successful and typically rely on improved technology (see **Control Charts and Process Capability**). In multivariate process control, the problem is amplified since measurements can be derived from several measurement systems that have different levels of uncertainty and operate at difference cycle times (see **Multivariate Control Charts Overview; Statistical Process Control, Multivariate**). In clinical trials, participants who drop out of the study for any reason cause missing data (see **Case-Control Studies; Disease and Clinical Trial Modeling**). In designed experiments, some replicates or experimental runs might have not been executed causing another type of missing data (see **Factorial Experiments; Screening Designs; Response Surface Methodology**). Methods of dealing with missing data typically cannot overcome the impact of missing data and yield less accurate results than what can be achieved with datasets collected on all the intended units and items. Indeed, as mentioned by one of the leaders in the research on missing data "the most important step in dealing with missing data is to try to avoid it during the data-collection stage" [1].

In many cases, the transformation of the sets of data containing missing observations into a complete dataset precedes the analyses aimed at assessing the parameters of the data. The two basic methods for transforming the dataset into a rectangular array to be analyzed are as follows:

1. discard the incomplete cases and then analyze the units with complete information (complete-case analysis), or
2. impute estimated values instead of the missing ones and then analyze the filled in data.

In the longitudinal study of consumer satisfaction example, for each survey separately, missing data have to be considered in the analysis, usually by

2 Missing Data and Imputation

properly weighting the observed items. Alternatively, incomplete data can be analyzed by more sophisticated methods which do not require a complete dataset. In some cases, the common practice is to perform a basic comparison between the characteristics of the units (say, respondents) with complete observations and the characteristics of the units with incomplete data.

The pattern of the missing data defines the observed and the missing values in the dataset. Especially in longitudinal studies, when the missing data are due to attrition and dropout, the pattern of missing data contains important relevant information for analysis. Another important pattern is the univariate item nonresponse, i.e. missing data on a single variable with complete data on all the others.

The Process that Causes Missing Data and the Data Analysis

The process that causes missing data is related to the possible relationship between the fact that an observation is missing and the values of observations. When analyzing the data, while we consider the observed pattern of missing data, we almost always ignore the process that caused the data. As Rubin [2] mentioned, “When making sampling distribution inferences about the parameter of the data, θ , it is appropriate to ignore the process that causes missing data if the missing data are ‘missing at random’ and the observed data are ‘observed at random’, but these inferences are generally conditional on the observed pattern of missing data”.

The terms *missing at random* (MAR) and *missing completely at random* (MCAR) are fundamental in defining the possible relationship between the fact that an observation is missing and the values of observations. In this context, values of observations refer both to the values of the observed data (call them Y_{obs}) as well as to the potential values of the missing data (call them Y_{miss}). Let Y denote all the observations, observed or missing, with Y_{ij} being the variable j on the i th observation. The terms *MAR* and *MCAR* concern the relationships between those values and the distribution of the random matrix M of missing data indicators, (i.e., $m_{ij} = 1$ is variable j on observation i is missing, and 0 otherwise).

Missing data are MCAR, if the conditional distribution of M may depend on the unknown

parameters of the data Θ but not on Y_{obs} or Y_{miss} . Then, if the data are MCAR, the conditional distribution of M satisfies

$$f(M|Y, \Theta) = f(M|\Theta) \text{ for all } Y \text{ and } \Theta \quad (1)$$

In the consumer survey example, data on specific items will be MCAR if, say, individuals may decide by a random drawing whether to reply to some questions, with possible different probabilities of replying in the various subpopulations. Also, unit nonresponse may be completely at random, if again, individuals may decide by a random drawing whether to participate in the study, with possible different probabilities of participation in the various subpopulations.

Missing data are MAR, if the conditional distribution of M may depend on Θ and Y_{obs} but not on Y_{miss} . Thus, if the data are MAR, the conditional distribution of M satisfies

$$f(M|Y, \Theta) = f(M|Y_{\text{obs}}, \Theta) \text{ for all } Y_{\text{miss}} \text{ and } \Theta \quad (2)$$

In the longitudinal consumer study example, MAR will occur if say, the customers’ decision regarding participation in the survey at time t , depends on the grade that they assigned in the previous surveys, but not on the grades that they would have assigned at this survey. Again, the probabilities of participations may also vary in the various subpopulations. Obviously, since in the unit nonresponse there are no Y_{obs} , there is no MAR for an entire unit.

All other cases of missing data are not missing at random (NMAR). As an example, let two units have equal observations on all but one variable, with one unit having data and the other not having data on the remaining variable. If missing data are NMAR, the two units differ systematically on that remaining variable. In terms of the conditional distribution of M , the fact that the data is missing is affected by the unobserved value.

Handling of Missing Data – Listwise Deletion and Weighting

In most of the cases, the process that causes missing data is unknown. However, by using various methods of handling missing data, which ignore either only the process that causes missing data or both the process as well as the pattern of the observed data, the investigator implicitly assumes a certain mechanism that caused the missing data.

The simplest method of handling missing data is to discard the incomplete cases and then analyze the units with complete information “as is”. This complete-case analysis is thus preceded by what is known in the computer software as a *{listwise deletion}*. This method of handling missing data is indeed simple and quite common. As an example, in a study [3] the investigators collected data on more than 500 companies from the United Kingdom and aimed at creating an algorithm that can be used in the formulation of a knowledge-based decision system. They mentioned that “prior to conducting the statistical analysis, cases with missing data were eliminated from the dataset”.

This type of complete-case analysis ignores both the mechanism that created the missing data as well as the observed pattern of missing data. When the observed pattern of missing data is ignored, MAR is not a sufficient assumption. The required assumption in this case is MCAR. Under the MCAR mechanism, for each variable, the complete data can be considered as a random subsample from the intended sample. Moreover, even if we assume that the MCAR assumption holds, deletion of collected data in the complete-case analysis is a waste of information, which is rarely justifiable. Furthermore, if the data originates from a designed balanced experiment, the listwise deletion will create an undesirable unbalanced design.

A different type of complete-case analysis considers the pattern of the observed pattern of missing data and weighs differentially the respondents while still discarding the incomplete cases. The underlying assumption is that the data are MAR. The weights are usually based on poststratification, either direct or based on the response propensity stratification. In the direct poststratification the allocation of the observations to strata is based on the respondents’ recorded values (e.g., geographical or socioeconomic data). For each stratum, given the prior sampling probabilities, respondents’ weights are the inverse of the product of those probabilities and the estimated probability of participation in that stratum. In the consumer satisfaction example, with respondents grouped by the duration since becoming customers, the weight w_k for an observation in the k th group is given by $w_k = 1/(p_k r_k)$ where p_k is the prior sampling probability and r_k is the response rate. In many cases, in the sample design the sampling rates for strata are proportional to the population

rates and equal to the prior sampling probabilities. The corresponding weights will have only to account for the estimated probability of participation in each stratum.

When the weighting is based on the response propensity stratification, the definition of the strata is based on categories of variables (“adjustment cells”) with homogeneous estimated rates of response. The rates of response are estimated by regressing the indicator variable m_i (for response–nonresponse) on the background variables. Little [4] showed theoretically and in a simulation study that when weighting adjustments are employed (in contrast to imputation) adjustment cells based on the response propensity are effective in controlling bias.

Handling of Missing Data – Available-Case Analysis

The listwise deletion method discards information collected in the study and is thus inefficient. The available-case analysis is a relatively simple method to estimate the means and the correlations matrix from all the information available for estimating each parameter separately. Thus, the means of variable Y_k is estimated from all the cases with observed data on the k th variable, while the **correlation** between Y_k and Y_j is estimated from all the cases with observed data on the k th and the j th variable.

The advantage of making use of all the available information is counterbalanced by the fact that the sample base changes from the **estimation** of one correlation to the other, and especially in regression problems, when the data is highly correlated, simulations performed by Haitovsky [5] found that the available data analysis can yield considerably inferior results than the listwise deletion. Later simulations [6] found the available-case analysis superior to the listwise deletion in regression models with weak correlation. Little [1] mentions that the available-case method cannot be generally recommended.

Imputation – Unconditional and Conditional Mean Imputation

Imputation is the replacement of unknown, unmeasured, or missing data with a particular value. As in

the available-case analysis, all the incomplete cases are retained, and the dataset is transformed into a complete (rectangular) dataset by imputing values instead of the missing ones.

The simplest form of imputation is to replace all missing values with the average of that variable. It has been shown that, as expected, the unconditional mean imputation yields inconsistent estimates of parameters both for the variances as well as for the parameters based on the **covariance** matrix (e.g. Affifi and Elashoff [7]). By imputing a single average value for all the missing data, the resulting estimated variance clearly underestimates the actual parameter. Also, the **measures of association** are distorted. For this case as well, Little [1] mentions that the unconditional case imputation cannot be generally recommended.

The conditional mean imputation uses the correlation structure among observed variables. The value imputed for a missing value on a specific variable is the value resulting from the regression prediction equation obtained from the set of complete cases (CCs), with the variable to be predicted as the response variable and all the others as predictor variables. The estimates resulting from the conditional mean imputation improves the estimation of the parameters, which are linear in the data (like regression equations), but can still yield distorted estimates for variances, correlations, **quantiles**, and other parameters, which are not linear in the data.

Imputation – Hot Deck and Predictive Mean Matching

The hot-deck imputation strategy replaces the missing data on one observation with the available data from another matched observation. The method is used by many surveys in the official statistics. Among the versions of hot-deck imputation strategies, the univariate hot-deck strategy is probably still the most commonly used. Under this strategy, given a unit with missing data on one or more items, the method assigns for each variable separately, a value from a record with similar characteristics on the recorded variables. For example, the current population survey (CPS) performed by the Census Bureau [8, 9], use different hot decks for the households, demographic, labor force, industry and occupation, earning, and school enrollment edits. Hot decks are always defined

by age, race, and gender. The hot decks for the major labor force recode (MLR), which classifies adults as either employed, unemployed, or not in the labor force, are defined by age, race, and gender and, on occasion, by a related labor force characteristic. The number of cells in labor force item hot decks is relatively small, perhaps less than 100. On the other hand, the weekly earnings hot deck is defined by age, race, gender, usual hours, occupation, and educational attainment. This hot deck has several thousand cells. David *et al.* [10] compared the performance of the hot-deck imputation in the CPS data to the performance of a structured linear regression model.

The univariate hot-deck strategy handles many imputation problems well but handles more complex problems especially in the multivariate context poorly. One of the main problems is the shrunk correlation when the imputation on each variable is processed individually. As an example, Thibaudeau *et al.* [11] mention the distortion of the multivariate properties of the data obtained by using the univariate hot-deck strategy in the Wealth topical module for the Survey on Income and Program Participation (SIPP). In the sixth wave of the 1996 panel, the correlation coefficient between property value (assets) and mortgage(liability) was 0.39 in the reported (not imputed) data, but only 0.16 in the data imputed by the present univariate hot-deck strategy (i.e., less than half).

The joint hot-deck strategy preserves the correlation between jointly missing data, by imputing two or more items from the same donor. For a unit Y with missing data on several items whose correlation is to be analyzed (say, both assets and liabilities) the joint hot-deck imputation strategy finds a corresponding unit Z which matches Y as well as possible on the items on which both Y and Z are recorded. The missing items in Y are then filled in with the corresponding items in Z . The searching method is sequential. The search starts with all the items, and if no matching unit is found, the search continues for a subset of the most important variables or by reducing the number of categories in variables through collapsing.

The predictive mean imputation and the predictive mean matching can be conceptually considered as extensions to the hot-deck strategies by progressively expanding the number of cells while keeping at least one donor in each cell. In the analyses of paired variables (as assets and liabilities), for each missing observation on one out of the two variables, the

value imputed by the predictive mean imputation strategy is from a donor, such that both the donor and the recipient have matched values on the other variable. On the basis of results of tests performed, Thibaudeau *et al.* [11] suggest replacing gradually the currently used univariate hot-deck strategy in SIPP, first by a bivariate hot-deck strategy for specific pairs of variables, and then by a predictive mean methodology.

The predictive mean matching is a further extension, which selects the donor whose value is to be imputed instead of the missing data on the basis of a distance function rather than directly on the values of the covariates. For each variable Y and based on all other variables, the predictions on Y are computed from the complete data. For a recipient with missing value on Y , the set of potential donors are the units with recorded values on Y . The imputed value by predictive mean matching method is the recorded value of the nearest-neighbor potential donor, as computed from the distance between the donor and the recipient on the predictions on Y . In ordinary predictive mean matching the predictions are computed through a linear regression model.

The methods presented above implicitly assume that the missing data are MAR. Hot-deck procedures and models have been proposed also for cases when the missing data are NMAR. In the hot-deck procedures, instead of imputing the actual donor's value, for the various variables the values may be multiplied by appropriate inflation or factors. By definition, the observed data in a set with NMAR missing do not provide information to allow us to cross-validate the validity of the models' assumption. A cross-validation method that has been suggested and tested is to fit the models under various assumptions and to test the sensitivity of the resulting conclusions.

Multiple Imputation and Bootstrapping

While the biases in estimation caused by missing data are widely recognized, the imputation is not a universally accepted solution. In a glossary of the National Institute of Standards and Technology (<http://ciks.cbt.nist.gov>), it is mentioned that "... imputation is most common in surveys of human populations. It is also used in certain computer experiment applications". Even in "surveys of human populations" the imputation and estimation is not

taken lightly. The international standards code for market, opinion, and social research [12] requires that "... no data shall be imputed/assumed without the knowledge of the project manager. Comparison to the original source data shall be the first step in the process. Any imputation process shall be documented and available to client on request."

It is clear that the imputation raises concern with the practitioners. This concern can be attributed in part to the feeling that imputation is in some sense "inventing data" [1], and that even the sophisticated imputation methods cannot account for all the uncertainty in the actual value of the missing data.

Multiple imputation (MI) developed by Rubin [13] is a practical solution, which can alleviate the problem. For the estimation of an unknown parameter vector Θ , if we have a model for predicting the missing (or erroneous) values conditional on all observed data, we can use this model to make m independent simulated draws for the missing data, each containing different estimated values of the missing values from their predictive distribution. Rubin [13] showed that for a fixed number of draws, $m \geq 2$, the estimate of Θ obtained by averaging the individual estimates is a consistent estimator. Furthermore, a consistent estimate of the covariance matrix can be constructed from the averages of the individual covariance matrices estimated for each draw, and an estimate of the covariance of the estimates between the draws. In particular, the variance of the estimate is the average of the variances from the m draws plus $1 + 1/m$ times the sample variance of the estimates between the m datasets.

The MI can enhance the accuracy of estimated critical values and **confidence intervals** and bands as compared to other methods of dealing with missing data. An additional technique, which was compared in a series of articles with the MIs (e.g. [14]), and which also relies on repeated sampling from empirical datasets and associated estimates is bootstrapping. Both are computationally intensive methods and their development has been supported by the increase in computing **power**. Especially in econometrics, they contribute to the improvement of the estimates of precision of the statistical estimates when standard analytical estimates are biased or even unavailable. Furthermore, the bootstrapping and MIs can be combined, with the **bootstrap** replica being used in the MIs (e.g., Serrat *et al.* [15]).

In a recent industrial design of experiment [16] the bootstrap method was used to impute missing data at **center point** in an industrial **design of experiments**. Indeed, although “imputation is most common in surveys of human populations” the industrial experimental designs are also faced with missing or extra experimental runs, resulting in an unplanned nonrectangular dataset. Imputation can be highly effective in those circumstances.

Despite the concern related to “inventing data” in imputations, both theoretical results and simulation clearly demonstrate that, when performed properly, imputation can considerably improve the accuracy and the efficiency of the analyses. It is our opinion that they probably should be used more often than they are in quality control settings.

Parameter Estimation from Incomplete Data – The EM algorithm

The disadvantage of the imputation techniques mentioned above regarding “invented data” is alleviated in the direct analysis of the nonrectangular, incomplete dataset by the maximum-likelihood analysis. Let us assume that the underlying statistical model whose parameters Θ we have to estimate and the distribution of the data under the model are specified. Furthermore, let us assume that the missing pattern is ignorable, i.e. that in the maximum-likelihood approach, the parameter estimates do not depend on the pattern of the missing values and that pattern can thus be ignored. Little and Rubin [17] called those *ignorable models*.

Under those conditions the expectation-maximization (EM) algorithm [18], is a general and vastly used iterative algorithm, which computes the maximum-likelihoods estimates for Θ and in the intermediate steps performs pseudo-imputations as explained below. Specifically, at the completion of the t th iteration or step of the **EM algorithm**, the current estimated value of Θ is $\Theta^{(t)}$. The expectation step (E-step) $(t + 1)$ th step computes the expected value of the log-likelihood, over the conditional distribution of the missing data Y_{miss} , given the observed data Y_{obs} and the current estimate $\Theta^{(t)}$ of Θ . Now, if the distribution of the data is from the exponential family, the above expectation is equivalent to computing expected values for the sufficient statistics, given Y_{obs} and $\Theta^{(t)}$ [1, 19]. This is the pseudo-imputation

alluded to above. In the M-step of the algorithm, a new estimate $\Theta^{(t+1)}$ of Θ is calculated by maximizing with respect to Θ the expected value of the log-likelihood found in the E-step. The algorithm was proved to converge to the maximum-likelihood estimates of Θ . Since the algorithm was developed, many enhancements, which increase the rate of convergence for various problems with application in industry, have been proposed (e.g. Fessler and Hero [20, 21]).

In many situations, calculating the conditional expectation required in the E-step, may be infeasible. A possible alternative is the stochastic EM algorithm [22] in which the expectations are performed *via* simulation. At the $(t + 1)$ th iteration the missing data is imputed with a draw from the conditional distribution of the missing data given Y_{obs} and $\Theta^{(t)}$. Given the pseudocomplete data, one can maximize the log-likelihood to obtain $\Theta^{(t+1)}$.

The EM algorithm is widely used in a very large spectrum of applications, either in its direct version or in special versions and enhancements developed for particular situations (as the stochastic EM). Among the applications we can mention biomedical imaging, such as positron emission tomography [20, 21, 23, 24] as a robust inference in the multivariate **normal distribution** with missing data [25, 26]; log-linear models for contingency tables with missing data [27]; models for logistic regression with missing covariates [28]; and **survival analysis** with missing covariates [29].

Repeated Measurements in Longitudinal Studies

In the conclusion of this review we return to the example of a longitudinal consumer satisfaction study and present some issues specific to missing data in longitudinal studies. In many such studies (as in consumer satisfaction with a single response or in some clinical studies), the response variable is univariate, with the predictor variables assumed constant over time, and recorded at the baseline. More complex models assume variable predictors, which have to be recorded at each survey time separately. The missing data in longitudinal studies can come in various forms: from the common case of dropouts after one or more initial surveys, to more complex cases of incidental dropouts at no particular order, or even

more complex patterns of missing with unit nonresponse on some surveys and item nonresponse on others. For the sake of simplicity, we focus on the important case in longitudinal studies where the missing data is due to dropout.

The main approaches of dealing with missing data in longitudinal studies are equivalent to those used in single-time studies, with proper and important adjustments. The listwise deletion (or CC) approach performs the analysis on the subset with complete data and ignores the dropouts. The analysis yields valid estimates in the MCAR case, and yields biased estimates in all other cases.

Both single and multiple imputations are used in longitudinal studies. In the clinical trials the entire groups of subjects who started the study are known as the *intent-to-treat* (ITT) group. In those studies, one of the widely used single-imputation techniques (at least until recently) was the last observation carried forward (LOCF) technique. Under this procedure, the observation recorded before the dropout is “carried forward” and imputed for the missing in all subsequent times. LOCF is a model with simple assumptions behind it and it is indeed easy to understand and communicate.

However, the procedure does not give valid analyses if the dropout mechanism is not MCAR. Even then, it can yield biased estimates unless the pattern of response over time is constant, and if groups are compared, the pattern of dropout had to be the same for each group. The LOCF does not give a statistical estimator of any actual population parameter and yields underestimated variances. The procedure was considered conservative (and it is so, but only if a treatment is monotonically beneficial). A frequent argument in its favour is that it can be interpreted as “real world results”. However, this argument is not tenable since the results are likely to be misinterpreted as if they were from marginal inference, and the actual interpretation is very difficult to explain (e.g. Lane [30]). The performance of the procedure was compared in many studies under a variety of simulated types of missingness, and it was found to be under par even under most favorable situations (e.g. [30–34]). Indeed as Carpenter *et al.* [35] and Shao and Zhong [36] put it, “It is time to stop carrying it forward.”

Other single imputations, mentioned in the context of a single-wave study, as the hot deck and the predicted mean imputation are also used in longitudinal studies. Unlike the single-wave studies,

the data used for matching in the hot-deck procedure and for prediction in the predictive mean procedure, include the responses on the previous waves on the same variable.

The MI approach, which was also presented in the context of single-wave studies is effective in longitudinal studies as well. The MI approach carries out analysis under each set of imputation and combines analyses to reflect within-imputation and between-imputation variability. Several techniques involved in MI and used in longitudinal studies are mentioned by Rubin [37], Little and Rubin [38], Schafer and Olsen [39], and Schafer [40].

Assuming ignorable missing data as well as multivariate normal distribution of responses in longitudinal studies, the general mixed model can be used for analyzing the data with or without imputation. The model is also known under the name “mixed model for repeated measures” (MMRM). For the i th unit, $i = 1, \dots, n$, the model for the $(\tau \times p)$ response vector, $\mathbf{Y}_i$ is $\mathbf{Y}_i = \mathbf{X}_i\boldsymbol{\beta} + \mathbf{Z}_i\mathbf{b}_i + \boldsymbol{\varepsilon}_i$, where $\mathbf{b}_i \sim \mathbf{N}_q(\mathbf{0}, \mathbf{D})$, $\boldsymbol{\varepsilon}_i \sim \mathbf{N}_\tau(\mathbf{0}, \boldsymbol{\Sigma})$, and $\mathbf{b}_i$ and $\boldsymbol{\varepsilon}_i$ are statistically independent. $\mathbf{X}_i$ is a $(\tau \times p)$ design matrix for the fixed effects and $\mathbf{Z}_i$ is a $(\tau \times q)$ design matrix for the random effects. $\boldsymbol{\beta}$ is a $(p \times 1)$ vector of unknown fixed effects and $\mathbf{b}_i$ is an unknown $(q \times 1)$ random-coefficient vector (e.g. [41, 42]). Maximum-likelihood estimates for this model (e.g. [43]) can now be calculated either by standard software programs like PROC MIXED in SAS or by more specialized programs like HLM.

If the assumption of nonignorability is considered to be untenable, one can model the dropout as well as the response through selection models (SMs) or pattern-mixture models (PMMs). In the SMs, given the matrix of covariates X , the joint conditional distribution for the responses Y and matrix of missing M is

$$f(Y, M|X, \Theta, \Psi) = f(Y|X, \Theta)f(M|Y, X, \Psi) \quad (3)$$

where the model under no missing data is $f(Y|X, \Theta)$ and $f(M|Y, X, \Psi)$ is the model for the pattern of missing data.

In the PMMs the joint conditional distribution is

$$f(Y, M|X, \Theta, \Psi) = f(Y|X, M, \Theta)f(M|X, \Psi) \quad (4)$$

with the two models coinciding when M is independent of Y . The shortcoming of the models for dropout

is that the model assumptions cannot be tested from the data.

An issue of considerable practical importance is to assess theoretically or in simulations, the performance and robustness of various procedures both in the case when the underlying assumptions are met as well as when they are not met. Thus, one can generate various patterns of missing not at random (MNAR), and assess the robustness of the methods which implicitly assume a MAR mechanism. When the MMRM model is fitted, the fitting may or may not be preceded by a preliminary imputation step.

Several such studies were performed lately. The comparison between LOCF and MMRM [30] led to the conclusion that the use of the LOCF method is likely to seriously misrepresent the results of the study and that LOCF is “not a good choice for primary analysis”. In contrast, except for MNAR cases with substantially unequal dropouts in various groups, the MMRM is unlikely to result in serious misinterpretation. Sensitivity analysis with more complex models is required in dealing with such MNAR patterns of missing data. Simulation results obtained with MMRM, with and without imputations found the MIs to be a valuable technique, and superior to complete-case analysis [44]. The data were simulated under the MAR and MCAR mechanisms. Furthermore, the direct analysis of the incomplete dataset was effective as compared to the imputed data. The effectiveness of adding MIs to the MMRM and to the generalized estimating equations (GEE) was tested recently [45] in a simulated clinical trial setting aimed at determining the treatment effect at specific time epochs. They concluded that in this setting, the additional use of MIs is, in general, counterproductive both in terms of bias as well as in terms of variances.

The imputation procedures and MMRM modeling lately gained acceptance in industry (e.g. [30]). This trend is a positive result of the studies that showed that both imputations (especially MIs) and model fittings are highly beneficial in improving the accuracy of the performed analyses. The benefits to the performed studies far outweigh the shortcomings of the methods, i.e. the fact that imputation seem to be inventing data and the results of MMRM models are more difficult to communicate than analyses based either on CCs or on simple analysis of available data. The trend of increased acceptance of imputation and MMRM model fitting deserves to be encouraged.

References

- [1] Little, R.J.A. (1997). Biostatistical analysis with missing data, in *Encyclopedia of Biostatistics*, P. Armitage & T. Colton, eds, John Wiley & Sons, London.
- [2] Rubin, D.B. (1976). Inference and missing data, *Biometrika* **63**, 581–592.
- [3] Nookabadi, A.S. & Middle, J.E. (2001). A knowledge-based system as a design aid for quality assurance information systems, *International Journal of Quality and Reliability Management* **18**, 657–671.
- [4] Little, R.J.A. (1986). Survey nonresponse adjustments for estimates of means, *International Statistical Review* **54**, 137–139.
- [5] Haitovsky, Y. (1968). Missing data in regression analysis, *Journal of the Royal Statistical Society, Series B* **30**, 67–81.
- [6] Kim, J.O. & Curry, J. (1977). The treatment of missing data in multivariate analysis, *Sociological Methods and Research* **6**, 215–240.
- [7] Afifi, A. & Elashoff, R. (1967). Missing observations in: II. Point estimation in simple linear regression, *Journal of the American Statistical Association* **62**(317), 10–29.
- [8] Kostanich, D.L. & Dipbo, C.S. (supervisors) (2002). *The Current Population Survey: Design and Methodology*, Technical Paper 63RV, U.S. Bureau of the Census, Washington, DC.
- [9] Hanson, R.H. (1978). *The Current Population Survey: Design and Methodology*, Technical Paper 40, U.S. Bureau of the Census, Washington, DC.
- [10] David, M., Little, R.J.A., Samuhel, M.E. & Triest, R.K. (1986). Alternative methods for CPS income imputation, *Journal of the American Statistical Association* **81**, 29–41.
- [11] Thibaudeau, Y., Gottschalck, A. & Palumbo, T. (2006). *The Predictive-Mean Method of Imputation for Preserving Coupling Between Assets and Liabilities*, U.S. Census Bureau, Statistical Research Division Research Report Series.
- [12] ISO 20252 (2006). *Market, Opinion and Social Research – Vocabulary and Service Requirements*, International Organization for Standardization.
- [13] Rubin, D.B. (1987). *Multiple Imputation for Nonresponse in Surveys*, Wiley, New York.
- [14] Brownstone, D. & Valletta, R. (2001). The bootstrap and multiple imputations: harnessing increased computing power for improved statistical tests, *The Journal of Economic Perspectives* **15**, 129–141.
- [15] Serrat, C., Gomez, G., Garcia de Olallab, P. & Caylab, J. (1998). CD4+ lymphocytes and tuberculin skin test as survival predictors in pulmonary tuberculosis HIV-infected patients, *International Journal of Epidemiology* **27**, 703–712.
- [16] Kenett, R.S., Rahav, E. & Steinberg, D.M. (2006). Bootstrap analysis of designed, *Experiments Quality and Reliability Engineering International* **22**, 659–667.

- [17] Little, R.J.A. & Rubin, D.B. (1989). Missing data in social science data sets, *Sociological Methods and Research* **18**, 292–326.
- [18] Dempster, A.P., Laird, N.M. & Rubin, D.B. (1977). Maximum likelihood from incomplete data via the EM algorithm, *Journal of the Royal Statistical Society* **B39**, 1–38.
- [19] Sundberg, R. (1974). Maximum likelihood theory for incomplete data from an exponential family, *Scandinavian Journal of Statistics* **1**, 49–58.
- [20] Fessler, J.A. & Hero, A.O. (1994). Space-alternating generalized expectation-maximization algorithm, *IEEE Transactions on Signal Processing* **42**, 2664–2677.
- [21] Fessler, J.A. & Hero, A.O. (1995). Penalized maximum-likelihood image reconstruction using space alternating generalized expectation-maximization algorithm, *IEEE Transactions on Image Processing* **4**, 1417–1438.
- [22] Celuex, G. & Diebolt, J. (1985). The SEM algorithm: a probabilistic teacher algorithm derived from the EM algorithm for the mixture problems, *Computational Statistics Quarterly* **2**, 73–82.
- [23] Lange, F. & Carson, R. (1984). EM reconstruction algorithms for emission and transmission tomography, *Journal of Computer-Assisted Tomography* **8**, 306–316.
- [24] Shepp, L.A. & Vardi, Y. (1982). Maximum likelihood reconstruction for emission tomography, *IEEE Transactions on Image Processing* **2**, 113–122.
- [25] Lange, K., Little, R.J.A. & Taylor, J. (1989). Robust statistical inference using the T distribution, *Journal of the American Statistical Association* **84**, 881–896.
- [26] Little, R.J.A. (1988). Robust estimation of the mean and covariance matrix from data with missing values, *Applied Statistics* **37**, 23–38.
- [27] Fuchs, C. (1982). Maximum likelihood estimation and model selection in contingency tables with missing data, *Journal of the American Statistical Association* **77**, 270–278.
- [28] Vach, W. (1994). *Logistic Regression with Missing Values in the Covariates*, Springer-Verlag, New York.
- [29] Schluchter, M.D. & Jackson, K.L. (1989). Loglinear analysis of censored survival data with partially observed covariates, *Journal of the American Statistical Association* **84**, 42–52.
- [30] Lane, P. (2007). Handling drop-out in longitudinal clinical trials: a comparison of the LOCF and MMRM approaches, *Pharmaceutical Statistics*, Wiley InterScience.
- [31] Heyting, A., Tolboom, J.T.B.M. & Essers, J.G.A. (1992). Statistical handling of drop-outs in longitudinal clinical trials, *Statistics in Medicine* **11**, 2043–2061.
- [32] Mallinckrodt, C.H., Clark, W.S. & David, S.R. (2001). Type I error rates from mixed effects model repeated measures compared with fixed effects ANOVA with missing values imputed via LOCF, *Drug Information Journal* **35**, 1214–1225.
- [33] Mallinckrodt, C.H., Kaiser, C.J., Watkin, J.G., Molenberghs, G. & Carroll, R.J. (2004). The effect of correlation structure on treatment contrasts estimated from incomplete clinical trial data with likelihood-based repeated measures compared with last observation carried forward ANOVA, *Clinical Trials* **1**, 477–489.
- [34] Cook, R.J., Zeng, L. & Yi, G.Y. (2004). Marginal analysis of incomplete longitudinal binary data; a cautionary note on LOCF imputation, *Biometrics* **60**, 820–828.
- [35] Carpenter, J., Kenward, M., Evans, S. & White, I. (2004). Last observation carry forward and last observation analysis (letter to the editor), *Statistics in Medicine* **23**, 3241–3244.
- [36] Shao, J. & Zhong, B. (2003). Last observation carry forward and last observation analysis, *Statistics in Medicine* **22**, 2429–2441.
- [37] Rubin, D.B. (1986). Statistical matching and file concatenation with adjusted weights and multiple imputations, *Journal of Business and Economic Statistics* **4**, 87–94.
- [38] Little, R.J.A. & Rubin, D.B. (1987). *Statistical Analysis with Missing Data*, John Wiley & Sons, New York.
- [39] Schafer, J.L. & Olsen, M.K. (1998). Multiple imputation for multivariate missing-data problems: a data analyst's perspective, *Multivariate Behavioral Research* **33**, 545–571.
- [40] Schafer, J.L. (1997). *Analysis of Incomplete Multivariate Data*, Chapman & Hall, London.
- [41] Molenberghs, G., Kenward, M.G. & Lesaffre, E. (1977). The analysis of longitudinal ordinal data with nonrandom drop-out, *Biometrika* **84**(1), 33–44.
- [42] Verbeke, G. & Molenberghs, G. (2000). *Linear Mixed Models for Longitudinal Data*, Springer, New York.
- [43] Laird, N.M. & Ware, J.H. (1982). Random-effects models for longitudinal data, *Biometrics* **38**, 963–974.
- [44] Shieh, Y.-Y. (2003). *Imputation Methods on General Linear Mixed Models of Longitudinal Studies*, Federal Committee on Statistical Methodology (FCSM), Washington, DC.
- [45] Fong, D.Y.T., Rai, S.N. & Lam, K.S.L. (2006). *Use of Multiple Imputation on Linear Mixed Model and Generalized Estimating Equations for Longitudinal Data Analysis: A Simulation Study*, SCRA 2006-FIM XIII, Lisbon-Tomar, Portugal.

Further Reading

- Cook, R. (1986). Assessment of local influence, *Journal of the Royal Statistical Society. Series B (Statistical Methodology)* **48**, 133–169.
- Lange, K. (1995a). A gradient algorithm locally equivalent to the EM algorithm, *Journal of the Royal Statistical Society. Series B* **57**, 425–437.
- Lange, K. (1995b). A quasi-Newtonian acceleration of the EM algorithm, *Statistica Sinica* **5**, 1–18.
- Molenberghs, G. & Verbeke, G. (2005). *Models for Discrete Longitudinal Data*, Springer, New York.

CAMIL FUCHS AND RON S. KENETT

Multivariate Analysis

Introduction

The body of statistical methodology used to analyze simultaneous measurements on many variables is called *multivariate analysis*. Many multivariate methods are based on an underlying probability model known as the *multivariate normal* (see **Normal Distribution**). Other methods are *ad hoc* in nature and are justified by logical or common-sense arguments. Regardless of their origin, multivariate techniques invariably must be implemented on a computer. Current sophisticated statistical software packages make the implementation step easier.

Multivariate analysis is a “mixed bag”. It is difficult to establish a classification scheme for multivariate techniques that is both widely accepted and also indicates the appropriateness of the techniques. One classification distinguishes techniques designed to study interdependent relationships from those designed to study dependent relationships. Another classifies techniques according to the number of populations and the number of sets of variables being studied. Yet another classification may distinguish those methods applicable to metric data from those applicable to nonmetric data. This chapter is divided into sections according to inferences about means, inferences about **covariance** structure, and techniques for classification or grouping. This should not, however, be regarded as an attempt to place each method into a slot. Rather, the choice of methods and the types of analyses employed are determined largely by the objectives of the investigation.

The objectives of scientific investigations, for which multivariate methods most naturally lend themselves, include the following:

1. data reduction or structural simplification,
2. sorting and grouping,
3. investigation of the dependence among variables,
4. forecasting and prediction, and
5. hypothesis construction and testing.

A Historical Perspective

Many current multivariate statistical procedures were developed during the first half of the twentieth century. A reasonably complete list of the developers

would be voluminous. However, a few individuals can be cited as making important initial contributions to the theory and practice of multivariate analysis. Francis Galton and Karl Pearson did pioneering work in the areas of *correlation* (see **Correlation**) and regression analysis. R. A. Fisher’s derivation of the exact distribution of the sample correlation coefficient and related quantities provided the impetus for the multivariate distribution theory. C. Spearman and K. Pearson were among the first to work in the area of factor analysis. Significant contributions to multivariate analysis were made during the 1930s by S. S. Wilks (general procedures for testing certain multivariate hypotheses), H. Hotelling (see **Hotelling’s T^2**), principal component analysis, canonical correlation analysis, R. A. Fisher (discrimination and classification), and P. C. Mahalanobis (generalized distance, **hypothesis testing**). J. Wishart derived an important joint distribution of sample variances and covariances that bears his name. Later, M. S. Bartlett and G. E. P. Box contributed to the large-sample theory associated with certain multivariate test statistics.

Many multivariate methods evolved in concert with the development of electronic computers. Specifically, ingenious graphical methods for displaying multivariate data (e.g., Chernoff faces, Andrews plots) can only be conveniently implemented on a computer. Multidimensional scaling and many clustering procedures were not feasible before the advent of fast computers. Several people have exploited the power of the computer to develop procedures for extracting information from very large data sets, and to create informative lower-dimensional representations of multivariate data (see, for example [1], and the references therein).

Notation

The description of multivariate data and the computations required for their analysis are greatly facilitated by the use of matrix algebra. Consequently, the subsequent discussion will rely heavily on the following notation:

$\mathbf{X}$, a $p \times 1$ random vector,

$\mathbf{x}_j$, a $p \times 1$ multivariate observation on $\mathbf{X}$

2 Multivariate Analysis

$$\mathbf{X} = \begin{bmatrix} x_{11} & x_{12} & \cdots & x_{1p} \\ x_{21} & x_{22} & \cdots & x_{2p} \\ \vdots & \vdots & \cdots & \vdots \\ x_{n1} & x_{n2} & \cdots & x_{np} \end{bmatrix}$$

$$= [\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_p] \quad (1)$$

an $n \times p$ matrix. (Each column of $\mathbf{X}$ represents a multivariate observation.)

$$\bar{\mathbf{x}} = [\bar{x}_1, \bar{x}_2, \dots, \bar{x}_p]'$$

$$= \left\{ \bar{x}_i = \frac{1}{n} \sum_{j=1}^n x_{ij}; i = 1, 2, \dots, p \right\} \quad (2)$$

a $p \times 1$ vector of sample means

$$\mathbf{S} = \left\{ s_{ik} = \frac{1}{n-1} \sum_{j=1}^n (x_{ij} - \bar{x}_i)(x_{kj} - \bar{x}_k); \right.$$

$$i, k = 1, 2, \dots, p \left. \right\} \quad (3)$$

a $p \times p$ symmetric matrix of sample variances and covariances.

$$\mathbf{R} = \left\{ r_{ik} = \frac{s_{ik}}{\sqrt{s_{ii}}\sqrt{s_{kk}}} \right\} \quad (4)$$

a $p \times p$ symmetric matrix of sample correlation coefficients

$$\boldsymbol{\mu} = \{\mu_i\} = E(\mathbf{X}) \quad (5)$$

a $p \times 1$ vector of population means. ($E(\cdot)$ is the expectation (see **Expectation**) operator.)

$$\boldsymbol{\Sigma} = \{\sigma_{ij}\} = E(\mathbf{X} - \boldsymbol{\mu})(\mathbf{X} - \boldsymbol{\mu})' \quad (6)$$

a $p \times p$ symmetric matrix of population variances and covariances.

$$\boldsymbol{\rho} = \left\{ \rho_{ij} = \frac{\sigma_{ij}}{\sqrt{\sigma_{ii}}\sqrt{\sigma_{jj}}} \right\} \quad (7)$$

a $p \times p$ symmetric matrix of population correlation coefficients.

$$\mathbf{Z} = \left\{ Z_i = \frac{X_i - \mu_i}{\sqrt{\sigma_{ii}}} \right\} \quad (8)$$

a $p \times 1$ vector of standardized variables.

Multivariate Normal Distribution

A generalization of the familiar bell-shaped normal density to several dimensions plays a fundamental role in multivariate analysis. In fact, many multivariate techniques assume that the data were generated from a *multivariate* normal distribution. Although real data are never *exactly* multivariate normal, the normal density is often a useful approximation to the “true” population distribution. Therefore, the *normal distribution* (see **Normal Distribution**) serves as a *bona fide* population model in some instances. Also, the sampling distributions of many multivariate statistics are approximately normal, regardless of the form of the parent population, because of a *central limit effect*.

The p -dimensional normal density for the random vector $\mathbf{X} = [X_1, X_2, \dots, X_p]'$, evaluated at the point $\mathbf{x} = [x_1, x_2, \dots, x_p]'$, is given by

$$f(\mathbf{x}) = (2\pi)^{-p/2} |\boldsymbol{\Sigma}|^{-1/2}$$

$$\times \exp \left\{ -\frac{1}{2} (\mathbf{x} - \boldsymbol{\mu})' \boldsymbol{\Sigma}^{-1} (\mathbf{x} - \boldsymbol{\mu}) \right\} \quad (9)$$

where $-\infty < x_i < \infty, i = 1, 2, \dots, p$. Here $\boldsymbol{\mu}$ is the population mean vector, $\boldsymbol{\Sigma}$ is the population's variance-covariance matrix, $|\boldsymbol{\Sigma}|$ is the determinant of $\boldsymbol{\Sigma}$, and $\exp(\cdot)$ stands for the exponential function. We denote this p -dimensional normal density by $N_p(\boldsymbol{\mu}, \boldsymbol{\Sigma})$.

Contours of constant density for the p -dimensional normal distribution are ellipsoids defined by $\mathbf{x}$ such that

$$(\mathbf{x} - \boldsymbol{\mu})' \boldsymbol{\Sigma}^{-1} (\mathbf{x} - \boldsymbol{\mu}) = c^2 \quad (10)$$

These ellipsoids are centered at $\boldsymbol{\mu}$ and have axes $\pm c\sqrt{\lambda_i} \mathbf{e}_i$, where $\boldsymbol{\Sigma} \mathbf{e}_i = \lambda_i \mathbf{e}_i, i = 1, 2, \dots, p$. That is, $\lambda_i, \mathbf{e}_i$ are the eigenvalue (normalized) eigenvector pairs associated with $\boldsymbol{\Sigma}$.

The following are true for a random vector $\mathbf{X}$ having a multivariate normal distribution.

1. Linear combinations of the components of $\mathbf{X}$ are normally distributed.
2. All subsets of the components of $\mathbf{X}$ have a (multivariate) normal distribution.
3. Zero covariance implies that the corresponding components are distributed independently.
4. The conditional distributions of the multivariate components are (multivariate) normal.

(See references [2, Chapter 2] and [3, Section 8a] for more discussion of the multivariate normal distribution.) Because these properties make the normal distribution easy to manipulate, it has been overemphasized as a population model.

To some degree, the *quality* of inferences made by some multivariate methods depend on how closely the true parent population resembles the multivariate normal form. It is imperative, then, that procedures exist for detecting cases in which the data exhibit moderate to extreme departures from what is expected under multivariate normality.

Sometimes nonnormal data can be made to look more normal by considering transformations (*see Box and Cox Transformation*) of the data. Normal theory analyses can then be carried out with the suitably transformed data [4, Section 4.8; 5, Section 4.2]. Latest advances in the theory of *discrete* multivariate analysis are contained in reference [6].

Sampling Distributions of $\bar{\mathbf{X}}$ and $\mathbf{S}$

The tentative assumption that the columns of the data matrix, treated as random vectors, constitute a random sample from a normal population, with mean $\boldsymbol{\mu}$ and covariance $\boldsymbol{\Sigma}$, completely determines the sampling distribution of $\bar{\mathbf{X}}$ and $\mathbf{S}$. We now summarize the sampling distribution results.

Let $\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_n$ be a random sample of size n from a p -variate *normal* distribution with mean $\boldsymbol{\mu}$ and covariance matrix $\boldsymbol{\Sigma}$. Then (see [2, Sections 3.3 and 7.2; 3, Section 8b]):

1. $\bar{\mathbf{X}}$ is distributed as $N_p(\boldsymbol{\mu}, (1/n)\boldsymbol{\Sigma})$.
2. $(n-1)\mathbf{S}$ is distributed as a Wishart random matrix with $n-1$ **degrees of freedom** (df).
3. $\bar{\mathbf{X}}$ and $\mathbf{S}$ are independent.

Because $\boldsymbol{\Sigma}$ is unknown, the distribution of $\bar{\mathbf{X}}$ cannot be used directly to make inferences about $\boldsymbol{\mu}$. However, $\mathbf{S}$ provides independent information about $\boldsymbol{\Sigma}$ and the distribution of $\mathbf{S}$ does not depend on $\boldsymbol{\mu}$. This allows one to construct a statistic for making inferences about $\boldsymbol{\mu}$.

Selected Problems about Means

Single-Population Mean Vector

An immediate objective in many multivariate studies is to make statistical inferences about population

mean vectors. As an initial example, consider the problem of testing whether a multivariate normal population mean vector has a particular value $\boldsymbol{\mu}_0$.

Let the null and alternative hypotheses be $H_0 : \boldsymbol{\mu} = \boldsymbol{\mu}_0$ and $H_1 : \boldsymbol{\mu} \neq \boldsymbol{\mu}_0$, respectively. Once the sample is in hand, a sample mean vector far from $\boldsymbol{\mu}_0$ tends to discredit H_0 . A test of H_0 based on the (statistical) distance of $\bar{\mathbf{x}}$ from $\boldsymbol{\mu}_0$ (assuming the population covariance matrix $\boldsymbol{\Sigma}$ is unknown) can be carried out using Hotelling's T^2 ,

$$T^2 = n(\bar{\mathbf{x}} - \boldsymbol{\mu}_0)' \mathbf{S}^{-1} (\bar{\mathbf{x}} - \boldsymbol{\mu}_0) \quad (11)$$

Hotelling's T^2 is a distance measure that takes account of the joint variability of the p measured variables. Under H_0 , T^2 is distributed as

$$[(n-1)p/(n-p)]F_{p, n-p} \quad (12)$$

where n is the **sample size** and $F_{p, n-p}$ denotes an F random variable with p and $n-p$ df. Let $F_{p, n-p}(\alpha)$ be the upper 100α th percentage point of this F distribution (*see Probability Density Functions*). The hypothesis $H_0 : \boldsymbol{\mu} = \boldsymbol{\mu}_0$ is rejected at the α level of significance if

$$T^2 > [(n-1)p/(n-p)]F_{p, n-p}(\alpha) \quad (13)$$

Example 1 The quality-control department of a microwave oven manufacturer is concerned about the radiation emitted by the ovens. They record radiation measurements with the oven doors opened and closed. Measurements for a random sample of microwave ovens could then be compared with the standards for radiation emission set by the manufacturer.

Other principles of test construction (e.g., likelihood ratio, union–intersection principle) also lead to the use of Hotelling's T^2 in this testing situation.

Hotelling's T^2 statistic has numerous other applications. There are multivariate analogs of the paired and two-sample univariate t statistics. For applications, see references [4, Chapter 5; 7, Chapter 2; 8, Chapter 5].

Several-Population Means, Multivariate Analysis of Variance* (MANOVA)

Multivariate analysis of variance (MANOVA) is concerned with inferences about several population means. It is a direct generalization of the *analysis*

4 Multivariate Analysis

of variance (ANOVA) (see **Analysis of Variance**) to the case of more than one response variable. In its simplest form, one-way MANOVA, random samples are collected from each of g populations and arranged as

$$\begin{array}{llll} \text{Population 1} & \mathbf{X}_{11}, \mathbf{X}_{12}, & \dots, & \mathbf{X}_{1n_1} \\ \text{Population 2} & \mathbf{X}_{21}, \mathbf{X}_{22}, & \dots, & \mathbf{X}_{2n_2} \\ \vdots & \vdots & & \vdots \\ \text{Population } g & \mathbf{X}_{g1}, \mathbf{X}_{g2}, & \dots, & \mathbf{X}_{gn_g} \end{array}$$

MANOVA is used first to investigate whether the population mean vectors are the same, and if not, which mean components differ significantly. It is assumed that

1. $\mathbf{X}_{l1}, \mathbf{X}_{l2}, \dots, \mathbf{X}_{ln_l}$ is a random sample of size n_l from a population with mean $\boldsymbol{\mu}_l$, $l = 1, 2, \dots, g$. The random samples from different populations are independent.
2. All populations have a common covariance matrix $\boldsymbol{\Sigma}$.
3. Each population is multivariate normal.

Condition 3 can be relaxed by applying the central limit theorem when the sample sizes n_l are large. The model states that the mean $\boldsymbol{\mu}_l$ consists of a common part $\boldsymbol{\mu}$ plus an amount $\boldsymbol{\tau}_l$ due to the l th treatment. According to the model, *each component* of the observation vector $\mathbf{X}_{lj}$ satisfies the univariate

model. The errors for the components of $\mathbf{X}_{lj}$ are correlated, but the covariance matrix $\boldsymbol{\Sigma}$ is the same for all populations (see Table 1).

A vector of observations may be decomposed as suggested by the model. Thus,

$$\begin{aligned} \mathbf{x}_{lj} &= \bar{\mathbf{x}} + (\bar{\mathbf{x}}_l - \bar{\mathbf{x}}) + (\mathbf{x}_{lj} - \bar{\mathbf{x}}_l) \\ (\text{observation}) &= \begin{bmatrix} \text{overall} \\ \text{sample} \\ \text{mean } \hat{\boldsymbol{\mu}} \end{bmatrix} + \begin{bmatrix} \text{estimated} \\ \text{treatment} \\ \text{effect } \hat{\boldsymbol{\tau}}_l \end{bmatrix} \\ &\quad + \begin{pmatrix} \text{residual} \\ \hat{\mathbf{e}}_{lj} \end{pmatrix} \end{aligned} \quad (14)$$

which leads to a decomposition of the **sum of squares** and cross-products matrix

$$\sum_{l=1}^g \sum_{j=1}^{n_l} (\mathbf{x}_{lj} - \bar{\mathbf{x}})(\mathbf{x}_{lj} - \bar{\mathbf{x}})' \quad (15)$$

and the MANOVA table (see Table 2). One test of $H_0: \boldsymbol{\tau}_1 = \boldsymbol{\tau}_2 = \dots = \boldsymbol{\tau}_g = \mathbf{0}$ involves *generalized variances*. We reject H_0 if the ratio of generalized variances

$$\begin{aligned} \Lambda &= \frac{|\mathbf{W}|}{|\mathbf{B} + \mathbf{W}|} \\ &= \left| \frac{\sum_{l=1}^g \sum_{j=1}^{n_l} (\mathbf{x}_{lj} - \bar{\mathbf{x}}_l)(\mathbf{x}_{lj} - \bar{\mathbf{x}}_l)'}{\sum_{l=1}^g \sum_{j=1}^{n_l} (\mathbf{x}_{lj} - \bar{\mathbf{x}})(\mathbf{x}_{lj} - \bar{\mathbf{x}})'} \right| \end{aligned} \quad (16)$$

Table 1 MANOVA model for comparing g population mean vectors

$\mathbf{X}_{lj} = \boldsymbol{\mu} + \boldsymbol{\tau}_l + \mathbf{e}_{lj}$, $j = 1, 2, \dots, n_l$, and $l = 1, 2, \dots, g$, where $\mathbf{e}_{lj}$ are independent $N_p(\mathbf{0}, \boldsymbol{\Sigma})$ variables. Here the parameter vector $\boldsymbol{\mu}$ is an overall mean (level) and $\boldsymbol{\tau}_l$ represents the l th treatment effect with, for instance,

$$\sum_{l=1}^g n_l \boldsymbol{\tau}_l = \mathbf{0}$$

Table 2 MANOVA table for comparing population mean vectors

Source of variation	Matrix of (SSP) ^(a)	Degrees of freedom (df)
Treatment	$\mathbf{B} = \sum_{l=1}^g n_l (\bar{\mathbf{x}}_l - \bar{\mathbf{x}})(\bar{\mathbf{x}}_l - \bar{\mathbf{x}})'$	$g - 1$
Residual (error)	$\mathbf{W} = \sum_{l=1}^g \sum_{j=1}^{n_l} (\mathbf{x}_{lj} - \bar{\mathbf{x}}_l)(\mathbf{x}_{lj} - \bar{\mathbf{x}}_l)'$	$\left(\sum_{l=1}^g n_l \right) - g$
Total (corrected for the mean)	$\mathbf{B} + \mathbf{W} = \sum_{l=1}^g \sum_{j=1}^{n_l} (\mathbf{x}_{lj} - \bar{\mathbf{x}})(\mathbf{x}_{lj} - \bar{\mathbf{x}})'$	$\left(\sum_{l=1}^g n_l \right) - 1$

^(a) SSP: sum of square and cross product

is too small. The quantity $\Lambda = |\mathbf{W}|/|\mathbf{B} + \mathbf{W}|$ is called *Wilks' lambda* after its proposer S. Wilks.

A comparison of the MANOVA table with the familiar $p = 1$ ANOVA table reveals that they are of the same structure. For the multivariate generalization, squares $(\bar{x}_i - \bar{x})^2$ are replaced by sums-of-squares and cross-products matrices $(\bar{\mathbf{x}}_i - \bar{\mathbf{x}})(\bar{\mathbf{x}}_i - \bar{\mathbf{x}})'$. The same type of replacement holds for any fixed ANOVA, so MANOVA tables can be constructed easily for any of the common designs (see, for example, references [2, Chapter 8; 4, Chapter 6 and 9; 9, Chapter 11]).

Summary Remarks

Multivariate analysis takes into account the joint variation of several responses. One noticeable difference from the univariate situation is that rejection of a null hypothesis $H_0 : \boldsymbol{\mu} = \boldsymbol{\mu}_0$ must be followed by a determination of which component(s) led to the rejection. Technically, it is at least one linear combination $a_1\mu_1 + \dots + a_p\mu_p = \mathbf{a}'\boldsymbol{\mu}$ that is different from $\mathbf{a}'\boldsymbol{\mu}_0$, but this class typically includes some individual μ_i . As we proceed to several treatments, rejection of the null hypothesis $H_0 : \boldsymbol{\mu}_1 = \boldsymbol{\mu}_2 = \dots = \boldsymbol{\mu}_g$ must be followed by a comparison of the $\boldsymbol{\mu}_i$ to determine which treatments are different and then to determine which components contribute to the difference.

Multivariate Multiple Regression

Regression analysis is the statistical methodology for predicting values of one or more *response* (dependent) variables from a collection of *predictor* (independent) variable values. It can also be used for assessing the effects of the predictor variables on the responses. In its simplest form, it applies to the fitting of a straight line to data.

The classical linear regression model states that the response Y is composed of a mean, which depends in a linear fashion on the predictor variables z_i and random error ϵ which accounts for measurement error and the effects of other variables not explicitly considered in the model (see **Mean Square Error**). The values of the predictor variables recorded from the experiment or set by the investigator are treated as *fixed*. The error (and hence the response) is viewed as a random variable, the behavior of which is characterized by a set of distributional assumptions.

Specifically, the linear regression model with a single response and n measurements on Y and the associated predictors $z_1, z_2, \dots, z_r$ can be written in matrix notation as

$$\mathbf{Y} = \mathbf{Z}\boldsymbol{\beta} + \boldsymbol{\epsilon} \quad (17)$$

with $E(\boldsymbol{\epsilon}) = \mathbf{0}$ and $\text{Cov}(\boldsymbol{\epsilon}) = \sigma^2\mathbf{I}$.

The *least-squares* (see **Least-Squares Estimation**) estimator of $\boldsymbol{\beta}$ is given by $\hat{\boldsymbol{\beta}} = (\mathbf{Z}'\mathbf{Z})^{-1}\mathbf{Z}'\mathbf{Y}$ and σ^2 is estimated by $(\mathbf{Y} - \mathbf{Z}\hat{\boldsymbol{\beta}})'(\mathbf{Y} - \mathbf{Z}\hat{\boldsymbol{\beta}})/(n - r - 1)$. The literature on multiple linear regression is vast; see the numerous books on the subject including references [10, 11].

Multivariate multiple regression is the extension of multiple regression to several response variables. Each response is assumed to follow its own regression model but with the same predictors, so that

$$\begin{aligned} Y_1 &= \beta_{01} + \beta_{11}z_1 + \dots + \beta_{r1}z_r + \epsilon_1 \\ Y_2 &= \beta_{02} + \beta_{12}z_1 + \dots + \beta_{r2}z_r + \epsilon_2 \\ &\vdots \\ Y_m &= \beta_{0m} + \beta_{1m}z_1 + \dots + \beta_{rm}z_r + \epsilon_m \end{aligned} \quad (18)$$

The error term $\boldsymbol{\epsilon} = [\epsilon_1, \epsilon_2, \dots, \epsilon_m]'$ has $E(\boldsymbol{\epsilon}) = \mathbf{0}$ and $\text{Var}(\boldsymbol{\epsilon}) = \boldsymbol{\Sigma}$. Thus, the error terms associated with different responses may be correlated.

To establish notation conforming to the classical linear regression model, let $[z_{j0}, z_{j1}, \dots, z_{jr}]$ denote the values of the predictor variables for the j th trial, $\mathbf{Y}_j = [Y_{j1}, Y_{j2}, \dots, Y_{jm}]'$ the responses, and $\boldsymbol{\epsilon}_j = [\epsilon_{j1}, \epsilon_{j2}, \dots, \epsilon_{jm}]'$ the errors. In matrix notation, the design matrix

$$\mathbf{Z}_{(n \times (r+1))} = \begin{bmatrix} z_{10} & z_{11} & \dots & z_{1r} \\ z_{20} & z_{21} & \dots & z_{2r} \\ \vdots & \vdots & \ddots & \vdots \\ z_{n0} & z_{n1} & \dots & z_{nr} \end{bmatrix} \quad (19)$$

is the same as that for the single-response regression model. The other matrix quantities have multivariate counterparts. Set

$$\mathbf{Y}_{(n \times m)} = \begin{bmatrix} Y_{11} & Y_{12} & \dots & Y_{1m} \\ Y_{21} & Y_{22} & \dots & Y_{2m} \\ \vdots & \vdots & \ddots & \vdots \\ Y_{n1} & Y_{n2} & \dots & Y_{nm} \end{bmatrix} \quad (20)$$

$$\begin{aligned} &= [\mathbf{Y}_{(1)} \quad \mathbf{Y}_{(2)} \quad \dots \quad \mathbf{Y}_{(m)}] \\ \boldsymbol{\beta}_{((r+1) \times m)} &= \begin{bmatrix} \beta_{01} & \beta_{02} & \dots & \beta_{0m} \\ \beta_{11} & \beta_{12} & \dots & \beta_{1m} \\ \vdots & \vdots & \ddots & \vdots \\ \beta_{r1} & \beta_{r2} & \dots & \beta_{rm} \end{bmatrix} \end{aligned} \quad (21)$$

$$= [\boldsymbol{\beta}_{(1)} \quad \vdots \quad \boldsymbol{\beta}_{(2)} \quad \vdots \quad \cdots \quad \vdots \quad \boldsymbol{\beta}_{(m)}]$$

and

$$\begin{aligned} \boldsymbol{\epsilon}_{(n \times m)} &= \begin{bmatrix} \epsilon_{11} & \epsilon_{12} & \cdots & \epsilon_{1m} \\ \epsilon_{21} & \epsilon_{22} & \cdots & \epsilon_{2m} \\ \vdots & \vdots & & \vdots \\ \epsilon_{n1} & \epsilon_{n2} & \cdots & \epsilon_{nm} \end{bmatrix} \\ &= [\boldsymbol{\epsilon}_{(1)} \quad \vdots \quad \boldsymbol{\epsilon}_{(2)} \quad \vdots \quad \cdots \quad \vdots \quad \boldsymbol{\epsilon}_{(m)}] \end{aligned} \quad (22)$$

Stated simply, the i th response $Y_{(i)}$ follows the linear regression model (see Table 3)

$$\mathbf{Y}_{(i)} = \mathbf{Z}\boldsymbol{\beta}_{(i)} + \boldsymbol{\epsilon}_{(i)}, \quad i = 1, 2, \dots, m \quad (23)$$

with $\text{Cov}(\boldsymbol{\epsilon}_{(i)}) = \sigma_{ii}\mathbf{I}$. However, the errors for *different* responses on the *same* trial can be correlated.

The m observations on the j th trial have covariance matrix $\boldsymbol{\Sigma} = \{\sigma_{ik}\}$, but observations from different trials are uncorrelated. Here $\boldsymbol{\beta}$ and $\boldsymbol{\Sigma}$ are matrices of unknown parameters and the design matrix $\mathbf{Z}$ has the j th row $[z_{j0}, \dots, z_{jr}]$.

Given the outcomes $\mathbf{Y}$ and the values of the predictor variables $\mathbf{Z}$, we determine the least-squares estimates $\hat{\boldsymbol{\beta}}_{(i)}$ exclusively from the observations, $\mathbf{Y}_{(i)}$, on the i th response. Since

$$\hat{\boldsymbol{\beta}}_{(i)} = (\mathbf{Z}'\mathbf{Z})^{-1}\mathbf{Z}'\mathbf{Y}_{(i)} \quad (24)$$

$$\begin{aligned} \hat{\boldsymbol{\beta}} &= [\hat{\boldsymbol{\beta}}_{(1)} \quad \vdots \quad \hat{\boldsymbol{\beta}}_{(2)} \quad \vdots \quad \cdots \quad \vdots \quad \hat{\boldsymbol{\beta}}_{(m)}] \\ &= (\mathbf{Z}'\mathbf{Z})^{-1}\mathbf{Z}'\mathbf{Y} \end{aligned} \quad (25)$$

Using the least-squares estimates $\hat{\boldsymbol{\beta}}$, we can form the matrices of

$$\text{Predicted values } \hat{\mathbf{Y}} = \mathbf{Z}\hat{\boldsymbol{\beta}} = \mathbf{Z}(\mathbf{Z}'\mathbf{Z})^{-1}\mathbf{Z}'\mathbf{Y} \quad (26)$$

$$\text{Residuals } \hat{\boldsymbol{\epsilon}} = \mathbf{Y} - \hat{\mathbf{Y}} = \mathbf{I} - \mathbf{Z}(\mathbf{Z}'\mathbf{Z})^{-1}\mathbf{Z}'\mathbf{Y} \quad (27)$$

Table 3 Multivariate linear regression model

$\mathbf{Y}_{(n \times m)}$	$=$	$\mathbf{Z}_{(n \times (r+1))}$	$\boldsymbol{\beta}_{((r+1) \times m)}$	$+$	$\boldsymbol{\epsilon}_{(n \times m)}$
with					
$E(\boldsymbol{\epsilon}_{(i)}) = \mathbf{0};$					
$\text{Cov}(\boldsymbol{\epsilon}_{(i)}, \boldsymbol{\epsilon}_{(k)}) = \sigma_{ik}\mathbf{I}, \quad i, k = 1, 2, \dots, m.$					
The m observations on the j th trial have covariance matrix $\boldsymbol{\Sigma} = \{\sigma_{ik}\}$, but observations from different trials are uncorrelated. Here $\boldsymbol{\beta}$ and $\boldsymbol{\Sigma}$ are matrices of unknown parameters and the design matrix $\mathbf{Z}$ has the j th row $[z_{j0}, \dots, z_{jr}]$.					

Example 2 Companies considering the purchase of a computer must first assess their future needs in order to determine the proper equipment. Data from several similar company sites can be used to develop a forecast equation of computer hardware requirements for, say, inventory management. The independent variables might include z_1 = customer orders and z_2 = add–delete items. The multivariate responses might include Y_1 = central processing unit (CPU) time and Y_2 = disc input/output capacity [10, Chapter 7].

Summary Remarks

Anderson [2, Section 8.7] derives test statistics and discusses the distribution theory for multivariate regression.

By allowing $\mathbf{Z}$ to have less than full rank, all fixed-effects MANOVA can be incorporated into the multivariate multiple regression framework. This unifying concept is also valuable in connecting ANOVA with the classical multiple linear regression model.

A regression model in which the $((r+1) \times m)$ coefficient matrix $\boldsymbol{\beta}$ is less than full rank is known as a *reduced-rank regression model*. This can arise in a variety of contexts, and can also be linked to the multivariate techniques of principal components and canonical correlation analysis. See reference [12] for a comprehensive discussion of multivariate reduced-rank regression.

Analysis of Covariance Structure

Principal Components

A principal component analysis is concerned with explaining the variance–covariance structure through a few *linear* combinations of the original variables. Its general objectives are (a) data reduction and (b) interpretation.

Although p components are required to reproduce the total system variability, often, much of this variability can be accounted for by a small number k of the principal components. If so, there is (almost) as much information in the k components as there is in the original p variables. The k principal components can then replace the initial p variables, and the original data set is reduced to one consisting of n measurements on k principal components.

Analyses of principal components are more of a means to an end than an end in themselves because they frequently serve as intermediate steps in much larger investigations. For example, principal components may be inputs to a multiple linear regression (as incorporated in the *partial least-squares* approach).

Algebraically, principal components are particular linear combinations of the p random variables $X_1, X_2, \dots, X_p$. Geometrically, these linear combinations represent the selection of a new coordinate system obtained by rotating the original system with $X_1, X_2, \dots, X_p$ as the coordinate axes. The new axes represent the directions with maximum variability and provide a simpler and more parsimonious description of the covariance structure.

Principal components depend solely on the covariance matrix Σ (or the correlation matrix ρ) of $X_1, X_2, \dots, X_p$. Their development does not require a multivariate normal assumption.

The first principal component is the linear combination with maximum variance. That is, it maximizes $\text{Var}(Y_1) = \mathbf{a}_1' \Sigma \mathbf{a}_1$, where the coefficient vector $\mathbf{a}_1$ is restricted to be of unit length. Therefore, we define first principal component = linear combination $\mathbf{a}_1' \mathbf{X}$ that maximizes

$$\text{Var}(\mathbf{a}_1' \mathbf{X}) \quad \text{subject to} \quad \mathbf{a}_1' \mathbf{a}_1 = 1 \quad (28)$$

second principal component = linear combination $\mathbf{a}_2' \mathbf{X}$ that maximizes

$$\text{Var}(\mathbf{a}_2' \mathbf{X}) \quad \text{subject to} \quad \mathbf{a}_2' \mathbf{a}_2 = 1 \quad (29)$$

and

$$\text{Cov}(\mathbf{a}_1' \mathbf{X}, \mathbf{a}_2' \mathbf{X}) = 0 \quad (30)$$

At the i th step, i th principal component = linear combination $\mathbf{a}_i' \mathbf{X}$ that maximizes

$$\text{Var}(\mathbf{a}_i' \mathbf{X}) \quad \text{subject to} \quad \mathbf{a}_i' \mathbf{a}_i = 1 \quad (31)$$

and

$$\text{Cov}(\mathbf{a}_i' \mathbf{X}, \mathbf{a}_k' \mathbf{X}) = 0 \quad \text{for} \quad k < i \quad (32)$$

Let Σ be the covariance matrix associated with the random vector $\mathbf{X}' = [X_1, X_2, \dots, X_p]$. Let Σ have the eigenvalue–eigenvector pairs $(\lambda_1, \mathbf{e}_1), (\lambda_2, \mathbf{e}_2), \dots, (\lambda_p, \mathbf{e}_p)$, where $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_p \geq 0$. The i th principal component is given by

$$Y_i = \mathbf{e}_i' \mathbf{X} = e_{i1}X_1 + e_{i2}X_2 + \dots + e_{ip}X_p, \quad i = 1, 2, \dots, p \quad (33)$$

With these choices,

$$\text{Var}(Y_i) = \mathbf{e}_i' \Sigma \mathbf{e}_i = \lambda_i, \quad i = 1, 2, \dots, p \quad (34)$$

$$\text{Cov}(Y_i, Y_k) = \mathbf{e}_i' \Sigma \mathbf{e}_k = 0, \quad i \neq k \quad (35)$$

Therefore, the principal components are uncorrelated and have variances equal to the eigenvalues of Σ . If some λ_i are equal, the choices of the corresponding coefficient vectors $\mathbf{e}_i$, and hence of $\mathbf{Y}_i$ are not unique.

How do we summarize the sample variation in n measurements on p variables with a few judiciously chosen linear combinations?

Assume the data $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ represent n independent drawings from some p -dimensional population with mean vector μ and covariance matrix Σ . These data yield the sample mean vector $\bar{\mathbf{x}}$, the sample covariance matrix $\mathbf{S}$, and the sample correlation matrix $\mathbf{R}$. These quantities are substituted for the corresponding population quantities above to get *sample principal components*.

Example 3 The weekly rates of return for five stocks (Allied Chemical, DuPont, Union Carbide, Exxon, and Texaco) listed on the New York Stock Exchange were determined for the period January 1975–December 1976. The weekly rates of return are defined as (current Friday closing price – previous Friday closing price)/(previous Friday closing price) adjusted for stock splits and dividends. The observations in 100 successive weeks appear to be distributed independently, but the rates of return *across* stocks are correlated, since, as one expects, stocks tend to move together in response to general economic conditions.

Let $x_1, x_2, \dots, x_5$ denote observed weekly rates of return for Allied Chemical, DuPont, Union Carbide, Exxon, and Texaco, respectively. Then,

$$\bar{\mathbf{x}}' = [0.0054, 0.0048, 0.0057, 0.0063, 0.0037] \quad (36)$$

$$\mathbf{R} = \begin{bmatrix} 1.000 & 0.577 & 0.509 & 0.387 & 0.462 \\ 0.577 & 1.000 & 0.599 & 0.389 & 0.322 \\ 0.509 & 0.599 & 1.000 & 0.436 & 0.426 \\ 0.387 & 0.389 & 0.436 & 1.000 & 0.523 \\ 0.462 & 0.322 & 0.426 & 0.523 & 1.000 \end{bmatrix} \quad (37)$$

8 Multivariate Analysis

We note that $\mathbf{R}$ is the covariance matrix of the standardized observations

$$\begin{aligned} z_1 &= \frac{x_1 - \bar{x}_1}{\sqrt{s_{11}}}, \quad z_2 = \frac{x_2 - \bar{x}_2}{\sqrt{s_{22}}}, \quad \dots, \\ z_5 &= \frac{x_5 - \bar{x}_5}{\sqrt{s_{55}}} \end{aligned} \quad (38)$$

The eigenvalues and corresponding normalized eigenvectors of $\mathbf{R}$ were determined by a computer and are

$$\hat{\lambda}_1 = 2.857, \hat{\mathbf{e}}'_1 = [0.464, 0.457, 0.470, 0.421, 0.421] \quad (39)$$

$$\hat{\lambda}_2 = 0.809, \hat{\mathbf{e}}'_2 = [0.240, 0.509, 0.260, -0.526, -0.582] \quad (40)$$

$$\hat{\lambda}_3 = 0.540, \hat{\mathbf{e}}'_3 = [-0.612, 0.178, 0.335, 0.541, -0.435] \quad (41)$$

$$\hat{\lambda}_4 = 0.452, \hat{\mathbf{e}}'_4 = [0.387, 0.206, -0.662, 0.472, -0.382] \quad (42)$$

$$\hat{\lambda}_5 = 0.343, \hat{\mathbf{e}}'_5 = [-0.451, 0.676, -0.400, -0.176, 0.385] \quad (43)$$

Using the standardized variables, we obtain the first two sample principal components

$$\hat{y}_1 = \hat{\mathbf{e}}'_1 \mathbf{z} = 0.464z_1 + 0.457z_2 + 0.470z_3 + 0.421z_4 + 0.421z_5 \quad (44)$$

$$\hat{y}_2 = \hat{\mathbf{e}}'_2 \mathbf{z} = 0.240z_1 + 0.509z_2 + 0.260z_3 - 0.526z_4 - 0.582z_5 \quad (45)$$

These components, which account for

$$\begin{aligned} \left(\frac{\hat{\lambda}_1 + \hat{\lambda}_2}{p} \right) 100\% &= \left(\frac{2.857 + 0.809}{5} \right) 100\% \\ &= 73\% \end{aligned} \quad (46)$$

of the total (standardized) sample variance, have interesting interpretations. The first component is a (roughly) equally weighted sum or “index”, of the five stocks. This component might be called a *general stock-market component* or simply a *market component*. (In fact, at the time the data were

collected, these five stocks were included in the Dow–Jones Industrial Average.)

The second component represents a contrast between the chemical stocks (Allied Chemical, DuPont, and Union Carbide) and the oil stocks (Exxon and Texaco). It might be called an *industry component*. Thus, most of the variation in these stock returns is due to market activity and uncorrelated industry activity.

The remaining components are not easy to interpret and, collectively, represent variation that is probably specific to each stock. In any event, they do not explain much of the total sample variance.

Factor Analysis

The essential purpose of factor analysis is to describe, if possible, the covariance relationships among many variables in terms of a few underlying, but unobservable, random quantities called *factors*. Basically, the factor model is motivated by the following argument. Suppose variables can be grouped by their correlations. That is, all variables within a particular group are highly correlated among themselves but have relatively small correlations with variables in a different group. It is conceivable that each group of variables represents a single underlying construct or factor that is responsible for the observed correlations. For example, correlations from the group of test scores in Classics, French, English, Mathematics, and Music collected by Spearman suggested an underlying “intelligence” factor. A second group of variables, perhaps representing physical-fitness scores, might correspond to another factor. It is this type of structure that factor analysis seeks to confirm.

Factor analysis can be considered an extension of principal component analysis. Both can be viewed as attempts to approximate the covariance matrix Σ . The approximation based on the factor analysis model is more elaborate; the primary question is whether the data are consistent with a prescribed structure.

The observable random vector $\mathbf{X}$ with p components has mean μ and covariance matrix Σ . The factor model postulates that $\mathbf{X}$ is linearly dependent on a few unobservable random variables $F_1, F_2, \dots, F_m$, called *common factors* and p additional sources of variation $\epsilon_1, \epsilon_2, \dots, \epsilon_p$, called *errors* or, sometimes,

specific factors. In particular, the factor analysis model is

$$\begin{aligned} X_1 - \mu_1 &= l_{11}F_1 + l_{12}F_2 + \cdots + l_{1m}F_m + \epsilon_1 \\ X_2 - \mu_2 &= l_{21}F_1 + l_{22}F_2 + \cdots + l_{2m}F_m + \epsilon_2 \\ &\vdots \\ X_p - \mu_p &= l_{p1}F_1 + l_{p2}F_2 + \cdots + l_{pm}F_m + \epsilon_p \end{aligned} \quad (47)$$

or, in matrix notation,

$$\underset{(p \times 1)}{\mathbf{X}} - \underset{(p \times 1)}{\boldsymbol{\mu}} = \underset{(p \times m)}{\mathbf{L}} \underset{(m \times 1)}{\mathbf{F}} + \underset{(p \times 1)}{\boldsymbol{\epsilon}} \quad (48)$$

The coefficient l_{ij} is the *loading* of the i th variable on the j th factor, so the matrix $\mathbf{L}$ is the *matrix of factor loadings*. Note that the i th specific factor ϵ_i is associated only with the i th response X_i . The p deviations $X_1 - \mu_1, X_2 - \mu_2, \dots, X_p - \mu_p$ are expressed in terms of $p + m$ variables $F_1, F_2, \dots, F_m, \epsilon_1, \epsilon_2, \dots, \epsilon_p$, which are unobservable.

With so many unobservable quantities, a direct verification of the factor model from observations on $X_1, X_2, \dots, X_p$ is hopeless. However, with some additional assumptions about the random vectors $\mathbf{F}$ and $\boldsymbol{\epsilon}$, the preceding model implies certain covariance relationships that can be checked. It follows immediately from the factor model that

1. $\text{Cov}(\mathbf{X}) = \mathbf{L}\mathbf{L}' + \boldsymbol{\Psi}$, or

$$\text{Var}(X_i) = l_{i1}^2 + \cdots + l_{im}^2 + \psi_i \quad \text{and} \quad (49)$$

$$\text{Cov}(X_i, X_k) = l_{i1}l_{k1} + \cdots + l_{im}l_{km} \quad (50)$$

2. $\text{Cov}(\mathbf{X}, \mathbf{F}) = \mathbf{L}$, or $\text{Cov}(X_i, F_j) = l_{ij}$.

That portion of the variance of the i th variable contributed by the m common factors is called the i th *communality*. The portion of $\text{Var}(X_i) = \sigma_{ii}$ due to the specific factor is often called the *uniqueness* or *specific variance*. Denoting the i th communality by h_i^2 ,

$$\begin{aligned} \underbrace{\sigma_{ii}}_{\text{Var}(X_i)} &= \underbrace{l_{i1}^2 + l_{i2}^2 + \cdots + l_{im}^2}_{\text{communality}} + \underbrace{\psi_i}_{\text{specific variance}} \end{aligned} \quad (51)$$

or $h_i^2 = l_{i1}^2 + l_{i2}^2 + \cdots + l_{im}^2$, and $\sigma_{ii} = h_i^2 + \psi_i$, $i = 1, 2, \dots, p$. The i th communality is the sum of squares of the loadings of the i th variable on the m common factors.

Given observations $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ on p generally correlated variables, factor analysis seeks to answer the question: Does the factor model (see Table 4), with a small number of factors, adequately represent the data? In essence, we tackle this statistical model-building problem by trying to verify covariance relationships 1 and 2.

The sample covariance matrix $\mathbf{S}$ is an estimator of the unknown population covariance matrix $\boldsymbol{\Sigma}$. If the off-diagonal elements of $\mathbf{S}$ are small or those of the sample correlation matrix $\mathbf{R}$ essentially zero, the variables are not related and a factor analysis will not prove useful. In these circumstances, the *specific* factors play the dominant role, whereas the major aim of the factor analysis is to determine a few important *common* factors.

If $\boldsymbol{\Sigma}$ appears to deviate significantly from a diagonal matrix, then a factor model can be entertained. The initial problem is one of estimating the factor loadings l_{ij} and specific variances ψ_i . (Methods of **estimation** are discussed, for example, in references [5, Section 5.4; 7, Section 7.3; 13, Chapter 9].)

Table 4 Orthogonal factor model with m common factors

$\underset{(p \times 1)}{\mathbf{X}} = \underset{(p \times 1)}{\boldsymbol{\mu}} + \underset{(p \times m)}{\mathbf{L}} \underset{(m \times 1)}{\mathbf{F}} + \underset{(p \times 1)}{\boldsymbol{\epsilon}}$
$\mu_i = \text{mean of variable } i$
$\epsilon_i = i\text{th specific factor}$
$F_j = j\text{th common factor}$
$l_{ij} = \text{loading of the } i\text{th variable on the } j\text{th factor}$
The unobservable random vectors $\mathbf{F}$ and $\boldsymbol{\epsilon}$ satisfy
$\mathbf{F}$ and $\boldsymbol{\epsilon}$ are independent.
$E(\mathbf{F}) = \mathbf{0}, \text{Cov}(\mathbf{F}) = \mathbf{I}.$
$E(\boldsymbol{\epsilon}) = \mathbf{0}, \text{Cov}(\boldsymbol{\epsilon}) = \boldsymbol{\Psi},$ where $\boldsymbol{\Psi}$ is a diagonal matrix.

All factor loadings obtained from the initial loading by an orthogonal transformation have the same ability to reproduce the covariance (or correlation) matrix. From matrix algebra, we know that an orthogonal transformation corresponds to a rigid rotation (or reflection) of the coordinate axes. For this reason, an orthogonal transformation of the factor loadings and the implied orthogonal transformation of the factors is called *factor rotation*.

If $\hat{\mathbf{L}}$ is the $p \times m$ matrix of estimated factor loadings obtained by any method, then

$$\hat{\mathbf{L}}^* = \hat{\mathbf{L}}\mathbf{T}, \quad \text{where} \quad \mathbf{T}\mathbf{T}' = \mathbf{T}'\mathbf{T} = \mathbf{I} \quad (52)$$

is a $p \times m$ matrix of “rotated” loadings. Moreover, the estimated covariance (or correlation) matrix remains unchanged, since

$$\hat{\mathbf{L}}\hat{\mathbf{L}}' + \hat{\Psi} = \hat{\mathbf{L}}\mathbf{T}\mathbf{T}'\hat{\mathbf{L}} + \hat{\Psi} = \hat{\mathbf{L}}^*\hat{\mathbf{L}}^{*'} + \hat{\Psi} \quad (53)$$

The residual matrix $\mathbf{S} - \hat{\mathbf{L}}\hat{\mathbf{L}}' - \hat{\Psi} = \mathbf{S} - \hat{\mathbf{L}}^*\hat{\mathbf{L}}^{*'} - \hat{\Psi}$ also remains unchanged after rotation. Moreover, the specific variances ψ_i and hence the communalities $\hat{h}_i^2$ are unaltered. Therefore, from a mathematical viewpoint, it is immaterial whether $\hat{\mathbf{L}}$ or $\hat{\mathbf{L}}^*$ is obtained.

Since the original loadings may not be readily interpretable, the usual practice is to rotate them until a “simple structure” is achieved. The rationale is very much akin to sharpening the focus of a microscope in order to see the detail more clearly.

Ideally, we should like to see a pattern of loadings such that each variable loads highly on a single factor and has small-to-moderate loadings on the remaining factors. It is not always possible to get this simple structure.

Kaiser [14] has suggested an analytical measure of simple structure known as the *varimax* (or *normal varimax*) criterion. Define $\tilde{l}_{ij}^* = \hat{l}_{ij}^*/\hat{h}_i$ to be the final rotated coefficients scaled by the square root of the communalities. The (normal) varimax procedure selects the orthogonal transformation $\mathbf{T}$ that makes

$$V = \frac{1}{p} \sum_{j=1}^m \left[\sum_{i=1}^p \tilde{l}_{ij}^{*4} - \left(\sum_{i=1}^p \tilde{l}_{ij}^{*2} \right)^2 / p \right] \quad (54)$$

as large as possible.

Effectively maximizing V corresponds to “spreading out” the squares of the loadings on each factor as much as possible. Therefore, we hope to find groups

of large and negligible coefficients in any *column* of the rotated loadings matrix $\mathbf{L}^*$.

In factor analysis, interest is usually centered on the parameters in the factor model. However, the estimated values of the common factors, called *factor scores*, may also be required. Often, these quantities are used for diagnostic purposes as well as inputs to a subsequent analysis. (See references [4, Section 9.5; 7, Section 7.8; 8, Section 13.6] for further discussion of factor scores.)

Example 4 Beginning with correlations between the scores of the Olympic decathlon events, a factor analysis can be employed to see if the 10 events can be explained in terms of two, three, or four underlying “physical” factors. One interesting study of this kind [15] found that the four factors “explosive arm strength”, “explosive leg strength”, “running speed”, and “running endurance” represented several years of decathlon data quite well.

Canonical Correlations and Variables

Canonical correlation analysis seeks to identify and quantify the associations between two sets of variables. Harold Hotelling [16], who initially developed the technique, provided the example of relating arithmetic speed and arithmetic power to reading speed and reading power. Other examples include relating governmental policy variables to economic goal variables and relating college “performance” variables with precollege “achievement” variables.

Canonical correlation analysis focuses on the correlation between a linear combination of the variables in one set and a linear combination of the variables in another set. When the association between the two sets is expected to be unidirectional – from one set to the other – we label one set the *independent* or *predictor* variables and the other set the *dependent* or *criterion* variables.

The idea of canonical correlation analysis is to first determine the pair of linear combinations having the largest correlation. Next, one determines the pair of predictor set/criterion set linear combinations having the largest correlation among all pairs uncorrelated with the initially selected pair. The process continues by selecting, at each stage, the pair of predictor set/criterion set linear combinations having largest correlation among all pairs that are uncorrelated with the preceding choices. The pairs of linear

combinations are the *canonical variables* and their correlations are *canonical correlations*. The following discussion gives the necessary details for obtaining the canonical variables and their correlations. In practice, sample covariance matrices are substituted for the corresponding population quantities, yielding sample canonical variables and sample canonical correlations.

Suppose $p \leq q$, and let the random vectors

$$\begin{matrix} \mathbf{X}_1 & \text{and} & \mathbf{X}_2 \\ (p \times 1) & & (q \times 1) \end{matrix}$$

have

$$\text{Cov}(\mathbf{X}_1) = \boldsymbol{\Sigma}_{11}, \text{Cov}(\mathbf{X}_2) = \boldsymbol{\Sigma}_{22} \quad (55)$$

$(p \times p) \quad (q \times q)$

and

$$\text{Cov}(\mathbf{X}_1, \mathbf{X}_2) = \boldsymbol{\Sigma}_{12} \quad (56)$$

$(p \times q)$

For coefficient vectors,

$$\begin{matrix} \mathbf{a} & \text{and} & \mathbf{b} \\ (p \times 1) & & (q \times 1) \end{matrix}$$

form the linear combinations $U = \mathbf{a}'\mathbf{X}_1$ and $V = \mathbf{b}'\mathbf{X}_2$. Then

$$\text{Max}_{\mathbf{a}, \mathbf{b}} \text{Corr}(U, V) = \rho_1^* \quad (57)$$

attained by the linear combination (first canonical variate pair)

$$U_1 = \mathbf{a}'_1 \mathbf{X}_1 \quad \text{and} \quad V_1 = \mathbf{b}'_1 \mathbf{X}_2 \quad (58)$$

The k th pair of canonical variates, $k = 2, 3, \dots, p$,

$$U_k = \mathbf{a}'_k \mathbf{X}_1 \quad \text{and} \quad V_k = \mathbf{b}'_k \mathbf{X}_2 \quad (59)$$

maximize

$$\text{Corr}(U_k, V_k) = \rho_k^* \quad (60)$$

among those linear combinations uncorrelated with the preceding $1, 2, \dots, k-1$ canonical variables.

The canonical variates have the following properties:

$$\text{Var}(U_k) = \text{Var}(V_k) = 1 \quad (61)$$

$$\text{Cov}(U_k, U_l) = \text{Corr}(U_k, U_l) = 0, \quad k \neq l \quad (62)$$

$$\text{Cov}(V_k, V_l) = \text{Corr}(V_k, V_l) = 0, \quad k \neq l \quad (63)$$

$$\text{Cov}(U_k, V_l) = \text{Corr}(U_k, V_l) = 0, \quad k \neq l \quad (64)$$

for $k, l = 1, 2, \dots, p$.

In general, canonical variables are artificial; they have no physical meaning. If the original variables $\mathbf{X}_1$ and $\mathbf{X}_2$ are used, the canonical coefficients $\mathbf{a}$ and $\mathbf{b}$ have units proportional to those of the $\mathbf{X}_1$ and $\mathbf{X}_2$ sets. If the original variables are standardized to have zero means and unit variances, the canonical coefficients have no units of measurement, and they must be interpreted in terms of the standardized variables.

Classification and Grouping Techniques

Discriminant Analysis and Classification

Discriminant analysis and classification are multivariate techniques concerned with *separating* distinct sets of objects (or observations) and with *allocating* new objects (observations) to previously defined groups. Discriminant analysis is rather exploratory in nature. As a separating procedure, it is often employed on a one-time basis in order to investigate observed differences when causal relationships are not well understood. Classification procedures are less exploratory in the sense that they lead to well-defined rules, which can be used for assigning new objects. Classification ordinarily requires more problem structure than discrimination.

Thus, the immediate goals of discrimination and classification, respectively, are

1. To describe either graphically (in three or fewer dimensions) or algebraically, the differential features of objects (observations) from several known collections (populations). We try to find “discriminants” whose numerical values are such that the collections are separated as much as possible.
2. To sort objects (observations) into two or more labeled classes. The emphasis is on deriving a rule that can be used to assign a *new* object to the labeled classes optimally.

To fix ideas, we will list situations in which one may be interested in (a) separating two classes of objects, or (b) assigning a new object to one of the two classes (or both). It is convenient to label the classes π_1 and π_2 . The objects are ordinarily separated or classified on the basis of measurements, for instance, on p associated random variables $\mathbf{X}' = [X_1, X_2, \dots, X_p]$. The observed values of $\mathbf{X}$ differ to some extent from one class to the other. We can

Populations π_1 and π_2	Measured variables $\mathbf{X}$
Solvent and distressed property-liability insurance companies	Total assets, cost of stocks and bonds, market value of stocks and bonds, loss expenses, surplus, amount of premiums written
Federalist papers written by James Madison and those written by Alexander Hamilton	Frequencies of different words and length of sentences
Purchasers of a new product and laggards (those “slow” to purchase)	Education, income, family size, amount of previous brand switching
Alcoholics and nonalcoholics	Activity of monoamine oxidase enzyme, activity of adenylate cyclase enzyme

think of the totality of values from the first class as being the population of $\mathbf{x}$ values for π_1 and those from the second class as the population of $\mathbf{x}$ values for π_2 . These two populations can then be described by probability density functions $f_1(\mathbf{x})$ and $f_2(\mathbf{x})$, and, consequently, we can talk of assigning observations to populations or objects to classes interchangeably.

You may wonder at this point how it is we *know* some observations belong to a particular population but we are unsure about others. (This, of course, is what makes classification a problem!) There are several conditions that can give rise to this apparent anomaly.

- Incomplete knowledge of future performance.
- “Perfect” information requires destroying object.
- Unavailable or expensive information.

It should be clear that classification rules cannot usually provide an error-free method of assignment. This is because there may not be a clear distinction between the measured characteristics of the populations; that is, the groups may overlap. It is then possible, for example, to incorrectly classify a π_2 object as belonging to π_1 or a π_1 object as belonging to π_2 .

For a discussion of discriminant analysis and subsequent classification procedures, see references [2, Chapter 6; 4, Chapter 11; 5, Chapter 6]. Since the literature on this subject is large, we simply display

below the “best” allocation rule for two multivariate normal populations with a common covariance matrix Σ . In practice, sample quantities replace the corresponding population quantities.

Let μ_1 and μ_2 be the two population mean vectors, $c(1|2)$ the cost of incorrectly assigning a population 2 observation to population 1, $c(2|1)$ the cost of incorrectly assigning a population 1 observation to population 2, p_1 the “prior” probability of population 1, and p_2 the “prior” probability of population 2, then we allocate $\mathbf{x}$ to π_1 if

$$(\mu_1 - \mu_2)' \Sigma^{-1} \mathbf{x} - \frac{1}{2}(\mu_1 - \mu_2)' \Sigma^{-1} (\mu_1 + \mu_2) \geq \ln \left[\frac{c(1|2)}{c(2|1)} \left(\frac{p_2}{p_1} \right) \right] \quad (65)$$

Allocate $\mathbf{x}$ to π_2 otherwise.

The first term, $y = (\mu_1 - \mu_2)' \Sigma^{-1} \mathbf{x}$, above is *Fisher's linear discriminant function*. (Fisher actually developed the sample version $\hat{y} = (\bar{\mathbf{x}}_1 - \bar{\mathbf{x}}_2)' \mathbf{S}_{\text{pooled}}^{-1} \mathbf{x}$, where $\mathbf{S}_{\text{pooled}}$ is the pooled sample covariance matrix.) Assuming a common population covariance matrix, it is the linear function $\mathbf{a}'\mathbf{x}$ with $\mathbf{a} \propto (\mu_1 - \mu_2)' \Sigma^{-1}$ that maximizes the separation between the two populations, as measured by

$$\frac{(\mu_{1y} - \mu_{2y})^2}{\sigma_y^2} = \frac{(\mathbf{a}'\mu_1 - \mathbf{a}'\mu_2)^2}{\mathbf{a}'\Sigma\mathbf{a}} \quad (66)$$

One important way of judging the performance of any classification procedure is to calculate its “error rates”, or misclassification probabilities. When the forms of the parent populations are known completely, misclassification probabilities can be calculated with relative ease. Because parent populations are rarely known, one must concentrate on the error rates associated with the sample classification function.

Finally, it should be intuitive that good classification (low error rates) will depend on the separation of the populations. The farther apart the groups, the more likely it is that a *useful* classification rule can be developed.

Fisher also proposed a several-population extension of his discriminant method. The motivation behind Fisher discriminant analysis is the need to obtain a reasonable representation of the populations that involves only a *few* linear combinations of the observations, such as $\mathbf{a}_1'\mathbf{x}$, $\mathbf{a}_2'\mathbf{x}$, and $\mathbf{a}_3'\mathbf{x}$. His approach

has several advantages when one is interested in *separating* several populations for visual inspection or graphically descriptive purposes. It allows for the following:

1. Convenient representations of the g populations that reduce the dimension from a very large number of characteristics to a relatively few linear combinations. Of course, some information – needed for optimal classification – may be lost unless the population means lie completely in the lower-dimensional space selected.
2. Plotting of the means of the first two or three linear combinations (discriminants). This helps display the relationships and possible groupings of the populations.
3. Scatterplots of the sample values of the first two discriminants, which can indicate **outliers*** or other abnormalities in the data. (See references [1, Chapter 4; 4, Section 11.7; 17, Chapter 2] for examples of low-dimensional representations.)

Summary Remarks

The linear discriminant functions that we have presented can arise from a multivariate multiple linear regression model with the response variable vector containing binary categorical variables representing the different groups.

Other, typically computer intensive, procedures are available for discrimination and classification. These include logistic regression, classification trees, **neural networks**, and support vector machines. A good reference for these and related methods is [1].

Clustering and Graphical Procedures

Rudimentary, exploratory procedures are often quite helpful in understanding the complex nature of multivariate relationships. Searching the data for a structure of “natural” groupings is an important exploratory technique. Groupings can provide an informal means for assessing dimensionality, identifying outliers, and suggesting interesting hypotheses concerning relationships.

Grouping, or clustering, is distinct from the classification methods discussed earlier. Classification pertains to a *known* number of groups, and the operational objective is to assign new observations to one

of these groups. Cluster analysis is a more primitive technique in that no assumptions are made concerning the number of groups or the group structure. Grouping is done on the basis of similarities or distances (dissimilarities). The inputs required are similarity measures or data from which similarities can be computed. Good references on clustering include [1, Section 14.3; 4, Chapter 7; 18, 19].

Clustering and Lower-Dimensional Representations

Even without the precise notion of a natural grouping, we are often able to cluster objects in two- or three-dimensional scatter plots by eye. To take advantage of the mind’s ability to group similar objects, several graphical procedures have been developed for depicting high-dimensional observations in two dimensions. Boxes, glyphs, stars, Andrews plots, and Chernoff faces are among the methods that produce two-dimensional “pictures” of multivariate data.

Two-dimensional scatter plots of high-dimensional data can be constructed using principal component axes, or discriminant score axes, or canonical variate axes. In doing so, there is always some information lost. In contrast to this activity, multidimensional scaling seeks to directly “fit” the original data into a low-dimensional space such that any distortion caused by a lack of dimensionality is minimized.

A biplot is a two-dimensional graphical representation of the information in the $n \times p$ data matrix. Correspondence analysis is a graphical procedure for representing association in a two-way table of frequencies or counts.

Extensions of two-dimensional principal component representations of data such as self-organizing maps and principal curves and surfaces are discussed in reference [1, Chapter 14].

Multivariate Methods and Data Mining

Data mining (ref. to Data Mining Review article in CIM section) refers to the process of discovering associations and relationships in very large data sets (perhaps several terabytes of data). Data mining is not possible without appropriate software and fast computers. Many of the multivariate methods discussed here, along with algorithms developed in the machine learning and artificial intelligence fields, play important roles in data mining. Data mining

has helped to identify new chemical compounds for prescription drugs, detect fraudulent claims and purchases, create and maintain individual customer relationships, design better engines, improve process control, and develop effective credit scoring rules.

References

- [1] Hastie, T., Tibshirani, R. & Friedman, J. (2001). *The Elements of Statistical Learning: Data Mining, Inference and Prediction*, Springer-Verlag, New York.
- [2] Anderson, T.W. (2003). *An Introduction to Multivariate Statistical Analysis*, 3rd Edition, John Wiley & Sons, New York.
- [3] Rao, C.R. (1973). *Linear Statistical Inference and its Applications*, 2nd Edition, John Wiley & Sons, New York.
- [4] Johnson, R.A. & Wichern, D.W. (2002). *Applied Multivariate Statistical Analysis*, 5th Edition, Prentice-Hall, Upper Saddle River.
- [5] Seber, G.A.F. (2004). *Multivariate Observations*, John Wiley & Sons, New York.
- [6] Agresti, A. (2002). *Categorical Data Analysis*, 2nd Edition, John Wiley & Sons, New York.
- [7] Morrison, D.F. (2005). *Multivariate Statistical Methods*, 4th Edition, Brooks/Cole Thompson Learning, Belmont.
- [8] Rencher, A.C. (2002). *Methods of Multivariate Analysis*, 2nd Edition, John Wiley & Sons, New York.
- [9] Lattin, J.M., Carroll, J.D. & Green, P.E. (2003). *Analyzing Multivariate Data*, Brooks/Cole Thompson Learning, Pacific Grove.
- [10] Draper, N.R. & Smith, H. (1998). *Applied Regression Analysis*, 3rd Edition, John Wiley & Sons, New York.
- [11] Seber, G.A.F. & Lee, A.J. (2003). *Linear Regression Analysis*, 2nd Edition, John Wiley & Sons, New York.
- [12] Reinsel, G.C. & Velu, R.P. (1998). *Multivariate Reduced-Rank Regression: Theory and Applications*, Springer-Verlag, New York.
- [13] Mardia, K.V., Kent, J.T. & Bibby, J.M. (1979). *Multivariate Analysis*, Academic Press, New York.
- [14] Kaiser, H.F. (1958). The varimax criterion for analytic rotation in factor analysis, *Psychometrika* **23**, 187–200.
- [15] Linden, M. (1977). A factor analytic study of Olympic decathlon data, *Res. Quar.* **48**, 562–568.
- [16] Hotelling, H. (1936). Relations between two sets of variates, *Biometrika* **28**, 321–377.
- [17] Everitt, B. & Dunn, G. (2001). *Applied Multivariate Data Analysis*, 2nd Edition, Oxford University Press, New York.
- [18] Everitt, B., Landau, S. & Leese, M. (2001). *Cluster Analysis*, 4th Edition, Oxford University Press, New York.
- [19] Hartigan, J.A. (1975). *Clustering Algorithms*, John Wiley & Sons, New York.

Further Reading

- Basilevsky, A.T. (1994). *Statistical Factor Analysis and Related Methods: Theory and Applications*, John Wiley & Sons, New York.
- Berry, M.J.A. & Linoff, G.S. (2004). *Data Mining Techniques*, 2nd Edition, John Wiley & Sons, New York.
- Bishop, C.M. (1995). *Neural Networks for Pattern Recognition*. Oxford University Press, New York.
- Bollen, K.A. (1989). *Structural Equation Models with Latent Variables*, John Wiley & Sons, New York.
- Breiman, L., Friedman, J., Olshen, R. & Stone, C. (1984). *Classification and Regression Trees*, Wadsworth, Belmont.
- Gower, J.C. & Hand, D.J. (1995). *Biplots*, Chapman & Hall, London.
- Greenacre, M.J. (1984). *Theory and Applications of Correspondence Analysis*, Academic Press, London.
- Hand, D., Mannila, H. & Smyth, P. (2001). *Principles of Data Mining*, MIT Press, Cambridge.
- Jolliffe, I.T. (2002). *Principal Component Analysis*, 2nd Edition, Springer-Verlag, New York.
- Mason, R.L. & Young, J.C. (2002). *Multivariate Statistical Process Control with Industrial Applications*, ASA-SIAM, Philadelphia.
- Reinsel, G.C. (1997). *Elements of Multivariate Time Series Analysis*, 2nd Edition, Springer-Verlag, New York.

RICHARD A. JOHNSON AND DEAN W.
WICHERN

Article originally published in Encyclopedia of Statistical Sciences, 2nd Edition (2005, John Wiley & Sons, Inc.). Minor revisions for this publication by Jeroen de Mast.

Probability Theory

Introduction

Probability theory provides a mathematical framework for the description of phenomena whose outcomes are not deterministic but rather are subject to chance. This framework, historically motivated by a frequency interpretation of probability, is provided by a formal set of three innocent-looking axioms and the consequences thereof. These are far-reaching, and a rich and diverse theory has been developed. It provides a flexible and effective model of physical reality and furnishes, in particular, the mathematical theory needed for the study of the discipline of statistics.

The Formal Framework

Elementary (nonmeasure theoretic) treatments of probability can readily deal with situations in which the set of distinguishable outcomes of an experiment subject to chance (called *elementary events*) is finite. This encompasses most simple gambling situations, for example.

Let the set of all elementary events be Ω (called the *sample space*) and let $\mathcal{F}$ be the Boolean field of subsets of Ω (meaning just that (a) $A \in \mathcal{F}$ implies that $\bar{A} \in \mathcal{F}$, the bar denoting complementation with respect to Ω and (b) $A_1, A_2 \in \mathcal{F}$ implies that $A_1 \cup A_2 \in \mathcal{F}$). Then the probability P is a set function on $\mathcal{F}$ satisfying the axioms:

Axiom 1 $P(A) \geq 0, A \in \mathcal{F}$.

Axiom 2 If $A_1, A_2 \in \mathcal{F}$ and $A_1 \cap A_2 = \emptyset$ (the empty set), then

$$P(A_1 \cup A_2) = P(A_1) + P(A_2) \quad (1)$$

Axiom 3 $P(\Omega) = 1$.

These axioms enable probabilities to be assigned to events in simple games of chance. For example, in the toss of an ordinary six-sided die, Ω consists of six equiprobable elementary events and for $A \in \mathcal{F}$, $P(A) = n(A)/6$, where $n(A)$ is the number of elementary events in A .

The formulation above can be extended straightforwardly to the case where the number of elementary

events in the sample space is at most countable, by extending Axiom 2 to deal with countable unions of disjoint sets. However, a full measure-theoretic formulation is necessary to deal with general sample spaces Ω and to provide a framework which is rich enough for a comprehensive study of limit results on combinations of events (sets).

The measure-theoretic formulation has a very similar appearance to the elementary version. We start from a general sample space Ω and let $\mathcal{F}$ be the Borel σ -field of subsets of Ω (meaning that (a) $A \in \mathcal{F}$ implies that $\bar{A} \in \mathcal{F}$ and (b) $A_i \in \mathcal{F}, 1 \leq i < \infty$, implies that $\bigcup_{i=1}^{\infty} A_i \in \mathcal{F}$). Then as axioms for the probability measure P we have Axioms 1 and 3 as before, while Axiom 2 is replaced by:

Axiom 2' If $A_i \in \mathcal{F}, 1 \leq i < \infty$, and $A_i \cap A_j = \emptyset, i \neq j$, then

$$P\left(\bigcup_{i=1}^{\infty} A_i\right) = \sum_{i=1}^{\infty} P(A_i) \quad (2)$$

The triplet $(\Omega, \mathcal{F}, P)$ is called a *probability space*. This framework is due to Kolmogorov [1] in 1933 and is very widely accepted, although some authors reject the countable additivity Axiom 2' and rely on the finite additivity version (Axiom 2) (e.g., Dubins and Savage [2]). An extension of the Kolmogorov theory in which unbounded measures are allowed and conditional probabilities are taken as the fundamental concept has been provided by Rényi [3, Chapter 2]. For $A, B \in \mathcal{F}$ with $P(B) > 0$, the conditional probability of A given B is defined as $P(A|B) = P(A \cap B)/P(B)$.

Suppose that the sample space Ω has points ω . A real-valued function $X(\omega)$ defined on the probability space $(\Omega, \mathcal{F}, P)$ is called a *random variable* (or, in the language of measure theory, $X(\omega)$ is a *measurable function*) if the set $\{\omega : X(\omega) \leq x\}$ belongs to $\mathcal{F}$ for every $x \in \mathbb{R}$ (the real line). Sets of the form $\{\omega : X(\omega) \leq x\}$ are usually written as $\{X \leq x\}$ for convenience and we shall follow this practice. Random variables and their relationships are the principal objects of study in much of probability and statistics.

Important Concepts

Each random variable X has an associated *distribution function* defined for each real x by

$$F(x) = P(X \leq x) \quad (3)$$

2 Probability Theory

The function F is nondecreasing, right continuous, and has at most a countable set of points of discontinuity. It also uniquely specifies the probability measure that X induces on the real line.

Random variables are categorized as having a discrete distribution if the corresponding distribution function is a step function and as having a continuous distribution if it is continuous. Most continuous distributions of practical interest are absolutely continuous, meaning that the distribution function $F(x)$ has a representation of the form

$$F(x) = \int_{-\infty}^x f(u) du \quad (4)$$

for some $f \geq 0$; f is called the **probability density function**.

A relatively small number of families of distributions are widely used in applications of probability theory (see **Probability Density Functions**). Among the most important are the binomial, which is discrete and for which

$$P(X = r) = \binom{n}{r} p^r (1-p)^{n-r}, \quad 0 \leq r \leq n, \quad 0 < p < 1 \quad (5)$$

and the (unit) normal, which is absolutely continuous and for which

$$P(X \leq x) = (2\pi)^{-1/2} \int_{-\infty}^x e^{-(1/2)u^2} du \quad -\infty < x < \infty \quad (6)$$

The binomial random variable represents the numbers of successful outcomes in n trials, repeated under the same conditions, where the probability of success in each individual trial is p . Normal distributions, or close approximations thereto, occur widely in practice, for example, in such measurements as heights or weights or examination scores, where a large number of individuals are involved.

Two other particularly important distributions are the Poisson, for which

$$P(X = r) = e^{-\lambda} \lambda^r / r! \quad r = 0, 1, 2, \dots, \quad \lambda > 0 \quad (7)$$

and the exponential, for which

$$P(X \leq x) = \lambda \int_0^x e^{-\lambda u} du, \quad x > 0, \quad \lambda > 0 \\ = 0 \text{ otherwise} \quad (8)$$

These commonly occur in situations of rare events and as waiting times between phenomena, respectively.

A sequence of random variables $\{X_1, X_2, \dots\}$ defined on the same probability space is called a *stochastic process*. For a finite collection of such random variables $\{X_1, X_2, \dots, X_m\}$ we can define the joint distribution function

$$F(x_1, x_2, \dots, x_m) \\ = P(X_1 \leq x_1, X_2 \leq x_2, \dots, X_m \leq x_m) \quad (9)$$

for $x_i \in \mathbb{R}$, $1 \leq i \leq m$, and the random variables are said to be *independent* if

$$P(X_1 \leq x_1, \dots, X_m \leq x_m) \\ = P(X_1 \leq x_1) \cdots P(X_m \leq x_m) \quad (10)$$

for all $x_1, \dots, x_m$ (see **Dependence**). Independence is often an essential basic property and much of probability and statistics is concerned with random sampling, namely independent observations (random variables) with identical distributions.

One step beyond independence is the Markov chain $\{X_1, X_2, \dots\}$, in which case

$$P(X_m \leq x_m | X_1 = i_1, \dots, X_n = i_n) \\ = P(X_m \leq x_m | X_n = i_n) \quad (11)$$

for any $n < m$ and all $i_1, \dots, i_n$ and x_m . That is, dependence on a past history involves only the most recent of the observations. Such random sequences with a natural timescale belong to the domain of stochastic processes; see **Markov Processes**.

Considerable summary information about random variables and their distributions is contained in quantities called *moments* (see **Moments**). For a random variable X with distribution function F , the r th *moment* ($r > 0$) is defined by

$$EX^r = \int_{-\infty}^{\infty} x^r dF(x) \quad (12)$$

provided that $\int_{-\infty}^{\infty} |x|^r dF(x) < \infty$. $E|X|^r$ is called the r th *absolute moment*. The operator E is called *expectation* (see **Expectation**) and the first moment ($r = 1$) is termed the *mean*. Moments about the mean $E(X - EX)^r$, $r = 2, 3, \dots$ are widely used and $E(X - EX)^2$ is termed the *variance* (see **Variance**) of X , while $(E(X - EX)^2)^{1/2}$ is its *standard deviation*. The mean and variance of X are, respectively, measures of the location and spread of its distribution.

Convergence

Much of probability theory is concerned with the limit behavior of sequences of random variables or their distributions. The great diversity of possible probabilistic behavior when the **sample size** is small is often replaced by unique and well-defined probabilistic behavior for large sample sizes. In this lies the great virtue of the limit theory: statistical regularity emerging out of statistical chaos as the sample size increases.

A sequence of random variables $\{X_n\}$ is said to converge *in distribution* (or *in law*) to that of a random variable X (written $X_n \xrightarrow{d} X$) if the sequence of distribution functions $\{F_n(x) = P(X_n \leq x)\}$ converges to the distribution function $F(x) = P(X \leq x)$ at all points x of continuity of F .

Convergence in distribution does not involve the random variables explicitly, only their distributions. Other commonly used modes of convergence are the following:

1. $\{X_n\}$ converges *in probability* to X (written $X_n \xrightarrow{P} X$) if, for every $\epsilon > 0$,

$$\lim_{n \rightarrow \infty} P(|X_n - X| > \epsilon) = 0 \quad (13)$$

2. For $r > 0$, $\{X_n\}$ converges *in the mean of order r* to X (written $X_n \xrightarrow{r} X$) if

$$\lim_{n \rightarrow \infty} E|X_n - X|^r = 0 \quad (14)$$

3. $\{X_n\}$ converges *almost surely* to X (written $X_n \xrightarrow{\text{a.s.}} X$) if for every $\epsilon > 0$,

$$\lim_{n \rightarrow \infty} P\left(\sup_{k \geq n} |X_k - X| > \epsilon\right) = 0 \quad (15)$$

It is not difficult to show that 2 implies 1, which in turn implies convergence in distribution while 3 implies 1. None of the converse implications hold in general.

Convergence in distribution is frequently established by the method of characteristic functions. The *characteristic function* of a random variable X with distribution function F is defined by

$$f(t) = Ee^{itX} = \int_{-\infty}^{\infty} e^{itx} dF(x) \quad (16)$$

where $i = \sqrt{-1}$. If $f_n(t)$ and $f(t)$ are, respectively, the characteristic functions of X_n and X , then $X_n \xrightarrow{d} X$ is equivalent to $f_n(t) \rightarrow f(t)$. Characteristic functions are particularly convenient for dealing with sums of independent random variables as a consequence of the multiplicative property

$$Ee^{it(X+Y)} = Ee^{itX} Ee^{itY} \quad (17)$$

for independent X and Y . They have been extensively used in most treatments of convergence in distribution for sums of independent random variables (the classical central limit problem). The basic central limit theorem (of which there are many generalizations) deals with independent and identically distributed random variables X_i with mean μ and variance σ^2 . It gives, upon writing $S_n = \sum_{i=1}^n X_i$, that

$$\begin{aligned} \lim_{n \rightarrow \infty} P(\sigma^{-1}n^{-1/2}(S_n - n\mu) \leq x) \\ = (2\pi)^{-1/2} \int_{-\infty}^x e^{-(1/2)u^2} du \end{aligned} \quad (18)$$

the limit being the distribution function of the unit normal law.

Convergence in probability results are often obtained with the aid of Markov's inequality

$$P(|X_n| > c) \leq c^{-r} E|X_n|^r \quad (19)$$

for $r > 0$ and a corresponding convergence in the r th mean result. The useful particular case $r = 2$ of Markov's inequality is ordinarily called *Chebyshev's inequality*. For example, in the notation of equation (20),

$$P(|n^{-1}S_n - \mu| > \epsilon) \leq \epsilon^{-2}n^{-2}E(S_n - n\mu)^2 \rightarrow 0 \quad (20)$$

4 Probability Theory

as $n \rightarrow \infty$, which is a version of the weak (see **Laws of Large Numbers**), $n^{-1}S_n \xrightarrow{P} \mu$.

Almost sure convergence results rely heavily on the Borel–Cantelli lemmas. These deal with a sequence of events (sets) $\{E_n\}$ and $E = \bigcap_{r=1}^{\infty} \bigcup_{n=r}^{\infty} E_n$, the set of elements common to infinitely many of the E_n . Then

1. $\sum P(E_n) < \infty$ implies that $P(E) = 0$.
2. $\sum P(E_n) = \infty$ and the events E_n are independent implies that $P(E) = 1$.

With the aid of 1, equation (21) can be strengthened to the strong *law of large numbers* (see **Laws of Large Numbers**),

$$n^{-1}S_n \xrightarrow{\text{a.s.}} \mu \quad (21)$$

as $n \rightarrow \infty$.

The strong law of large numbers embodies the idea of a probability as a strong limit of relative frequencies; if we set $X_i = I_i(A)$, the indicator function of the set A at the i th trial, then $\mu = P(A)$. This property is an important requirement of a physically realistic theory.

The law of large numbers, together with the central limit theorem, provides a basis for a fundamental piece of statistical theory. The former yields an estimator of the location parameter μ and the latter enables approximate *confidence intervals* (see **Confidence Intervals**) for μ to be prescribed and tests of hypotheses (see **Hypothesis Testing**) constructed.

Other Approaches

The discussion above is based on a frequency interpretation of probability, and this provides the richest theory. There are, however, significant limitations to a frequency interpretation in describing uncertainty in some contexts, and various other approaches have been suggested. Examples are the development of subjective probability by de Finetti and Savage (e.g., references [4] and [5]) and the so-called logical probability of Carnap [6]. The former is the concept of probability used in Bayesian inference and the latter is concerned with objective assessment of the degree to which evidence supports a hypothesis. A detailed comparative discussion of the

principal theories that have been proposed is given in Fine [7].

The emergence of probability theory as a scientific discipline dates from around 1650.

The Literature

Systematic accounts of probability theory have been provided in book form by many authors. These are necessarily at two distinct levels since a complete treatment of the subject has a measure-theory prerequisite. Elementary discussions, not requiring measure theory, are given for example, in Chung [8], Feller [9], Gnedenko [10], and Parzen [11]. A full treatment can be found in Billingsley [12], Breiman [13], Chow and Teicher [14], Chung [15], Feller [16], Hennequin and Tortrat [17], Loève [18], and Rényi [19].

References

- [1] Kolmogorov, A. (1933). *Grundbegriffe der Wahrscheinlichkeitsrechnung*, Springer-Verlag, Berlin. (English trans.: N. Morrison, in *Foundations of the Theory of Probability*, Chelsea, New York, 1956.)
- [2] Dubins, L.E. & Savage, L.J. (1965). *How to Gamble if you Must*, McGraw-Hill, New York.
- [3] Rényi, A. (1970). *Foundations of Probability*, Holden-Day, San Francisco.
- [4] de Finetti, B. (1974). *Theory of Probability*, John Wiley & Sons, New York, Vols. 1–2.
- [5] Savage, L.J. (1962). *The Foundations of Statistical Inference: A Discussion*, John Wiley & Sons, New York.
- [6] Carnap, R. (1962). *The Logical Foundations of Probability*, 2nd Edition, University of Chicago Press, Chicago.
- [7] Fine, T.L. (1973). *Theories of Probability. An Examination of the Foundations*, Academic Press, New York.
- [8] Chung, K.L. (1974). *Elementary Probability Theory with Stochastic Processes*, Springer-Verlag, Berlin.
- [9] Feller, W. (1968). *An Introduction to Probability Theory and its Applications*, 3rd Edition, John Wiley & Sons, New York, Vol. 1.
- [10] Gnedenko, B.V. (1967). *The Theory of Probability*, 4th Edition, Chelsea Publishing, New York. (Translated from the Russian by B.D. Seckler.)
- [11] Parzen, E. (1960). *Modern Probability Theory and its Applications*, John Wiley & Sons, New York.
- [12] Billingsley, P. (1979). *Probability and Measure*, John Wiley & Sons, New York.
- [13] Breiman, L. (1968). *Probability*, Addison-Wesley, Reading.

- [14] Chow, Y.S. & Teicher, H. (1978). *Probability Theory. Independence, Interchangeability, Martingales*, Springer-Verlag, New York.
- [15] Chung, K.L. (1974). *A Course in Probability Theory*, 2nd Edition, Academic Press, New York.
- [16] Feller, W. (1971). *An Introduction to Probability Theory and its Applications*, 2nd Edition, John Wiley & Sons, New York, Vol. 2.
- [17] Hennequin, P.L. & Tortrat, A. (1965). *Théorie des Probabilités et Quelques Applications*, Masson, Paris.
- [18] Loève, M. (1977). *Probability Theory*, 4th Edition, Springer-Verlag, Berlin, Vols. 1–2.
- [19] Rényi, A. (1970). *Probability Theory*, North-Holland, Amsterdam. (Earlier versions of this book appeared as *Wahrscheinlichkeitsrechnung*, VEB Deutscher Verlag, Berlin, 1962; and *Calcul des probabilités*, Dunod, Paris, 1966.)

CHRISTOPHER. C. HEYDE

Article originally published in Encyclopedia of Statistical Sciences, 2nd Edition (2005, John Wiley & Sons, Inc.). Minor revisions for this publication by Jeroen de Mast.

Significance Level

The customary (“classical”) meaning of *significance level* (also called the α level or the *probability* of a type I error) is a property of a statistical test procedure – namely, the probability of rejecting the hypothesis tested when it is, in fact, valid (*see Hypothesis Testing; Power*). Most standard tests are constructed to control the level of significance at a prespecified value, commonly 0.01 or 0.05 (1% or 5%). Statistically significant results are sometimes reported in the scientific literature by adding a star (*) to the effect.

More recently (see, e.g. Gibbons and Pratt [1]), the term *level of significance* of a test has been applied to *data*, being the least numerical level of

significance at which the hypothesis tested would be rejected by the test if applied to the data. This is also called the *P value* (*see P Values*) for the data. Note that the lower the level of significance, in this sense, the more “highly significant” of departure from the hypothesis – in conventional phraseology – is the result of the test.

The distinction between *significance* – which relates to evidence of departure – and *importance* or *relevance* – which relates to the size and nature of the departure – is widely emphasized but often forgotten.

Reference

- [1] Gibbons, J.D. & Pratt, J.W. (1975). P-values: interpretation and methodology. *The American Statistician*, **29**, 20–25.

Overview of Statistics

Introduction

The term statistics is used in different ways. “Accident statistics” or “sales statistics” refer to numerical data in these areas. (The corresponding theoretical terminology defines statistics to be any functions of observable random variables.) Common problems encountered in the work with such data and those collected by scientists, engineers, government officials, lawyers, doctors, and so on have led to the development of general methods and principles concerning the collection, presentation, and analysis of data. The term statistics also denotes the discipline concerned with such methods.

This article considers statistics broadly as the field comprising all of the above concerns. It might be described as the enterprise dealing with the collection of data sets, and extraction and presentation of the information they contain.

New methods and formulations are constantly being developed, and new types of application come into view and in turn give rise to new problems. This flow of research is disseminated through a multitude of journals being published around the world. A listing of annual contributions is available in the *Current Index to Statistics: Applications, Methods, and Theory*, published since 1975. Much of the material prior to 1970 is listed in the five-volume *Index to Statistics and Probability* (Ross and Tukey, eds. R and D Press).

What establishes statistics as a discipline is that the same kind of data requiring the same kind of concepts and methods turn up in many different fields. On the other hand, special areas of application also may require some specialization and adaptation to particular needs. The *Encyclopedia of Statistics* gives a comprehensive overview in nine volumes and a number of updates.

Data Interpretation

Statistical Methodology

It is a central fact, underlying essentially all statistical thinking, that an actual data set typically is only one of many possible such sets that might have

been obtained under the given circumstances. (For a possible exception, see Diaconis [1].) Measurements vary when they are repeated. A store inventory, besides depending on the day on which it is taken, will be affected by bookkeeping and counting errors. In a survey, data will be obtained from different households if a new sample is to be drawn; even if the same households are visited on another occasion, different members may be at home and provide different answers; and, finally, even the same family member may answer the same questions differently on another day. As a consequence, interpretation of a data set depends not only on the actual data, but also on what (if anything) is assumed about the possible alternative observations that might have been obtained instead. The following sections consider three categories of such assumptions, and briefly indicate the kinds of statistical procedures that can be based on each.

Data Analysis

It is rare that a data set is studied without any preconceived notions. Consider, however, the idealized approach of pure data analysis in which the data are considered on their own terms. The statistical methods developed on this basis have as their primary aim (a) exploration of the data to uncover the features of principal interest and (b) presentation of the data in a manner that will bring out and highlight these features. The set of techniques dealing with (a) and (b) are called *exploratory data analysis* (see **Exploratory Data Analysis**) and *descriptive statistics* (see **Descriptive Statistics**), respectively. Both employ a great variety of numerical and *graphical* (see **Graphical Representation of Data**) techniques.

The simplest, most basic data-analytic methods concern a single batch of numbers, for example, the first-year sales of 12 novels of a successful author, 40 measurements of corrosion taken at different locations on a copper plate, or the ages of 350 000 cancer patients listed in a tumor registry. Histograms, stem and leaf displays, and one-dimensional scatter plots (see **Graphical Representation of Data**) are some of the many ways of presenting the numbers of a batch. From these one can get an impression of where the data are centered (the general level of their values), and how spread out they are. Observations that

lie far from the bulk of the data, the so-called *outliers* (see **Outliers**), may correspond to errors or exceptional cases and may deserve special attention. The display may exhibit unusual features such as bimodality, suggesting perhaps the possibility of a mixture of two batches, each unimodal but with different *modes*. If there is marked asymmetry, the possibility of making a transformation (see **Box and Cox Transformation**) (such as taking logarithms) to obtain a more symmetric data set arises.

Instead of a fairly detailed display of the data, authors frequently present only one or two summary statistics; for example, the mean or median to indicate the general level of the numbers (see **Measures of Location**), and a measure of their variability, such as the standard deviation (SD), median absolute deviation, or interquartile range (see **Measures of Scale**). A compromise is the *five-number summary*, consisting of the smallest, largest, and median observation, and the first and third *quartiles* (see **Quartiles**). These may be displayed graphically as a box plot (see **Graphical Representation of Data**).

Additional information concerning the numbers of a batch may be available and are important. If the order in which the 12 novels appeared is known, the data may show that the sales steadily increased, or that they increased up to a certain point and then leveled off, and so on, thus providing an indication of the author's changing reputation and success.

Example 1: Corrosion Data The exploratory use to which even very simple data can be put is illustrated by 40 corrosion readings taken at random over a metal plate (Campbell [2], Wolfowitz [3]). No particular pattern emerged when the observations were plotted according to their position on the plate. However, when they were plotted in the order in which they were taken, a bunching together in *runs* (see **Runs, Runs Tests**) of high and low observations suggested (as it turned out, correctly) a malfunctioning of the delicate measuring device.

Much of data analysis deals with batches of more complex units, such as pairs of numbers, more general vectors, matrices, or curves, and it need not be restricted to a single batch. With multivariate data, one may be interested in the separate features of the different variables, in relationships among the variables or sets of variables, or in aspects of the

overall pattern of the multidimensional data set. The development of graphical (including, in particular, computer-displayed) and numerical methods for these purposes is a very active area of statistical research (see **Multivariate Analysis**). It includes, among others, approaches such as cluster analysis, and multivariate classification, multidimensional scaling, pattern recognition, and factor analysis. When several batches are being considered simultaneously, comparisons of the batches will tend to be of primary interest. Given today's large volumes of data collected from industrial processes, much recent research focuses on the adoption of multivariate techniques for studying high-dimensional process data (see [4] and [5]).

Statistical Analysis Based on Probability Models

The data-analytic approach indicated in the preceding section can provide clarification of the phenomena represented by both simple and very complex data sets, and can lead to important new insights and hypotheses. In its simplest forms such as numerical summaries and histograms, it is the statistical presentation that is most frequently encountered by the general public in newspaper articles and magazine reports. (Through misleading use, it also lends itself to much mischief. See for example, Huff [6].) However, this approach lacks an often essential requirement of the resulting inferences: because of the fact mentioned earlier, that the same phenomena might have led to different observations, it is impossible to assess the reliability of the conclusions.

Such an assessment requires knowledge concerning the alternative data sets that might have been observed in the given situation, instead of the set that was observed. The crucial step underlying modern theories of statistical inference and decision making is to consider the observed data as realizations of random variables. The possible values of these variables are governed by probability distributions specifying the probabilities of observing the various possible data sets. These distributions are not assumed to be known, but only to belong to a postulated family of possible distributions. The lack of complete specification represents the unknown aspect of the situation for the clarification of which the data were collected.

For a single batch $X_1, \dots, X_n$, for example, a set of n measurements of some quantity, investigators often assume that the n observations are independent,

and that each has the same probability distribution. If this common distribution is denoted by F , the model is completed by specifying a family $\mathbb{F}$ to which F is assumed to belong. This may for example be the family of all possible distributions F (i.e., no further assumptions are made) or the family of all distributions that are symmetric with respect to a specified point of symmetry. Such broad families are called *nonparametric* in distinction to parametric families where F is known except for the values of some parameters, for example, the family of all normal, Poisson, or Weibull distributions (see **Probability Density Functions**). Once the model is specified, one can now ask, and answer, the type of question concerning a single batch considered in the preceding section, with more precision. In particular, the conclusions no longer refer to this particular data set, but to the underlying process that produced it.

Suppose for example that the aspect of interest is the overall level previously represented by the average of the observations. Suppose that $\mathbb{F}$ is the family of normal distributions with mean ξ and unit variance. Then ξ is the average value, not of these particular n measurements but rather of all potential measurements, each weighted according to its probability. The average $\bar{X} = (X_1 + \cdots + X_n)/n$, which was a descriptive measure of the general level of the batch earlier, now becomes an estimator of the unknown ξ : for example, the true value of the quantity being measured or the average value of the characteristic (e.g., height, age, or income) in the population from which the X 's were obtained as a sample.

The model assumptions make it possible to get an idea of the accuracy of an estimator such as $\bar{X}$, for example in terms of the expected closeness to the true value ξ . One commonly used measure of this closeness is the expected squared error, which for the case of n independent measurements with variance 1 equals

$$E(\bar{X} - \xi)^2 = \frac{1}{n} \quad (1)$$

This formula shows for instance how the accuracy improves with n , and enables one to determine the **sample size** n required to achieve a given accuracy.

This approach also provides a basis for comparing the accuracy of competing estimators. For the median $\tilde{X}$ of the X 's for example, one finds that approximately (if n is not too small) $E(\tilde{X} - \xi)^2 = 1.57/n$ – more than 50% larger than the corresponding value

$1/n$ for $\bar{X}$. The median is thus considerably less accurate than the mean. This conclusion depends strongly on the assumption that the X 's are normally distributed. For other distributions the result may be just the reverse (see **Estimation**).

Another way of describing the accuracy of $\bar{X}$ is obtained by noting that

$$P\left[|\bar{X} - \xi| \leq \frac{1.96}{\sqrt{n}}\right] = 0.95 \quad (2)$$

so that with probability 0.95 the estimator $\bar{X}$ will differ from the true value ξ by less than $1.96/\sqrt{n}$. The statement (1) can be paraphrased by saying that the random interval $(\bar{X} - 1.96/\sqrt{n}, \bar{X} + 1.96/\sqrt{n})$ will contain the unknown ξ with probability 0.95. Random intervals that cover an unknown parameter ξ with probability greater than or equal to some prescribed value γ are called *confidence intervals* (see **Confidence Intervals**) at confidence level γ .

Point **estimation** and estimation by confidence intervals provide two of the classical approaches to statistical inference. The third is *hypothesis testing* (see **Hypothesis Testing**).

Example 2: Extrasensory Perception Suppose the claim of a subject A to have extrasensory perception (ESP) is to be tested by tossing a coin 100 times at a location invisible to A, and recording A's perception (heads or tails) for each toss. Suppose that A obtains the correct result on 54 of the tosses. This clearly is not an indication of a strong ability at ESP. However, even a very slight ability would be of extraordinary interest. Is there support for such a finding, or is the result compatible with purely random guessing? Under the null hypothesis of pure guessing, the probability of calling a toss correctly is $1/2$ for each toss, and the probability of getting 54 or more right is then $1/4$. The result is therefore not particularly surprising under H , and a case for A's ability to do better than chance has not been made. Had A correctly identified 65 tosses rather than 54, the conclusion would have been quite different. The probability of 65 or more correct calls is only 0.002 when H is true. In the light of so extreme a result, one would have to give serious attention to A's claim.

Hypothesis testing, and point and interval estimation are all used extensively in applications, with each being more useful and popular in some areas than in others. The theory of these three methodologies

is concerned with the performance of proposed procedures (including the determination of sample size to achieve a desired performance), the comparison of different procedures, and the determination of optimal ones. An important consideration is the robustness of a given procedure under violation of the model assumptions. If the procedures are very sensitive to these assumptions, one may want to study the problem in a nonparametric setting of the kind described for a single batch at the beginning of this subsection. Nonparametric (and semiparametric) models are particularly important for the analysis of large data sets, which has become more practicable as a result of increased computer capabilities and availability.

A unified framework for the three areas is provided by Wald's decision theory. This very general theoretical approach deals with the choice of one of a set of possible decisions d on the basis of observations x_i . Let the chosen d be denoted by $\delta(\mathbf{x})$. The observed value $\mathbf{x}$ is assumed to be the realization of a random quantity $\mathbf{X}$ with probability density $p_\theta(\mathbf{x})$, θ unknown. A loss function $L(\theta, d)$ measures the loss resulting from decision d when θ is the true parameter value, and the performance of a decision procedure δ is measured by its *risk function*

$$R(\theta, \delta) = E_\theta [L(\theta, \delta(\mathbf{X}))] \quad (3)$$

that is, the average loss incurred by δ when θ is true. A principal concern of the theory is the determination of a δ for which the risk function is as small as possible in some suitable sense. Another problem is the characterization of all admissible procedures, i.e., all procedures whose risk cannot be uniformly improved.

Decision theory can also be made to encompass sequential analysis, the **design of experiments** (by letting the loss function take account of the cost of observations), and the choice of model (by imposing a penalty that increases with the complexity of the model). However, it is an abstract approach that has been useful primarily for exhibiting general relationships rather than for its impact on specific methods. An important consequence of Wald's theory has been its liberating effect on the formal consideration of new types of situations, among them multiple comparisons and other simultaneous inference procedures.

Exploration versus Verification

To illustrate both the relation and the difference between the approaches described in the preceding two sections, consider once more the ESP example.

Example 2: Extrasensory Perception (Continued)

When looking at the results of the 100 tosses, suppose the experimenter notices that of the 54 successes, 33 occurred during the last and only 21 during the first 50 tosses. The probability of 33 or more successes in 50 tosses with success probability $p = 1/2$ is only about 0.01. One might be tempted to explain away the poor performance in the first half of the experiment by the theory that it requires some warming up before the ability hits its stride, and to declare the success of the second half significant. However, such a conclusion would not be justified on the basis of this analysis since the calculation does not take account of the fact that the particular test (restricting oneself to the second half) was not originally planned but was suggested by the data.

Suppose that the situation had been reversed, that there had been 33 successes in the first and 21 in the second half. A possible explanation is that the exercise of ESP requires great concentration, and after 50 attempts the subject is likely to get tired. Similar explanations could have been found if success had been unusually high in the middle 50, or on the last 25, or the first 25, and so on. Instead of high concentration of successes in a particular segment of the sequence, other patterns might have struck an observer: for example, a gradual rise in the frequency of success, or a cyclic pattern of successes and failures, and so on. For each, an explanation could have been found.

This is the problem of *multiplicity*. Every set of observations – even a completely random one – will show some special features, and explanations can usually be found *post facto* to account for them. Unfettered examination of many different aspects of a data set is legitimate, and in fact a primary purpose, of exploratory data analysis. However, the results will then tend to look more impressive than they really are since attention is likely to focus on the extremes of a large number of possibilities. To legitimize a theory suggested by the data, one must test it, for example, on a separate part of the data not used at the exploratory stage or from new data specially obtained for this purpose. (A

somewhat more limiting alternative to such a two-stage procedure is to formulate a number of possible theories to be considered before any observations are taken. The simultaneous testing of such a number of possibilities provides a legitimate calculation for the probability of the most striking of the associated results.)

The two stages, exploration of the data leading to the formulation of a hypothesis, followed by an independent test of this hypothesis, constitute the basic pattern of scientific progress as described by scientists (see for example, Feynman [7, Chapter 7]) and discussed by philosophers of science. Before considering a third aspect in the next section, let us briefly mention another distinction, which relates to the purpose of a statistical investigation. This is the difference between inference and decision making. It may be illustrated by the problem faced by a doctor who wants to arrive at a diagnosis (inference) but must also select a treatment (decision).

Example 3: Medical Diagnosis Suppose there are k possible conditions (diagnoses) $\theta_1, \dots, \theta_k$ that might have led to the observed symptoms and test results $\mathbf{X}$ (the data), including for example measurements of temperature, blood pressure, and so on. Under condition θ , the observations $\mathbf{X}$ have a distribution $p_\theta(\mathbf{x})$. The problem of diagnosis is thus the statistical problem of using the observed value of x to determine the correct value of θ . (A standard procedure is to select the value $\hat{\theta}$ of θ that maximizes $p_\theta(\mathbf{x})$ for the given observation $\mathbf{x}$, the so-called maximum-likelihood estimate (see **Maximum Likelihood**) of θ .) The associated decision procedure might be to select the treatment that would be most appropriate for the chosen θ . The situation is however more complicated since the choice of treatment must also take account of the severity of the consequences in case of an incorrect diagnosis resulting in a nonoptimal treatment, the so-called loss function. The distinction between inference and decision making is reviewed in Barnett [8]. For a discussion of computer-based medical diagnosis and treatment choice, see Shortliffe [9] and Spiegelhalter and Knill-Jones [10].

Bayesian Inference and Decision Making

Example 3: Medical Diagnosis (Continued) Suppose a patient P being tested for the conditions

$\theta_1, \dots, \theta_k$ of the last example is told that the tests point to $\hat{\theta} (= \theta_1, \text{ say})$ as the most likely cause of the symptoms. Naturally, P wants to know just how likely it is that θ_1 is in fact the true cause. The doctor has to admit that the term *most likely* was used imprecisely; that $\hat{\theta} = \theta_1$ is the condition that assigns the highest probability to the observed test results, not necessarily the most likely of the conditions $\theta_1, \dots, \theta_k$, and that in fact no probability can be assigned to these conditions.

Actually, in this example it may be possible to make such an assignment. Suppose that π_i is the incidence of condition θ_i in the population of sufferers from the given symptoms, and hence is the probability that a patient drawn at random from the population of such sufferers has condition θ_i . The condition of such a patient is then a random variable Θ , with prior probability $P(\Theta = \theta_i) = \pi_i$ before the tests are taken, and posterior probability $P(\Theta = \theta_i | \mathbf{x})$ in the light of the test results $\mathbf{x}$. The latter probability can be calculated from the π_i and the $p_{\theta_i}(\mathbf{x})$ by Bayes' theorem.

On being presented with this probability, P may however still not be satisfied but complain to the doctor: "You have treated me for 20 years, you know my complete medical history, lifestyle, and habits. In the light of all this additional information, what is the probability of suffering from θ_i not for a random patient but for me personally?" Unfortunately, the interpretation of probability as frequency, which was tacitly assumed up to now, and which in particular applied to the prior probabilities π_i of the preceding paragraph, precludes assigning probabilities (other than 0 or 1) to unique events such as this particular patient's suffering from condition θ_i .

This difficulty is met head-on by the Bayesian approach according to which π_i can be chosen to represent the physician's probability that condition θ_i obtains for this particular patient. This meaning of probability is, however, different from the earlier one. Probability is no longer a frequency but the degree of belief attached to the event in question. (If the event is repetitive, this probability typically draws close to the observed frequency as the number of cases gets large.)

In general, the Bayesian approach assigns a *prior probability distribution* to a parameter θ before the observations $\mathbf{X}$ are taken. Once the values $\mathbf{x}$ of $\mathbf{X}$ are available, the prior distribution of θ is updated to

the posterior (i.e., conditional) distribution of θ given $\mathbf{X} = \mathbf{x}$, which shows how the prior beliefs regarding the chances of different θ values have changed in the light of the data.

In the decision theoretic terminology introduced earlier, the relevant assessment of the performance of a procedure δ from a Bayesian point of view is not the risk function $R(\Theta, \delta)$ but rather the posterior expected loss

$$r(\mathbf{x}, \delta) = E [L(\theta, \delta(\mathbf{x})|\mathbf{x})] \quad (4)$$

calculated according to the conditional distribution of θ given $\mathbf{x}$.

Ideally a Bayesian has a comprehensive, consistent view of the world with a probability attached to every unknown fact. These probabilities must satisfy an appealing set of axioms (the axioms of *coherence*), and must be updated as new information becomes available. For an individual's response to the world (or even a specific problem), a chief difficulty in implementing this program is the determination of the prior distribution. A considerable literature deals with methods for eliciting a person's degrees of belief with respect to a given situation, but there is also an *objective Bayesian school* in which the prior distribution represents "total lack of information".

Ideally, each person has a single correct personal probability regarding any unknown event. However, in practice, these probabilities "can never be quantified or elicited exactly (i.e., without error) especially in a finite amount of time" (Berger [11, p. 64]). This has led to the suggestion of a robust Bayesian viewpoint according to which one "should strive for Bayesian behavior which is satisfactory for all **prior distributions** which remain plausible after the prior elicitation process has been terminated" (Berger [11]).

A second difficulty arises in situations in which the decision or opinion concerns not a single person, but represents a joint problem for a group or occurs in the public domain, as for example in the publication of the analysis of a scientific investigation. Some aspects of this problem will be considered in the next section.

The Bayesian/Frequentist Controversy

The mutual criticism of Bayesians and frequentists has given rise to a lively (and sometimes acrimonious) debate, which has helped to clarify a number

of basic statistical issues. (There are many different variants of Bayesians and frequentists, not all of which will agree with the positions ascribed to these approaches here.) One of the central concerns is that of subjective *versus* objective data evaluation in scientific inference and reporting. Fisher, Neyman, and Pearson developed their frequentist theories in a deliberate effort to free statistics from the Bayesian dependence on a prior distribution, and this aspect has continued as the central frequentist objection. The Bayesian response to this criticism is twofold. On the one hand, it is pointed out that frequentist analysis involves similar types of specification. There is the choice of model and loss function, both of which must be chosen in the light of previous experience and involve judgments that are likely to vary from one person to another. In addition, there is the problem of selecting a frame of reference that forms the basis of the frequency calculations. In assessing the incidence of the conditions $\theta_1, \dots, \theta_k$ of the preceding section, for example, for what population should this be calculated: the population of the world, or the country, or the city in which the person lives; should the comparison be restricted to patients of the same age, sex, and so on? This is the problem of conditional inference, which so far has found no satisfactory frequentist solution. (For the Bayesian, the problem does not arise in this form since the probabilities will always refer to this particular patient, but the same issue arises when one must decide how to weigh the experience with other patients in forming an opinion about this one.)

As a more positive response, there have been Bayesian suggestions (for example, Dickey [12], Smith [13]), that in scientific inference the analysis should be reported under a variety of different priors that – it is hoped – will include the opinions of the readers. The view that "any approach to scientific inference which seeks to legitimize *an* answer to complex uncertainty is a totalitarian parody of a would-be rational human learning process" (Smith [13]) considerably narrows the gulf between the two approaches.

The reason for this narrowing can be found in the combination of two facts. The first is a basic result of Wald's (frequentist) decision theory to the effect that every admissible procedure is a Bayes solution or a limit of Bayes solutions. Secondly, frequentists tend not to believe in a unique correct approach, and may

therefore try a number of different solutions corresponding to different optimality principles, robustness properties, and so on. If these lead to similar conclusions, any one of them can be adopted. Otherwise, a careful examination of the differences may clarify the reason for the discrepancies and point to one as the most appropriate. Lacking such a resolution, one may instead prefer to report a number of different procedures.

The Bayesian and frequency approaches lead to different ways of assessing the performance of a decision procedure. From a strict Bayesian point of view, only the posterior distribution of θ given $\mathbf{x}$, and the posterior expected loss $r(\mathbf{x}, \delta)$, are relevant, while frequentists measure the performance of a procedure by its risk function. However, on this issue also, an accommodation to statistical practice and the need for communication has narrowed the gap by generating a Bayesian interest in risk functions. Thus Rubin [14, p. 116] writes, “Frequency calculations that investigate the operating characteristic of Bayesian procedures are relevant and justifiable for a Bayesian when investigating or recommending procedures for general consumption.” Similar considerations can be found in Berger [15].

In the other direction, the frequentist approach has been strongly influenced by Bayesian ideas, in particular, by recognizing that it is natural and useful to consider the prior distribution leading to a proposed admissible procedure. An example in which such a Bayesian interpretation has made an important contribution to a theory developed in a decision theoretic framework is that of Stein estimation.

Data Acquisition

Measuring Single Units

The first part of this entry dealt with the interpretation of data once they have been collected. In this and the following sections, we consider some of the processes that produce data, and the statistical problems arising at this earlier stage. (An extensive discussion of various types of data is provided in Hoaglin *et al.* [16].)

The basic data units are the numbers, symbols, words, or other entries making up a data set; for example, measurements of the height of a person or plant, or of the weight of a wagonload of fruit or a minute amount of some chemical; barometer or

temperature readings; or the scores of a student in an intelligence or aptitude test. Alternatively, the observations could be the information provided by a person answering a questionnaire or an interviewer, such as family size, last year’s income, or religious and political affiliation.

A concern for **data quality**, for their reliability and validity, is an important task preceding the collection of data (*see Measurement Systems Analysis, Overview*). Efforts must be made to eliminate **bias**, reduce variability, and eliminate sources of error. The repeatability and reproducibility of measuring gauges should be assessed, and if needed improved. In constructing a questionnaire, great care is required to avoid ambiguities. Checks on the reliability of responses can often be built into the data set, for example by asking for the same information in a number of different ways or in different contexts.

Another aspect of data improvement arises after the data have been obtained. Data editing (or “cleaning”) involves the deletion or modification of entries that do not appear to be in consonance with the rest of the data and that sometimes represent obvious errors (for example, in a series of monthly measurements of a baby’s head circumference when one month’s measurement is smaller than the preceding ones). A variety of statistical methods have been developed for this purpose. In addition, robust statistical procedures are available that satisfactorily control the effect of outlying observations (*see for example Hoaglin et al.* [17] and Hampel [18]). On the other hand, data cleaning by inspection, without clearly stated rules, runs the risk of introducing a subjective element. (For example, just which observations are to be singled out for such treatment?) Even with good rules, it may remove valuable evidence and destroy the basis for probability calculations. For this reason, it is typically better to consider the editing of data as part of the statistical analysis and, while perhaps indicating definite or suspected errors, not to change the original data.

The improvement of data quality mentioned above is not the only aspect of **data collection** to be considered before obtaining the observations for an investigation. Two questions, in particular, that are of crucial importance and that will be discussed in the next sections are how many and what kind of observations are required.

Assessing Population Characteristics

The target of most investigations is not a single

unit but a population of such units (or the process generating the units). A set of professors, apples, days, mice, or light bulbs is typically examined not because of interest in these particular specimens but in order to reach conclusions about the populations they represent. (It is this aspect and the associated quantification that has earned statistics its reputation of being antihumanistic.)

Efforts to obtain information about every member of the target population through a census go back to at least the third millennium B.C. in Babylonia, China, and Egypt. (For a nontechnical history of the census, see Alterman [19].) The purpose was usually to provide a basis for taxation or proscription. Today, population censuses seeking a great variety of information are carried out, many of them at regular intervals, in nearly all countries of the world.

However, taking a complete inventory is not the only way to obtain accurate statistical information about a population (and in the case of an infinite population, such as the one consisting of all products that a process manufactures, it is not even possible). The same end can usually be achieved more economically by collecting information only from the members of a sample, taken from the population by a suitable sampling method. The much smaller size of the sampling operation tends to make this procedure not only cheaper but also more accurate since it permits better control of the whole process and hence the quality of the resulting data. On the other hand, a census has the advantage – not shared by any sample – of providing information even for very small subpopulations.

Suppose a population Π consists of N units, each of which has a value v of some characteristic of interest, such as the age or income in the case of a human population, the number or weight of the apples on each of the N trees in an orchard, or the length of life or brightness of the lightbulbs in a shipment received from the factory. To obtain an estimate of the average v value $\bar{v} = (v_1 + \cdots + v_N)/N$, a sample is taken from the population. In the early days of sampling, investigators usually relied on judgment samples, in which judgment is used to obtain a sample “representative” of the population. It is now realized that such sampling tends to lead to biases, and does not provide a basis for calculating accuracy or the sample size needed to achieve a desired accuracy. The methods used instead are probability **sampling schemes** that select the units to be included in the sample according to stated probabilities. The simplest

such scheme is simple random sampling, according to which n units are chosen in such a way that every possible sample of size n has the same probability of being drawn. The average of the sample v values is then the natural estimate of $\bar{v}$.

Better accuracy can often be attained, and the needed sample size and resulting cost therefore reduced, by dividing the population to be sampled into strata within which the v values are more homogeneous than in the population as a whole but which differ widely among each other. (The population of school children of a city might for example be stratified by school, grade, and gender.) A stratified sample of size n is then obtained by drawing a simple random sample of size n_i from the i th stratum for each i , where $\sum n_i = n$.

In a different direction, the cost of sampling can often be reduced by combining units into clusters (for example, all the apartments in an apartment house, all houses in a city block, or all patients in a hospital ward), and obtaining the required information for each member of a sampled cluster. The two approaches can be combined into stratified cluster sampling, and many other designs are possible.

Sample surveys to obtain information about a population are in widespread use and have become familiar through election polls, market surveys, and surveys of television viewers to establish ratings. However, they are not very well understood. In particular, it appears puzzling how it is possible to obtain an accurate estimate of the opinions or intentions of many millions from information concerning just a few thousands.

Example 4: Election Poll To get some insight into this question, let us simplify the situation, and suppose that a population Π consists of a large number N of voters, each of whom supports either candidate A or B. A simple random sample of n voters is drawn from Π , and the preference of each member of the sample is ascertained. If X is the number of voters in the sample favoring A, then X/n , the proportion of A supporters in the sample, is the natural estimate of the proportion p of A supporters in the population.

The question at issue is how large a sample is required for X/n to achieve a prescribed accuracy as an estimate of p , for example, for the standard

deviation (SD) of X/n to satisfy

$$\text{SD}\left(\frac{X}{n}\right) \leq 0.01 \quad (5)$$

It is often felt intuitively that the required sample size n should be roughly proportional to the population size N . However, it turns out that this intuition is misleading and that in fact, for large populations, the required n is essentially independent of the value of N .

To see this, suppose for a moment that the sampling is done “with replacement”, i.e., that the members are drawn successively at random, with each – after giving the required information – being put back into the population before the next member is drawn, again at random. This method is slightly less efficient than the original simple random sampling (and therefore requires a larger sample size), because it allows the same unit to be drawn more than once. It is introduced here because it is particularly simple to analyze. Sampling with replacement is characterized by two properties: (a) on each draw, the probability of obtaining an A supporter is p ; (b) the results of the n draws are independent in the sense that the probability of getting an A supporter on any given draw does not depend on the results of the earlier draws.

Let us now compare this situation with a quite different one. Suppose a coin with a probability p of falling heads when spun on its edge (this probability may be far from $1/2$), is spun n times and that X is the number of times it falls heads. Then (a) the probability of heads is p on each spin, and (b) the results of the n spins are independent. The standard deviation of X/n is therefore the same for this coin problem as in the election sampling with replacement. The number n required to reduce this standard deviation to 0.01 is therefore also the same in both cases. Note however that the coin problem involves only n and p ; no population is involved. Therefore the required n cannot depend on N .

This argument depends of course crucially on the assumption that the sampling was random. It is the randomness that insures that with high probability the sample contains approximately the same proportion of A supporters as the population. In addition, it was tacitly assumed that the preference of each voter can be ascertained without error. In practice, the possibility of “response error” can rarely be ruled out.

Data from Experiments

A study is called an *experiment* if its data are produced by an intervention (i.e., do not occur naturally) for the purpose of gathering information. It is a *comparative experiment* if its purpose is to compare several ways of doing something (e.g., different machines, procedures, medical treatments, fertilizers, and so on) rather than to determine some absolute value.

Example 5: Weather Modification Consider a company’s claim to be able to increase precipitation by seeding the clouds of suitable storms. Here the comparison is between seeding and not seeding. How can one obtain data to test the claim and to provide an estimate of the amount of increase?

As one possibility, the company might seed all suitable storms in the given location, and compare the resulting rainfall with that during the corresponding period in the preceding years. Of course, if the results are very striking (for example, if an enormous downpour occurs immediately following each seeding) this may settle the issue. However, typically the results are less clear. Suppose, for example, that the total rainfall matches, or even slightly exceeds, that of the wettest of the last five years. This may be the result of the seeding; or it may just be the consequence of an exceptionally wet year.

This difficulty is inherent in studies in which there is no randomness in the assignment to the experimental units of the conditions or treatments being compared. A better basis for the establishment of a causal relationship – in this case that the increased rainfall is due to the seeding – is obtained if storms are compared within the same season, and if they are assigned to the two treatments (seeding and not seeding) according to a random mechanism. This can be done in a variety of ways, corresponding to different experimental designs.

As a simple possibility, suppose that the experiment is to extend to the first 20 storms that are suitable for seeding. Of these, 10 will be seeded and 10 not, the latter providing the *controls*. According to one design (**complete randomization**), 10 of the numbers $1, \dots, 20$ are selected at random; the storms bearing these numbers will be seeded, the remaining 10 will not. An alternative design (*paired comparisons*) pairs the storms $(1, 2), (3, 4), \dots, (19, 20)$, and within each pair assigns at random (e.g., by

tossing a coin) one storm to seeding, the other to control. This second design will be particularly effective if storms occurring close together in time are more likely to be similar in strength than storms separated by longer intervals.

It is interesting to note the close relationship of these designs to the sampling schemes of the preceding section. In the first design, the storms assigned to seeding are a simple random sample of the 20 available storms; in the second case, they constitute a stratified sample, with samples of size 1 from each of 10 strata of size 2.

Unfortunately, random assignment of the conditions being compared is not always possible. Consider for example a study that reports that married men tend to live longer than unmarried ones. It is tempting to draw the conclusion that marriage prolongs life, perhaps by providing a more regulated lifestyle. However, such a conclusion is not justified. Married and unmarried men constitute groups that differ in many ways. The latter for instance includes men with health problems that preclude marriage and that also tend to shorten life. It is thus not clear whether the observed effect is the result of the difference in marital status or of conditions leading to this difference. All that can be safely concluded is that marriage is associated with longer life expectancy. This may be enough for an insurance company but does not answer the sociological or public health question regarding the effect of marriage. Quite generally, *observational studies* (see **Observational Studies**) in which subjects have not been assigned to the conditions being compared (marital status, smoking habits, religious beliefs, and so on) can establish associations (such as that between marital status and longevity) but have great difficulty in validating causal relationships. To establish *causation* (see **Causality**) is a cherished goal of statistical methodology but tends to be rather elusive. (For some further discussion of this point, see for example Mosteller and Tukey [20, pp. 260–261] and Holland [21].)

Example 6: Headache Remedy Suppose that you wake up with a headache, take a headache remedy, and an hour later find that the pain is gone. On the basis of this isolated instance (which, in words attributed to R. A. Fisher, is an experience rather than an experiment), it is clearly not possible to conclude that the result is due to the remedy. The cause might instead have been another intervening event such as

breakfast, or possibly the headache had run its course and would have dissipated in any case.

The attribution of the cure to the medication becomes much stronger if the incident is not isolated. If you have suffered from headaches before, took the medicine sometimes soon after onset, at other times only after having waited in vain for the pain to subside on its own; if on different occasions you took the tablets at different times of day, and if under all these different circumstances the headache disappeared shortly after treatment and never (or only very rarely) before, the robustness of the effect over many different conditions would tend to carry conviction where a single instance would not. The finding could be strengthened further if your own experience could be merged with that of others. (However, even if the evidence for a treatment effect is convincing, observations of this kind cannot determine whether the ingredients of the medication are responsible for the improvement, or whether it is a *placebo effect*, i.e. the patient's belief in the effectiveness of the treatment, which relieves the pain.)

To see how to systematize the informal reasoning of the preceding paragraph, consider another example.

Example 7: New Traffic Signs Suppose that four-way stop signs have been installed at a dangerous intersection, and that four accidents had occurred in the month preceding the installment of the signs but only one in the month following it. These facts by themselves provide little basis for an inference. Suppose for example that the monthly accidental frequencies constitute a purely random sequence of which the value 4 was by chance high, which however caused the installment of the signs. Then even without this intervention, a decrease could have been expected in the succeeding month.

A more meaningful comparison can be obtained if the accident statistics are available not only for one month before and after, but for several months in both directions. One is then dealing with an interrupted times series, for which it is possible to compare the observations before with those after the “interruption” (the installment of the signs). Even the nonrandom choice of the time of intervention would then have only a relatively minor effect.

While the interrupted **time series** can establish that the accident rate is lower after installment than

it was before, it cannot establish the new signs as the cause of the decrease since other changes might have occurred at the same time. To mention only one possibility, the community may have been affected by editorials published at the time of the third and fourth accidents. Some control – although not as firm as that resulting from randomization – can be obtained by studying the accident rates during the same months at some other intersections also, a multiple time series design. If the other series do not show a corresponding decrease, this will clearly strengthen the causal argument.

Quasi-experiments (a better term might be “quasi-controlled experiments”) such as the interrupted time series and multiple time series described above, which try to identify and control the most plausible alternatives to the treatment being the cause of the change, are treated by Cook and Campbell [22] and Cochran [23] (for an elementary discussion see Campbell [24]). They can greatly strengthen the causal attribution although they cannot be expected to be as convincing as a controlled randomized experiment.

Serial Data

Most of the data considered so far were assumed to be collected on a one-shot basis. Often, they are obtained instead consecutively over a period of time. Such data are particularly useful in assessing changes over time. (Are winters getting colder? Is the birth rate in the United States falling? Is a process producing products of stable quality?)

An important application is provided by statistical **quality control**. An established production process is monitored by taking observations of the quality of the product at regular intervals. The process is said to be in control if the successive observations are independent and have the same distribution. A *control chart* (see **Control Charts, Overview**) on which the successive observations are plotted provides signals when the process appears to be going out of control. A similar approach is used in monitoring the cardiogram of a heart patient in **intensive care**, where the data consist of a continuous graph instead of a discrete sequence of points. Analogously, seismographs provide continuous data for the monitoring of seismic activity while persons watching their weight through regular weighings provide another illustration of the discrete case.

In these applications, a single process is followed to check that no changes are occurring and to alert us when they occur. A different situation arises when one is interested in assessing the changes occurring in a population as part of a developing pattern or as the result of some intervention. For example, to study the dynamic behavior of a chemical process, or stock prices, one collects a sequence of observations. To study the pattern of growth (the *growth curve*) of children, plants, or institutions, one observes each member of a sample from the population over a period of time. In such time series data and *longitudinal studies* the observations at different times often are dependent. Probability models for sequences of dependent observations are considerably more complicated than those for sequences of independent and identically distributed (i.i.d.) variables. Such models and their statistical analysis are treated in the theory of *time series* (see **Time Series Analysis**) or more generally of stochastic processes.

Observations arising serially, for instance on patients coming to a clinic, successive books by an author, or stock market performance in successive time periods, provide an opportunity to economize by letting the size of the study be determined by the data (according to a clearly specified rule) instead of fixing it in advance. Suppose for example that a shipment of goods is being sampled to determine whether its overall quality meets certain specifications. Items are being drawn at random and examined one by one. It may then be possible to stop early when the quality of the initial observations is either very satisfactory or very unsatisfactory, while one may wish to take a larger sample when the initial observations are mixed. The working out of economical stopping rules providing the desired statistical information, and the analysis of the resulting data, is the problem treated in *sequential analysis*.

Designing Experiments

The first part of this article was concerned with the interpretation of data once they have been obtained, and the preceding sections with various processes used to acquire the data. However, preceding even this stage, it is necessary to plan how many and what kind of data will be needed to answer the questions under consideration.

As an illustration, consider once more the ESP study of Example 2, based on 100 tosses of a coin. Suppose the investigator had decided, before the experiment was carried out, to give serious attention to the possibility of ESP, provided the number of subject A's correct calls would be at least 65. As pointed out in Example 2, under the hypothesis of pure guessing, the probability of 65 or more answers is 0.002. There is thus little danger of paying attention to the claim if it has no validity.

However, is this study giving subject A a fair shake? What is the probability of getting 65 or more of the tosses right if A really does possess the claimed ability? The answer depends of course on the extent of the ability, which can be measured by the probability of calling a toss correctly. Here are some values, computed under the simplifying assumptions that the 100 calls are independent, and that the probability p of a correct call is the same for each toss.

p	0.55	0.6	0.7	0.75
$P(\text{number of correct calls} \geq 65)$				
= power	0.027	0.179	0.884	0.990

The probability in this table measures the **power** (see **Power**) of the test of the hypothesis $H : p = 1/2$, i.e., the probability of following up A's claim against various alternative values of p . It shows that the power is quite satisfactory when p is 0.7 or more, but not, for example, when $p = 0.6$. In this case, despite the large discrepancy between pure guessing and an ability corresponding to $p = 0.6$, the probability is less than 20% that serious attention would be given to the claim. To get higher power would require a larger number of tosses. For example, to increase the power against $p = 0.6$ to 0.9, while keeping the probability of following up the claim when $p = 0.5$ at 0.002, would require a sample of about 425 tosses instead of the previous 100. A statement of the goals to be achieved, e.g., the required power of a test or accuracy of an estimate, and determination of the sample size needed for this purpose, are crucial aspects of planning a study.

Taking a fixed number of observations (100 or 425, etc.) is not necessarily the best design. Suppose for example that A calls all of the first 25 tosses correctly. This may be enough to provide the desired evidence of A's ability (or suggest that something is wrong with the experiment). Thus, a sequential

stopping rule of the type indicated at the end of the previous section might be more efficient.

Consider next the problem of determining an **optimal design** when different types of observations are involved. Suppose that we are concerned with the comparison of two treatments, and that the total number of observations is fixed at $N = 2k$. Then it will typically be best to assign k subjects to each treatment, since this will tend to maximize the power of the tests and minimize the variance of the estimate of the difference. This may of course not be the case if observations on one of the treatments have smaller variance than on the other.

Finding the optimum assignment becomes much more difficult when observations are taken sequentially, and a decision is required at each stage as to which treatment to apply next. What is the best (or even a good) rule depends strongly on the purpose of the procedure. If the only concern is to decide whether the treatments are equally effective and, if not, which is better and by how much, the treatments should be assigned in a balanced way, i.e., so that each is received by the same number of subjects. However, if one is comparing two treatments of a serious medical condition, an additional consideration arises: a desire to minimize the number of patients in the study who receive the inferior treatment. A sequential procedure that has been suggested for such a case is the "play the winner rule" in which the next patient receives the treatment that at that point looks better. (For details, see for example, Flehinger and Louis [25] and Siegmund [26].) It should be noted however, that such an unbalanced procedure will require more observations than the best balanced rule, and hence will take longer to lead to termination of the study and hence to a recommendation. Thus, the inferior treatment will be assigned to fewer patients within the experimental group but to a larger number outside of it.

Comparative experiments typically involve more than the study of just one difference. Consider an experiment to investigate the effect of a number of factors (to which experimental units can be assigned at random) on some observable response such as rainfall, crop yield, or length of life. Two possible designs are one-at-a-time experiments, in which each factor is studied in a separate experiment with all other factors held constant, and *factorial experiments* (see **Factorial Experiments**), which provide a joint

study of the effects of all factors simultaneously. The latter type of design has two important advantages.

Two or more effects may interact in the sense that the effect of one factor depends on the level of the other. In a study of the effect of textbooks and instructors on the performance of students, for example, it may turn out that a textbook that works very well for one instructor is quite uncongenial to another. The existence and size of such interactions are most easily investigated by studying the various factors simultaneously. When no interactions are present, joint experimentation has the advantage of requiring far fewer observations to estimate or test the various effects than would be needed if each factor were studied separately. (For further discussion of these and related issues see for example, Cox [27] and Box *et al.* [28].)

Much of the theory of factorial designs was originated by R. A. Fisher in the 1920s, in the context of agricultural field trials, where it continues to play a central role; in addition, its uses have since expanded to many other areas of application. The paper on **response surface** analysis by Box and Wilson [29] marked the starting point of their adoption in industry.

Conclusion

Statistics has been discussed in this article as dealing with the collection and interpretation of *data* to obtain information. We have been particularly concerned with the uncertainty caused by the fact that the observations might have taken on other values, and with the resulting variability of the data. An alternative view of statistics is obtained by noting that uncertainty attaches not only to data, but also to many other “chancy” events such as length of life, the occurrence of accidents, the quality of a manufactured article, or the fall of a die. Correspondingly, statistics has also been defined as the subject concerned with understanding, controlling, and reducing *uncertainty*. Actually, there is little difference between these two descriptions: data involve uncertainty, and the study of any particular uncertainty requires the collection of appropriate data. It is a matter of emphasis.

Earlier sections have given a general indication of the pervasiveness and impact of statistical considerations relating to both data and uncertainty. To illustrate the power and wide range of the statistical

Table 1 ^(a)

	Total number	Number with confirmed polio
Vaccinated	200 745	57
Placebo	201 229	142

^(a) Adapted from Meier

approach, we shall in this concluding section take a brief look at three specific studies.

Example 8: Polio Vaccine Trial To test whether a proposed polio vaccine (the Salk vaccine) was effective, a large scale study was carried out in 1954. The most useful part of the study was a completely randomized trial in which about half of approximately 400 000 school children were assigned at random to receive the vaccine, with the other half receiving a placebo (injection of an ineffective salt solution). The study was double blind, i.e., neither the children nor the physicians making the diagnosis knew to which of the two groups any given child belonged. The results in Table 1 show that vaccination has cut the rate of polio by about 60%. The probability of that marked a decrease under the hypothesis that the vaccine has no effect is about $1/10^7$, much too small to reasonably attribute the effect to pure chance. (Note that this example is of the same type as the hypothetical Example 2 (ESP).)

Example 9: Medical Progress In the preceding example we were concerned with the effectiveness of a single medical innovation – the Salk vaccine. The study reported in the present example (Gilbert *et al.* [30]) addresses a much broader question: What can be said about the effectiveness of present-day medical (or rather more specifically, surgical and anesthetic) innovations as a whole? Such an investigation of the combined implications of a whole area of studies is called a *meta-analysis*. (For an account of meta-analysis with many references to the literature, see Hedges and Olkin [31].)

A first step toward such an analysis is to decide what data to collect, in particular, which studies to include in the investigation. Ideally, one would like to have a list of all relevant studies for the period in question. For medical research, an approximation to this ideal is available in National Library of Medicine’s MEDLARS (MEDical Literature And

Table 2 ^(a)

	Innovation highly preferred	Innovation about standard			Standard highly preferred	Total
		Preferred	Equal	Preferred		
Randomized	5	7	14	6	4	36
Nonrandomized	5	2	3	1	1	12

^(a)© WW Norton & Company, Inc, 1978

Retrieval System), which provides exhaustive coverage of the world's medical literature since 1964. From MEDLARS, the authors obtained a set of 107 papers ("the sample") evaluating the success of specific innovative surgical and anesthetic treatments, and used these to assess the effectiveness of such treatments as a whole. The authors view this set of papers as a sample from the flow of such research studies, and therefore believe that the results should give a realistic idea of what to expect from future innovations, at least for the short term.

The sample contained a total of 48 comparisons of a new treatment with a standard (control). In 36 of the cases, the assignment to treatment and control was randomized, but not in the remaining 12. The results are summarized in Table 2.

A striking feature of the results for the randomized trials is that only 5 out of 36 or about 14% of the innovations were highly preferred. This suggests, the authors point out, that medical science is sufficiently well established so that substantial improvements over the standard treatments are difficult to achieve yet that it is not so settled that only a major theoretical advance will lead to any substantial further improvements.

A more sanguine assessment is obtained from the nonrandomized studies, where 5 out of 12 or about 42% of the innovations were highly preferred. Since randomization provides a surer foundation, the results in the first row may be deemed to be more reliable than those in the second. However, such a judgment overlooks the fact that the comparison of randomized with nonrandomized studies is itself not randomized, i.e., the assignment of studies to these two types was not, and was in fact far from, random. As a result, the two types of studies may not be directly comparable. For example, a surgeon who is strongly convinced of the great superiority of the new treatment, might for ethical or other reasons not be willing to apply the standard treatment, and would therefore have to

look for controls from an earlier period or from other surgeons, while no such qualms might arise for a less dramatic innovation where a randomized study would then be acceptable. Many other explanations could be imagined, and much more information would be needed before one could attempt to assign a cause.

The authors go beyond the frequencies displayed in Table 2 to estimate the size of the treatment effects. For this purpose, they utilize a sophisticated technique, **empirical Bayes**, which is particularly suited for meta-analysis.

Example 10: Literary Detection The Federalist papers are a historically important set of 85 short political essays published mostly anonymously in 1787–1788 by Alexander Hamilton, James Madison, and John Jay. The authorship of most of these papers was eventually established but that of 12 of them has remained in dispute between Madison (M) and Hamilton (H), with a recent historical research leaning toward, but not clearly deciding in favor of, Madison.

An effort to resolve the doubt by statistical methods was undertaken by Mosteller and Wallace [34]. The basic idea of such literary detection is to find aspects of the writing styles of the two (or more) authors in question that have good ability to discriminate between the various possibilities. This was particularly difficult in the present case because the two authors have very similar styles. However, by comparing known texts of M and H, Mosteller and Wallace were able to identify a number of words that were used with much higher frequency by one of the authors than the other. As an illustration, Table 3 shows the frequency of occurrence of the word "on" in blocks of about 200 words from texts by the two authors. (For example, the number of blocks in which the word "on" occurred exactly twice was 27 for the 247 H blocks, but 55 for the 262 M blocks.) The rate of occurrence of the word "on" per thousand words of text was 3.38 for H and 7.57 for M.

Table 3 (a)

	Word frequency of “on” in 200-word blocks							Total number of blocks
	0	1	2	3	4	5	6	
H	145	67	27	7	1	–	–	247
M	63	80	55	32	20	8	4	262

(a) Adapted from Mosteller and Wallace [34]

The present problem is similar to the diagnostic problem of Example 3, with θ taking on the two values H and M. (The statistical methodology dealing with the attribution of an item to one of two or more classes, a patient to a disease, or a piece of writing to an author, is called *classification* or *discrimination*.) Mosteller and Wallace’s first task was to decide on observations X that might provide them with the evidence needed to settle the question. They selected a total of 30 words (including the word “on”) with high discriminating ability, and used as their observations X the frequencies with which these words occurred in a given text of disputed authorship.

Next the distribution $p_\theta(x)$ for the frequency of a given word had to be specified for each author with the help of the known H and M texts. Bayesian analyses for each of the 30 words based on the chosen family of distributions were combined into overall odds for M and H. For all but two of the papers, these overwhelmingly (several hundred million to one) favored Madison, given any reasonable prior odds for M and H. For each of these papers, the analysis leaves little room for doubt. In the remaining two cases Madison is also favored, but the odds are more modest, and the statistical attribution therefore less certain.

Statistical considerations arise in nearly all fields of human endeavor. Many of the associated activities are carried out by large numbers of full- or part-time professional statisticians. The basic training in statistics, as in most other fields, occurs at universities and colleges where departments of statistics offer both undergraduate and graduate degrees in statistics. (At some institutions, the principal statistics courses are instead provided by the mathematics or economics department.) In addition, there may be degrees in biostatistics. Statistics courses and quantitative programs with a strong statistical component may also be offered in operations research, business

schools, demography, economics, education, psychology, and sociology. In addition to courses and programs preparing for a profession in statistics, statistics departments also provide “service courses” both at introductory and advanced levels to students in other fields who may need to use statistical methods in their work.

Another need of statistical education is filled by courses in statistical concepts as a part of general education. Acquiring the ability to think in statistical terms is of great importance, even for persons without a quantitative bent, in view of the pervasive occurrence of statistical ideas in newspapers and magazines, and in the terminology we use to describe and discuss the world around us. What do we mean by saying that women tend to live longer than men, or that cancer patients survive longer today than they did 10 years ago? Does it mean they live longer on the average, that the median length of their life is longer, or that they have a better chance to survive to any given age? And is the longer survival of persons diagnosed to have cancer primarily due to the availability of more effective treatments, including earlier diagnosis? In fact, is there even an increase in length of survival, or is the apparent increase just a statistical consequence of earlier diagnosis? If the disease is diagnosed a year earlier, the survival after diagnosis has increased by a year even if nothing else has changed.

To consider another example, what is meant by the “rate of unemployment”? Roughly speaking, it is the proportion among the people wishing to be employed who are not. But do the official rates include those past seekers for jobs who have given up in despair? A change in the figure for unemployment may be the result of a small change in the definition, or of some sociological change such as the increased entry of women into the workforce.

Many of the considerations involved in such issues are statistical. Fortunately, the basic ideas needed

for general discussion and comprehension can be communicated at a fairly nontechnical level, as is done, for example, in Tanur *et al.* [35], Mosteller *et al.* [33], and Freedman *et al.* [32]. Books such as these, and the increasing availability of first courses in probability and statistics in high schools, should have the effect of gradually raising the general level of statistical literacy.

Acknowledgments

I would like to thank David Cox, Persi Diaconis, David Freedman, William Kruskal, Frederick Mosteller, Juliet Shaffer, and the editors for reading a draft of this article, and providing me with comments and suggestions that resulted in many improvements.

References

- [1] Diaconis, P. (1985). In *Exploring Data Tables, Trends, and Shapes*, D.C. Hoaglin, F. Mosteller and J.W. Tukey, eds, John Wiley & Sons, New York.
- [2] Campbell, W.E. (1942). Monograph No. B-1350, Bell Telephone System Technical Publication.
- [3] Wolfowitz, J. (1943). On the theory of runs with some applications to quality control, *Annals of Mathematical Statistics* **14**, 280–288.
- [4] Nomikos, P. & MacGregor, J.F. (1995). Multivariate SPC charts for monitoring batch processes, *Technometrics* **37**(1), 41–59.
- [5] Nomikos, P. & MacGregor, J.F. (1995). Multiway partial least squares in monitoring batch processes, *Chemometrics and Intelligent Laboratory Systems* **30**, 97–108.
- [6] Huff, D. (1954). *How to Lie with Statistics*, Norton, New York.
- [7] Feynman, R. (1965). *The Character of Physical Law*, The British Broadcasting Corporation, London, England.
- [8] Barnett, V. (1982). *Comparative Statistical Inference*, 2nd Edition, John Wiley & Sons, New York.
- [9] Shortliffe, E.H. (1976). *Computer-Based Medical Consultations, MYCIN*, Elsevier, New York.
- [10] Spiegelhalter, D.J. & Knill-Jones, R.P. (1984). Statistical and knowledge-based approaches to clinical decision-support systems, *Journal of the Royal Statistical Society. Series A* **147**, 35–77.
- [11] Berger, J. (1984). In *Robustness of Bayesian Analysis*, J. Kadane, ed, North-Holland, Amsterdam, The Netherlands.
- [12] Dickey, J.M. (1973). Scientific reporting and personal probabilities: student's hypothesis, *Journal of the Royal Statistical Society. Series B* **35**, 285–305.
- [13] Smith, A.F.M. (1984). Present position and potential developments: some personal views. Bayesian statistics, *Journal of the Royal Statistical Society. Series A* **147**, 245–259.
- [14] Rubin, D.B. (1984). Bayesianly justifiable and relevant frequency calculations for the applied statistician, *Annals of Statistics* **12**, 1151–1172.
- [15] Berger, J. (1985). *Statistical Decision Theory and Bayesian Analysis*, 2nd Edition, Springer, New York.
- [16] Hoaglin, D.C., Light, R., McPeck, B., Mosteller, F. & Stoto, M.A. (1982). *Data for Decisions*, Abt Associates, Cambridge.
- [17] Hoaglin, D.C., Mosteller, F. & Tukey, J.W. (1983). *Understanding Robust and Exploratory Data Analysis*, John Wiley & Sons, New York.
- [18] Hampel, F.R., Ronchetti, E.M., Rousseeuw, P.J. & Stahel, W.A. (1986). In *Robust Statistics*, John Wiley & Sons, New York.
- [19] Alterman, H. (1969). *Counting People*, Harcourt, Brace and World, New York.
- [20] Mosteller, F. & Tukey, J.W. (1977). *Data Analysis and Regression*, Addison-Wesley, Reading, MA.
- [21] Holland, P.W. (1986). Statistics and causal inference, *Journal of the American Statistical Association* **81**(396), 945–960.
- [22] Cook, T.D. & Campbell, D.T. (1979). *Quasi-Experimentation*, Rand McNally, Chicago, IL.
- [23] Cochran, W.G. (1983). *Planning and Analysis of Observational Studies*, John Wiley & Sons, New York.
- [24] Campbell, D.T. (1978). In *Statistics: A Guide to the Unknown*, 2nd Edition, J. Tanur *et al.*, eds, Wadsworth, Belmont, CA.
- [25] Flehinger, B.J. & Louis, T.A. (1972). Sequential treatment allocation in clinical trials, *Biometrika* **58**, 419–426.
- [26] Siegmund, D. (1985). *Sequential Analysis*, Springer, New York.
- [27] Cox, D.R. (1958). *Planning of Experiments*, John Wiley & Sons, New York.
- [28] Box, G.E.P., Hunter, W.G. & Hunter, J.S. (1978). *Statistics for Experimenters*, John Wiley & Sons, New York.
- [29] Box, G.E.P. & Wilson, K.B. (1951). On the experimental attainment of optimum conditions, *Journal of the Royal Statistical Society. Series B* **8**(11), 1–45.
- [30] Gilbert, J.P., McPeck, B. & Mosteller, F. (1977). In *Costs, Risks and Benefits of Surgery*, J.P. Bunker, B.A. Barnes, F. Mosteller, eds, Oxford University Press, New York.
- [31] Hedges, L.V. & Olkin, I. (1985). *Statistical Methods for Meta-Analysis*, Academic, Orlando.
- [32] Freedman, D., Pisani, R. & Purves, R. (1978). *Statistics*, Norton, New York.
- [33] Mosteller, F., Kruskal, W.H., Link, R.F., Pieters, R.S. & Rising, G.R. (eds) (1973). *Statistics by Example*. Addison-Wesley, Reading, MA, 4 Vols.
- [34] Mosteller, F. & Wallace, D.L. (1984). *Applied Bayesian and Classical Inference*, Springer, New York.
- [35] Tanur, J., Mosteller, F., Kruskal, W.H., Link, R.F., Pieters, R.S., Rising, G.R. & Lehmann, E.L. (eds) (1978). *Statistics: A Guide to the Unknown*, 2nd ed, Wordsworth, Belmont, CA.

Further Reading

- Healy, M.J.R. (1973). The varieties of statistician, *Journal of the Royal Statistical Society. Series A* **136**, 71–74.
- Hoaglin, D.C., Mosteller, F. & Tukey, J.W. (1985). *Exploring Data Tables, Trends, and Shapes*, John Wiley & Sons, New York.
- Meier, P. (1978). In *Statistics: A Guide to the Unknown*, 2nd Edition, J. Tanur *et al.*, eds, Wordsworth, Belmont.

Moser, C.A. (1973). *Journal of the Royal Statistical Society. Series A* **136**, 75–88.

ERICH L. LEHMANN

Article originally published in Encyclopedia of Statistical Sciences, 2nd Edition (2005, John Wiley & Sons, Inc.). Minor revisions for this publication by Jeroen de Mast.

Variables

The notion of recording some feature that is subject to systematic and haphazard variation, and hence the idea of a variable, is all pervasive. In mathematical statistics, *random variable* has a technical meaning, defined in the context of measure theory: random variables are measurable functions from a probability space to a measurable space (see **Probability Theory**). In applied statistics and empirical studies, a variable is a measurable factor, characteristic, or attribute of a product, process, machine, or other experimental unit. When considering products of a production process as the experimental units of interest, one could measure variables such as the products' length, strength, and whether the products are free from scratches.

When a classification of types of variables is attempted, a rather bewildering variety soon emerges. There is the further complication that terminology sometimes varies considerably between different fields, for example between the natural and the social sciences and, within the latter, between psychometrics, sociology, and econometrics.

In some ways the most important classification of variables is *via* their purpose in the particular investigation under study. These distinctions are inevitably context dependent. Variables may be *explanatory*, *intermediate*, or *responses*. That is, the objective of the study is regarded as being the explanation of variability in the *response* in terms of the *explanatory* variables, and *intermediate* variables may be regarded as responses in one stage of the analysis and as explanatory in another stage of analysis. In econometrics the terms *exogenous* and *endogenous* are used, respectively, for variables that are regarded as given from outside and those whose variation is to be explained or modeled. These classes correspond broadly to the explanatory *versus* the others. The older terminology of *independent* instead of explanatory and *dependent* instead of response is probably best abandoned because of potential confusion with statistical independence (see **Dependence**). See also **Covariate**.

Another classification of variables is based on the mathematical structure of the possible values. Thus variables can be one dimensional (univariate) or essentially multidimensional (multivariate),

and forcing essentially multivariate concepts, for example service quality, into one dimension could possibly lead to dire consequences. In one dimension variables may be binary (two possible values), nominal (more than two unordered values), ordinal (ordered values, but without the concept of "distance" among the values), quantitative and essentially integer valued, or quantitative and essentially with a continuous range of possible values. Mixtures of the various types occur quite commonly and can lead to problems of formulation, especially when multidimensional distributions are required. Pass/fail inspection is an example of binary measurement. Classification of production faults into defect types (machine related, material related, operator related, etc.) gives a nominal variable. Classification of faults into ordered classes (severe, major, minor) gives an ordinal variable. Measuring the products' dimensions, strength, or another characteristic gives a quantitative variable. Binary, nominal, and ordinal variables are often called *categorical*, while quantitative variables are often called *numeric*. In industrial statistics and engineering, the terms *attribute* and *variable* are often used instead of categorical and numeric, but this usage of these words deviates from the meaning that the words attribute and variable usually have in the other statistical disciplines.

It is possible to pass down the hierarchy by collapsing, in the extreme case converting, a quantitative measurement into a binary one according to whether the quantitative variable exceeds or does not exceed a given threshold. Sometimes it is helpful either in interpretation or model selection to move in another direction, for example interpreting a binary or ordinal variable *via* a latent continuous variable.

Note that the notion of a continuous range of possible values is often an extremely useful idealization, but in fact all variables as actually recorded are discrete and in some contexts this discreteness must be recognized both in theoretical analyses and in applications. Campbell [1] gave a careful discussion of principles of measurement in the physical sciences. Stevens [2] introduced a classification of quantitative variables, according to differences and/or ratios of corresponding scale points that retain their meaning across the range of values (so-called interval and ratio scales).

References

- [1] Campbell, N.R. (1920). *Physics: The Elements*, Cambridge University Press.
- [2] Stevens, S.S. (1950). *Handbook of Experimental Psychology*, John Wiley & Sons, New York. Chapter 1.

Article originally published in Encyclopedia of Statistical Sciences, 2nd Edition (2005, John Wiley & Sons, Inc.). Minor revisions for this publication by Jeroen de Mast.

DAVID. R. COX

Variance Components

Although there is a vast amount of literature dealing with **variance components**, the basic concept is simple. Suppose an observable random variable Y , with variance σ_Y^2 , is the sum $A + E$ of two nonobservable independent variables A and E . Let σ_A^2 denote the variance of A and σ_E^2 denote the variance of E . Then, due to the independence of A and E , $\sigma_Y^2 = \sigma_A^2 + \sigma_E^2$. σ_A^2 and σ_E^2 are called *variance components* of σ_Y^2 .

In a variety of practical problem areas, it is important to make statistical inference statements relative to variance components. As early as 1861, the astronomer Airy [1] concerned himself with telescopic observations of the same phenomenon b_i times for the i th night, for $i = 1, 2, \dots, a$ nights. Airy desired information pertinent to the nightly variation σ_A^2 and the within-night variation σ_E^2 . The essence of Airy's problem is also the essence of problems familiar to many investigators. For illustrative purposes, we consider an example involving quality measurements in a production process. Let Y_{ij} be the quality of the j th product in batch i .

It is helpful to think of the observations Y_{ij} as variables created by adding realizations of variables associated with sampling from two hypothetical populations; one is a population of batch effects and the other a population of nonbatch-related individual products effects. Crucial to the **estimation** procedure is the fact that one realization of an A_i variable is added to several (in this case b) different values of E_{ij} variables. If one E value was added to each A value, then the effects would be confounded and the components σ_A^2 and σ_E^2 would be nonestimable. To describe this situation statistically, we assume the random-effects model $Y_{ij} = \mu + A_i + E_{ij}$, where all variables $A_i, i = 1, \dots, a$, and $E_{ij}, j = 1, \dots, b$, are mutually independent,

$\mathcal{E}[A_i] = 0, \mathcal{E}[E_{ij}] = 0, \text{var}(A_i) = \sigma_A^2$, and $\text{var}(E_{ij}) = \sigma_E^2$. In this setting A_i is the random effect produced by the i th batch and σ_A^2 is the variance component attributable to variation between batches. $E_{ij} = Y_{ij} - \mu - A_i$ is the residual effect or the random individual within-batch product effect and σ_E^2 is the variance component due to variation between products within batches. The problem is to make inference statements relative to σ_A^2 and σ_E^2 . We first focus on point estimates of the components.

To quote Searle [2], who has studied variance components problems in detail, "The basic principle for estimating variance components has been and to a large extent still is, that of equating quadratic functions of the observations to their expected values. Obvious candidates for such functions are those of the analysis of variance (AOV) table."

Suppose $a = 6$ batches are chosen at random from a very large number of batches. Consider the data given in Table 1.

The one-way classification *analysis of variance* (see **Analysis of Variance**) is shown in Table 2.

The estimation procedure is to set numerical values of the mean squares (MSs) equal to expected mean squares (EMS) and solve the resulting equations

$$\hat{\sigma}_E^2 + 5\hat{\sigma}_A^2 = 2.936 \quad (1)$$

$$\hat{\sigma}_E^2 = 0.153 \quad (2)$$

Table 1 Quality of 30 products

Batch	1	2	3	4	5	6
Quality of	4.17	4.83	5.95	4.43	3.86	4.98
individual	5.21	5.70	6.50	4.85	4.12	5.58
products	4.60	5.91	6.34	3.92	4.70	6.10
	4.35	4.92	6.09	4.25	3.90	5.75
	4.54	5.72	5.87	4.64	4.41	5.29

Table 2 AOV for quality data

Source	$df^{(a)}$	$SS^{(b)}$	$MS^{(c)}$	$EMS^{(d)}$
Total corrected	$ab - 1 = 29$	$SST = 18.34$		
Among batches	$a - 1 = 5$	$SSA = 14.68$	$MSA = 2.936$	$\sigma_E^2 + 5\sigma_A^2$
Within batches	$a(b - 1) = 24$	$SSE = 3.66$	$MSE = 0.153$	σ_E^2

^(a)df: degrees of degree ^(b)SS: Sum of Squares ^(c)MS: Mean Squares ^(d)EMS: Expand Value of mean square

2 Variance Components

We find $\hat{\sigma}_A^2 = 0.557$ and thus conclude that the component due to batches are quite large relative to the within-batches component.

Unbiasedness is among the desirable properties that AOV point estimators of variance components enjoy. Among their undesirable properties is the fact that estimates can be negative and indeed often are. Furthermore, no closed form exists for the distribution of estimators such as $\hat{\sigma}_A^2 = (MSA - MSE)/b$. The support for the distribution is all real numbers. AOV estimators are linear combinations of χ^2 related variables.

If the inference mode is testing instead of point estimation (see **Parametric Tests**), the AOV along with normality assumptions give MSs, which are distributed independently as multiples of central χ^2 variables. In particular, for the balanced one-way situation, which we have thus far focused upon,

$$\frac{SSA}{\sigma_E^2 + b\sigma_A^2} \sim \chi^2(a-1) \text{ independent of} \quad (3)$$

$$\frac{SSE}{\sigma_E^2} \sim \chi^2[a(b-1)] \quad (4)$$

Thus under the hypothesis $H_0: \sigma_A^2 = 0$, the ratio $MSA/MSE \sim F\{(a-1), a(b-1)\}$ and serves as a test statistic for testing H_0 versus $\sigma_A^2 > 0$. For the quality data, the F ratio is $F = 2.936/0.153 = 19.25$.

If the inference mode is that of a **confidence interval** (see **Confidence Intervals**), then the procedure varies, for a multitude of suggestions have been presented. The reason for this is that for most components only appropriate interval statements can be derived. An exception is the exact interval statement

$$P \left[\frac{SSE}{\chi_{1-\alpha/2}^2} < \sigma_E^2 < \frac{SSE}{\chi_{\alpha/2}^2} \right] = 1 - \alpha \quad (5)$$

This statement is an immediate consequence of the fact that under normality, SSE/σ_E^2 is distributed as a central χ^2 variable.

An exact interval can be obtained for $\sigma_E^2/(\sigma_A^2 + \sigma_E^2)$ from the fact that

$$\frac{\sigma_E^2}{\sigma_A^2 + \sigma_E^2} \frac{MSA}{MSE} \sim F\{a-1, a(b-1)\} \quad (6)$$

Manipulation of the statement for $\sigma_E^2/(\sigma_A^2 + \sigma_E^2)$ leads directly to probability interval statements for σ_A^2/σ_E^2 and $\sigma_A^2/(\sigma_A^2 + \sigma_E^2)$, the last function being known as the **intraclass correlation**. An overview of methods for the determination of confidence intervals for variance components is given in Burdick and Graybill [3].

The concepts presented for the balanced one-way classification situation can be extended in many ways. Consider three variance components that arise in conjunction with a random-effects model for balanced two-way **factorial** data. Let factor A be studied at randomly chosen levels and factor B at b levels. Denote the observation for the i th level of A and the j th level of B by Y_{ij} and suppose $Y_{ij} = \mu + A_i + B_j + E_{ij}$. The parameters of interest are the variances σ_A^2, σ_B^2 , and σ_E^2 for the respective populations from which A_i, B_j , and E_{ij} were selected. Provided we have a full $a \times b$ factorial experiment (see **Factorial Experiments**) (no missing observations), the AOV procedure is to let

$$MSA = \hat{\sigma}_E^2 + b\hat{\sigma}_A^2, MSB = \hat{\sigma}_E^2 + a\hat{\sigma}_B^2 \quad (7)$$

and $MSE = \hat{\sigma}_E^2$, from which we obtain

$$\hat{\sigma}_A^2 = \frac{1}{b}[MSA - MSE] \quad (8)$$

$$\hat{\sigma}_B^2 = \frac{1}{a}[MSB - MSE] \quad (9)$$

and $\hat{\sigma}_E^2 = MSE$.

Under normality assumptions, the F ratios MSA/MSE and MSB/MSE serve as test statistics for the hypotheses $\sigma_A^2 = 0$ and $\sigma_B^2 = 0$.

Extensions of the basic ideas are not always straightforward. There are complicating factors when:

1. the data are unbalanced,
2. the model is a mixed effects model, and when
3. interactions are present.

One complicating factor is that of determining EMS. Another is choosing among the often many possible quadratic forms to equate with their expectations. Henderson [4] presented three methods, which have been widely used. Method I uses quadratic forms (not necessarily sums of squares (SS)) analogous to those already discussed for the balanced data cases. Method I unfortunately gives biased

estimates of variance components in mixed-model situations. Although this is well known, the ease of the method has made it popular and for mixed-model data two approaches are commonly employed. The first is to ignore the fixed effects and the second is to treat the fixed effects as random effects. Both approaches have been investigated and deemed unsatisfactory.

Henderson's method II was designed to overcome the **bias** problem inherent in method I when applied to mixed models. The idea is to correct the data for the mixed effects and then use method I on the corrected data. The correction process produces correlated error terms, which then lead to a biased estimate of σ_E^2 , but this bias can also be corrected.

Henderson's method of correcting the data for fixed effects is but one of several ways that can and have been applied. This lack of uniqueness detracts from the usefulness of the method. There are other difficulties. If there are interactions between any of the fixed effects with a random effect, then one cannot transform to a model where method I can be directly applied. This is a serious limitation.

Henderson's method III (the fitting constants method) is often employable when methods I and II would be unsatisfactory. The method can be used for any situation. We will illustrate its use in an unbalanced two-way cross-classification situation with interaction. Suppose there are N observations. Let the model be $Y_{ijk} = \mu + \beta_j + A_i + (A\beta)_{ij} + E_{ijk}$. Factor B is fixed but factor A is random as are interaction and error. Suppose s cells contain observations. We desire to estimate the variance components σ_A^2 , σ_{AB}^2 , and σ_E^2 . Seeing there are three components, the method consists of setting three reductions in SS equal to their expectations. The method is not unique. The problem remains unsolved as to which quadratic forms should be used. One choice is to write an AOV for "A after B" (see Table 3).

Table 3 An AOV for two-way unbalanced data

Source	df	Reductions
Levels of B	$b - 1$	$R(\beta \mu)$
A adjusted for B	$a - 1$	$R(A \mu, \beta)$
Interaction adjusted	$s - a - b + 1$	$R(A \times B \mu, A, \beta)$
Error	$N - s$	SSE

$\hat{\sigma}_E^2$ is always estimated from $SSE/(n - s)$ but other choices are

$$R(A \times B|\mu, A, \beta) = C_1 \hat{\sigma}_{A\beta}^2 + (s - a - b - 1) \hat{\sigma}_E^2 \quad (10)$$

and

$$R(A, A \times B|\mu, \beta) = C_2 \hat{\sigma}_{A\beta}^2 + C_3 \hat{\sigma}_A^2 + (s - b) \hat{\sigma}_E^2 \quad (11)$$

Until 1967, Henderson's methods were the only methods used with unbalanced data. Since that year many alternative approaches have been presented. Among these are:

- ML: maximum likelihood.
- REML: restricted maximum likelihood.
- MINQUE: minimum norm quadratic unbiased estimation.

REML is an adaptation of ML (see **Maximum Likelihood**) where the idea is to maximize the likelihood of that part of the sufficient statistics, which is location invariant. It is well known that the ML estimate of the variance in a simple random sample of size b is biased, having divisor b rather than $b - 1$. In more complex models the resulting estimates of variance may be entirely unsatisfactory especially if the number of nuisance parameters is large. The problem is tackled by the REML procedure, which effectively adjusts for **degrees of freedom** (df) lost in estimating parameters in the expected value, and gives estimators of the remaining parameters with less bias and better consistency properties. The idea is to divide the data into two sets of contrasts – one between group means, and another with zero expectation (the error contrasts). The estimates are found maximizing the likelihood function of the error contrasts, excluding the first set of contrasts on the grounds that such contrasts cannot provide information about the variance components. See Searle *et al.* [5] for details.

References

- [1] Airy, G.B. (1861). *On the Algebraical and Numerical Theory of Errors of Observations and the Combination of Observations*, Macmillan Publishing, London.

4 Variance Components

- [2] Searle, S.R. (1971). *Linear Models*, John Wiley & Sons, New York.
- [3] Burdick, R.K. & Graybill, F.A. (1992). *Confidence Intervals on Variance Components*, Marcel Dekker, New York.
- [4] Henderson, C.R. (1953). Estimation of variance and covariance components, *Biometrics* **9**, 226–252.
- [5] Searle, S.R., Casella, G. & McCulloch, C.E. (1992). *Variance Components*, John Wiley & Sons, New York.

ROBERT HULTQUIST

Article originally published in Encyclopedia of Statistical Sciences, 2nd Edition (2005, John Wiley & Sons, Inc.). Minor revisions for this publication by Jeroen de Mast.

Control Charts, Overview

History

Control charts originated in the work of Dr. Walter Shewhart, the father of statistical process control, in the 1920s. Shewhart, a physicist by training, was employed by the Western Electric Company, which made early telephone equipment. He was concerned with finding an economical way to monitor and control manufacturing processes in order to yield high quality product. Shewhart knew that external factors beyond the control of production line workers sometimes precipitated process changes that necessitated corrective action. However, he also realized that it required discipline to avoid making frequent unnecessary adjustments that might do more harm than good.

Shewhart's groundbreaking insight, which was quite revolutionary at the time, was to apply statistical concepts to solve this very practical problem. He reasoned that the characteristic variability of a manufacturing process could best be represented in statistical terms. Process changes that required corrective action could then be detected by observing results that were inconsistent with these established statistical norms. The tool Shewhart devised for conducting this comparison was the control chart. He first communicated the concept behind control charts internally in 1924, and subsequently published his approach in various papers and books throughout the 1920s and 1930s (e.g. [1]).

Methodology

Working in an era of extremely rudimentary computational resources, and very limited education among production workers, Shewhart's techniques necessarily had to be simple ones. Consequently, so-called *Shewhart Charts* can be constructed from process data using only basic arithmetic, and tables of statistical constants which Shewhart (and subsequently other researchers) provided. The method was traditionally implemented graphically, with line workers plotting production data over time. The scale of the characteristic being monitored would be on the vertical axis, and the time sequence of the observations on

the horizontal axis (Figure 1). Horizontal lines ("control limits") superimposed on the plot would indicate the typical range of the data; corrective action would be taken only when the plotted values fell outside of these limits (later, other supplementary response rules were added, as well).

The sources of routine variability which are viewed as being intrinsic to a process have come to be called *common causes*. Root causes of atypical variability external to the proper operation of a process have been named *special causes*. A process that is operating properly at a stable level, exhibiting consistent variability in the absence of special causes, is said to be in "statistical control".

Viewed in these terms, the use of Shewhart control charts is intended to let common cause variability pass, but to detect and signal the need for corrective action (called *giving an out-of-control signal*) when special causes are present.

Evolution

Starting with Shewhart, the most common use of control charts throughout their history has been to monitor the level of a product or process property which is measured on a continuous numerical scale, often assumed to be normally distributed (*see Control Charts for the Mean*). Over the intervening decades, however, a vast literature has originated, which extends the uses of control charts to many other domains. Included among these are, for example, control charts for monitoring:

- the level of a continuous characteristic having an arbitrary distribution (*see Control Charts, Nonparametric*);
- the level of a process which is run in short or batch production runs (*see Control Charts for Short Production Runs; Control Charts for Batch Processes*);
- the variability of a process (*see Moving Range and R Charts; Control Charts for the Standard Deviation*);
- process characteristics measured using categorical (also known as *discrete* or *attribute*) data (*see Control Charts for Attributes*);
- processes which by their nature drift over time (*see Autocorrelated Data; Large Autoregressive Integrated Moving Average (ARIMA) Modeling*);

2 Control Charts, Overview

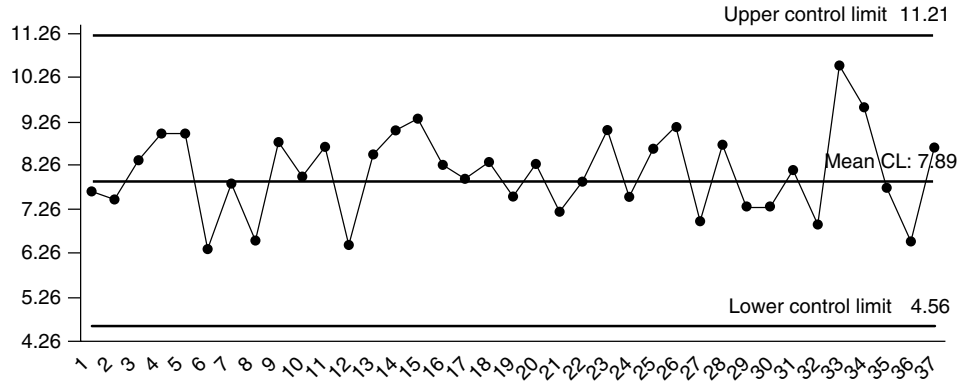


Figure 1 Typical Shewhart control chart for the level of a process

- the level or variability of multiple correlated (“multivariate”) product/process characteristics (*see* **Multivariate Control Charts Overview; Hotelling’s T^2 Chart; Multivariate Exponentially Weighted Moving Average (MEWMA) Control Chart; Multivariate Cumulative Sum (CUSUM) Chart; Multivariate Charts for Variability; Multivariate Control Charts, Interpretation of**);
- and even for monitoring the in-principle infinite-dimensional (e.g., spatial) profile of a property present throughout a material or manufactured part (*see* **Profile Monitoring**).

Many of these techniques make intrinsic use of the extraordinary growth of computational resources that has occurred since World War II. One artifact of this fundamental change in the production environment has been the rise of control charts which make use of computationally intensive techniques (*see*, for example, **Neural Networks: Construction and Evaluation**). Another has been a migration from the traditional manual graphical implementation, to an automated algorithmic approach. Moreover, these monitoring schemes have for some time now been applied to systems of processes, or even on a plant-wide scale [2].

On another dimension, statistical process control techniques have, for a number of years now, found rapidly increasing application outside of their tradition manufacturing origins. The control charting methods mentioned above are seeing use for monitoring and improving nearly all genres of business processes, including, for example, supply

chain and order fulfillment, human resource management, financial and accounting processes, call center management, and many others.

Application Issues

One very useful aspect of research into control charting is that many of the newer techniques (*see* **Exponentially Weighted Moving Average (EWMA) Control Chart; Cumulative Sum (CUSUM) Chart**) are considerably more powerful than those pioneered by Shewhart. As a result, it is now possible to detect process excursions more quickly than before.

It is this speed of detection – the average number of observations (“average run length” (ARL))” after onset of a special cause until it is detected – that is most commonly used as a metric for determining the power and sophistication of a control charting method (*see* **Average Run Lengths and Operating Characteristic Curves**). A common fallacy surrounding the use of control charts is that an out-of-control signal is given immediately at the point in time at which a special cause has occurred. Although this may be true in the case of extremely severe disruptions, control charts usually require some time to detect the process change attributable to a special cause. As one might logically expect, larger changes will generally be detected more quickly, while more subtle effects will take longer to produce an out-of-control signal. Thus the ARL is a useful measure of how promptly one might expect notification of a change, and is a good way of comparing the performance of alternative charting methods in a given situation.

When it is essential to identify when the change actually took place (for example, to aid in determining the identity of the special cause), control charts may be used in conjunction with statistical change-point methods (*see* **Change-Point Methods**).

Another consideration is that all control charts have some likelihood of signaling even when no special cause is present. With properly constructed charts, however, the frequency of such false alarms is generally quite small (typically, once in several hundred observations at most). As a practical matter, judicious use of common control chart types that have been well documented in the literature will allow selection of chart parameters that avoids undue risk of false alarms.

The increasing complexity of control charts also poses a number of challenges for users, however. One example of this is that signals from some chart types are more difficult to interpret than was the case with traditional Shewhart methods (*see*, for example, **Multivariate Control Charts, Interpretation of**). Moreover, the fact that there is now such a wealth of control charting techniques, which overall is a great benefit, nonetheless, makes the selection of which chart type to use, more difficult (*see* **Control Charts, Selection of**). Another consideration is that the power of modern control charting techniques is often such as to enable detection of process changes that are so small as to be of little practical interest.

Related Techniques

Research into process control theory and methods since Shewhart's time has also identified fundamentally different strategies for process control which complement or offer alternatives to control charting. Best known among these are algorithmic feedback or feed-forward techniques which adjust the process in real time with the goal of minimizing variability and thereby maximizing quality (*see* **Engineering Process Control; Feedback Control**). Considerable effort has gone into characterizing what types of process are best controlled using these techniques, or by using control charts, or a combination of both.

Conclusion

On the basis of Shewhart's seminal insight, control charts have evolved from simple beginnings into a

multifaceted process control discipline. Control charting techniques have been developed for monitoring properties of the overwhelming majority of processes encountered in common practice. As a result, statistical process control and related fields such as engineering process control provide business and industry with powerful and easy to use tools for optimizing performance and improving the quality of all kinds of business processes.

References

- [1] Shewhart, W. (1931). *Economic Control of Quality of Manufactured Product* D. Van Nostrand Company, Republished by Quality Press, Milwaukee, 1980.
- [2] Faltin, F. & Tucker, T. (1991). On-line quality control for the factory of the 90's and beyond, in *Chapter in Statistics and Design in Process Control: Keeping Pace with Automated Manufacturing*, J. Keats, ed Marcel Dekker, New York.

Related Articles

Autocorrelated Data; Average Run Lengths and Operating Characteristic Curves; Change-Point Methods; Control Charts for Attributes; Control Charts for Batch Processes; Control Charts for Short Production Runs; Control Charts for the Mean; Control Charts for the Standard Deviation; Control Charts, Nonparametric; Control Charts, Selection of; Cumulative Sum (CUSUM) Chart; Economic Design of Control Charts; Engineering Process Control; Exponentially Weighted Moving Average (EWMA) Control Chart; Feedback Control; Hotelling's T^2 Chart; Large Autoregressive Integrated Moving Average (ARIMA) Modeling; Moving Range and R Charts; Multivariate Charts for Variability; Multivariate Control Charts Overview; Multivariate Control Charts, Interpretation of; Multivariate Cumulative Sum (CUSUM) Chart; Multivariate Exponentially Weighted Moving Average (MEWMA) Control Chart; Neural Networks: Construction and Evaluation; Precontrol; Profile Monitoring; Quality Management, Overview; Regression Control Charts; Six Sigma; Statistical Quality Control; Total Quality Management (TQM); Variable Sampling Rate Control Charts.

FREDERICK W. FALTIN

Control Charts for the Mean

Introduction with Example

The survey results given in [1] showed that the $\bar{X}$ chart was the most frequently used **control chart**. Although those survey results are of course now out of date, the $\bar{X}$ chart is almost certainly still the most frequently used control chart.

As the name of the chart implies, this is a control chart of averages, and more specifically of subgroup averages that are time ordered. As with every Shewhart control chart for which 3σ limits are used, the control limits are of the general form $\hat{\theta} \pm 3\hat{\sigma}_{\hat{\theta}}$, with $\hat{\theta}$ denoting a function of the statistic that is plotted on the chart, which serves as the centerline on the chart, and $\hat{\sigma}_{\hat{\theta}}$ denotes the estimator of the standard deviation of $\hat{\theta}$.

The control limits for the $\bar{X}$ chart are thus computed as $\bar{\bar{X}} \pm 3\hat{\sigma}_{\bar{X}}$, with $\bar{\bar{X}}$ denoting the average of the subgroup averages $\bar{X}$, obtained ideally from recent data, and $\hat{\sigma}_{\bar{X}}$ denotes some estimator of the standard deviation of the subgroup averages. That estimator is typically either $(\bar{s}/c_4)/\sqrt{n}$, with $\bar{s}$ denoting the average of the subgroup standard deviations and c_4 the tabular value for a given subgroup size, or $(\bar{R}/d_2)/\sqrt{n}$, with $\bar{R}$ the average of the subgroup ranges and d_2 the tabular value for a given subgroup size. (The values of c_4 and d_2 are those that make the respective estimators unbiased estimators of σ_x . Some values of c_4 are given here in Table 1.)

To illustrate, using subgroup standard deviations in estimating σ_x , let us assume that 40 subgroups of size 5 are available from recent data, with $\bar{\bar{X}} = 36.79$, $\bar{s} = 3.3853$, and $c_4 = 0.9400$ from Table 1. The control limits would then be computed as $\bar{\bar{X}} \pm 3(\bar{s}/c_4)/\sqrt{n} = 36.79 \pm 3(3.3853/0.94)/\sqrt{5} = (31.96, 41.62)$. Thus, the lower control limit (LCL) is 31.96 and the upper control limit (UCL) is 41.62, with the midline on the chart being $36.79 (= \bar{\bar{X}})$.

The chart is shown in Figure 1 and it can be observed that all of the points are within the control limits. This suggests that the process was in control in Phase 1.

Note: Control charts are sometimes constructed using assumed or desired parameter values. This is

highly inadvisable because an out-of-control signal could simply mean that a process was out of control relative to desired parameter values that might not be attainable under current conditions. We want to see whether a process is stable or out of control relative to its current capability.

If any of the 40 subgroup averages had plotted outside the control limits, a search for an **assignable cause** would be performed in an attempt to determine the cause(s) of the aberrant point(s). If the cause(s) could be both identified and removed, the control limits would be recomputed and the process continued until either all of the plotted points are inside the control limits or causes of points outside the limits cannot be determined, or if determinable causes, are not removable.

The end of this process would be the end of Phase I, to be followed by real-time process monitoring, which is Phase II.

In order for the control limits that are determined at the end of Phase I to have good properties (especially properties that deviate only slightly from the properties when parameter values are assumed known), and not have sampling variability that is deemed unacceptably high, the limits should be computed from a reasonably large number of observations from the assumed in-control distribution (see, e.g., [2]). In this example, there were 200 observations, which is less than preferred. Therefore, it would be better, while ostensibly operating in Phase II, to add about 20 more subgroups of size 5 to the Phase I data after the data from the 20 subgroups have been collected in real time. The “semipermanent” control limits would then be recomputed using the 60 subgroups of size 5 if the Phase II monitoring showed the process to apparently be in control when the new points were plotted.

Table 1 Selected values of c_4

n	c_4
3	0.8862
4	0.9213
5	0.9400
6	0.9515
7	0.9594
8	0.9650
9	0.9693
10	0.9727

2 Control Charts for the Mean

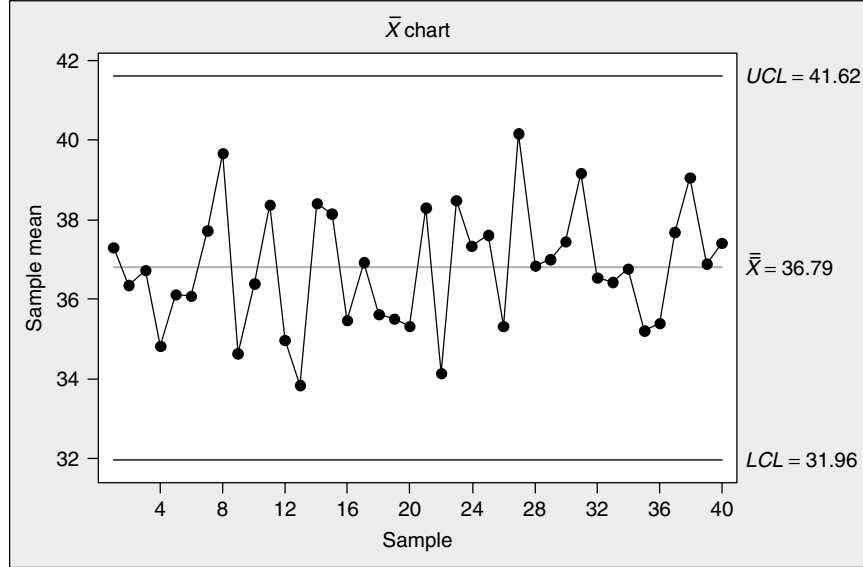


Figure 1 $\bar{X}$ -chart

Of course eventually the control limits will have to be recomputed as processes change over time and control limits should be based on recent data. When new control limits should be computed is a matter of judgment.

Since subgroups are involved, the subgroup size must be determined. Historically, subgroups of four or five observations have frequently been used. The results of a survey of 117 companies in the food manufacturing/processing and metalworking industries were reported in [3], with the respondents asked to provide the subgroup size that they typically used. Unfortunately, the author, Osborn, did not provide any summary statistics regarding subgroup size, but instead gave the average **power** (i.e., probability) of detecting, with the first plotted point, a (minimum) shift that the respondent believed was crucial. If we assume normality and known parameter values, it can be shown that the power of detecting a $1\sigma_x$ shift (either increase or decrease) on the first plotted point after the shift with a subgroup size of 5 is 0.22, and 0.45 for a subgroup size of 6. (Note that the shift is here stated in terms of σ_x , not $\sigma_{\bar{x}}$. If the shift were stated in terms of $\sigma_{\bar{x}}$, the power would be a function only of the amount of the shift and be independent of the subgroup size.) Since the average power reported by Osborn was 0.3426, the average subgroup size was

probably about 5, assuming that the shift was $1\sigma_x$ and was constant over the respondents.

A large subgroup size should be avoided (just as very large sample sizes should be avoided with any tool of statistical inference), as this would make the control chart too sensitive to very small mean shifts that are considered to be inconsequential. Indeed, Box in [4] stated, in responding to the question of whether statistical control can be achieved: "The truth is that we all live in a nonstationary world; a world in which external factors never stay still. Indeed the idea of stationarity – of a stable world in which, without our intervention, things stay put over time – is a purely conceptual one." If we accept this, then we should strive for processes being in a state of near statistical control and not be concerned by what might seem to be very small mean shifts.

As with any statistical procedure for detecting change, the user would have to determine the minimal shift that is important to detect as quickly as possible, and then decide if the selected subgroup size provides sufficient power for detecting that amount of shift. (Although the $\bar{X}$ chart is reasonably good at detecting large shifts, there are other process control procedures, such as cumulative sum (**CUSUM**) procedures and **exponentially weighted moving average** (**EWMA**) procedures that are better for detecting small mean shifts.)

Assumptions

With every statistical procedure, it is important to recognize the assumptions that underlie the procedure and to test those assumptions, not leaving anything to chance. For an $\bar{X}$ chart, the basic assumptions are that the data are approximately normally distributed and are independent, both within subgroups and between subgroups.

The normality assumption might be questioned by some, but if the distribution of the data is highly skewed and/or the population distribution is believed to be highly skewed based on the nature of the process characteristic, there will also be considerable **skewness** in the distribution of subgroup averages computed from five observations. In that case, it would be inappropriate to use control limits that are symmetric about $\bar{\bar{X}}$.

Nonnormality can be assessed through a normal **probability plot** of the data or of the subgroup averages if there is a sufficient number of them, and numerical tests for normality are also available. If the collected data should be approximately normal but instead are markedly nonnormal, this might be a signal that the process is out of control. In that case, bringing the process into a state of statistical control may remedy the problem. If a process is in control and the data are markedly nonnormal, it might be possible to transform the data to approximate normality, such as by using a Box–Cox transformation (see [5]). For more discussion of the normality issue for subgroup data, see [6] and [7].

The assumption of independence is also important since a lack of independence will, strictly speaking, invalidate the control limits computed under the assumption of independence. Lack of independence can be ascertained through a simple time sequence plot of the data and/or by constructing an **autocorrelation** plot.

As with nonnormality, a lack of independence (i.e., autocorrelated data) could signal that a process is out of control. A classic example of this was given in [8], which was originally published in *Industrial Quality Control* in 1949. Tooling data exhibited a very strong trend, almost like a sine curve. The control chart of subgroup averages indicated control, but many of the individual values did not even meet engineering specifications. The trend was not discovered until the individual observations were plotted. (It will sometimes be necessary to construct a

plot of individual observations even though a control chart for averages is being used.)

If autocorrelation is a normal part of a process, the control limits should be adjusted appropriately, as, for example, in [9]. One often-suggested remedy is to sample less frequently, thus removing or at least reducing the autocorrelation. This is not a good idea as tight control of a process could then be lost. It is best to maintain the normal sampling frequency.

The Possible Use of $k\sigma$ Limits

As with every other type of control chart, there is the option of using $k\sigma$ limits rather than 3σ limits. 3σ limits have historically been used because they offer a reasonable balance between false alarms and quick detection of a parameter change. In some fields of application, quick detection may be far more important than false alarms. For example, in [10] the author mused that 2σ limits might sometimes be more appropriate in clinical care applications.

Individuals Charts

It is not always possible or practical to form subgroups, so individual observations are often plotted instead of subgroup averages. When this is the case, the control limits are obtained from $\bar{X} \pm 3\hat{\sigma}_x$ and various labels have been used for the chart, including an individuals chart (*I* chart) and an *X* chart. Without subgrouping, there is no obvious way to measure short-term variability, however. The common way to overcome this problem is to compute “moving subgroups” of size two. (There will be $n - 1$ such subgroups for n observations.) The range is computed for each subgroup and $\hat{\sigma}_x = \overline{MR}/1.128$, with $\overline{MR}$ denoting the average of these (moving) ranges.

This is appropriate for Phase I. For the start of Phase II, it is better, as shown in [11], to use $\hat{\sigma}_x = \bar{s}/c_4$, with $\bar{s}$ as previously defined and c_4 the tabular value corresponding to the number of observations. (The data used to compute the estimate would be the data not discarded in Phase I.) Specifically, this estimator has a smaller variance than the estimator computed from the moving ranges. It is better to use the latter for Phase I, however, since that estimator is less influenced by extreme data points, thus giving

the chart greater power for detecting points that correspond to an out-of-control condition. Thus, different estimators should ideally be used in each Phase.

It is especially important that the assumption of (approximate) normality be checked when plotting individual observations, as well as checking the assumption of independence. When these assumptions are not met, the chart can have undesirable properties.

Software Usage

The control limits can be computed most efficiently using computer software, and much commercial software is available for control chart construction. Once the limits have been computed, however, points could easily be hand plotted on a chart, if desired, although a calculator would be necessary to compute each subgroup average. This obviously is not the easiest way to do it, however.

Summary

Two charts for controlling a process mean have been described in this article: an $\bar{X}$ chart and an X chart. For a mean shift of a given size, the $\bar{X}$ chart has the greater power of detection but it is not always possible to form subgroups. The normality assumption is especially important for the X chart but should also be checked for the $\bar{X}$ chart. The independence assumption is also important for both charts.

References

- [1] Saniga, E. & Shirland, L. (1977). Quality control in practice: A survey, *Quality Progress* **10**, 30–33.
- [2] Quesenberry, C.P. (1993). The effect of sample size on estimated limits for $\bar{X}$ and X control charts, *Journal of Quality Technology* **25**, 237–247.
- [3] Osborn, D.P. (1990). Statistical power and sample size for control charts – survey results and implications, *Production and Inventory Management Journal* **31**, 49–54.
- [4] Box, G.E.P. (1990). Must we randomize our experiment? *Quality Engineering* **2**, 497–502.
- [5] Box, G.E.P. & Cox, D.R. (1964). An analysis of transformations (with discussion), *Journal of the Royal Statistical Society. Series B* **26**, 211–246.
- [6] Moore, P.G. (1957). Normality in quality control charts, *Applied Statistics* **6**, 171–179.
- [7] Ryan, T.P. (2000). *Statistical Methods for Quality Improvement*, 2nd ed., John Wiley & Sons, New York.
- [8] McCoun, V.E. (1974). The case of the perjured control chart, *Quality Progress* **7**(10), 17–19 (Originally published in *Industrial Quality Control*, May 1949).
- [9] Wardell, D.G., Moscovitz, H. & Plante, R.D. (1994). Run length distributions of special-cause control charts for correlated processes, *Technometrics* **36**, 3–17.
- [10] Carey, R.G. (2003). *Improving Healthcare with Control Charts: Basic and Advanced SPC Methods and Case Studies*, ASQ Quality Press, Milwaukee.
- [11] Cryer, J.D. & Ryan, T.P. (1990). Estimation of sigma for an X -chart: $\overline{MR}/d_2$ or s/c_4 ? *Journal of Quality Technology* **22**, 187–192.

THOMAS P. RYAN

Exponentially Weighted Moving Average (EWMA) Control Chart

Introduction

The exponentially weighted moving average (EWMA) chart first appeared in the statistical **quality control** literature approximately 50 years ago. Roberts [1] introduced the procedure, which he called the *geometric moving average*, as an alternative to the standard Shewhart **control chart** with zone tests. He concluded that the EWMA was more effective than the standard control chart at detecting small changes in the process mean (*see Control Charts for the Mean*). The EWMA has since been studied extensively as a process monitoring tool used to detect both process mean and process variance shifts (*see Control Charts for the Standard Deviation*). More recently (see Vander Wiel [2], for example), the EWMA has appeared in the statistical quality control literature as a bridge between statistical process monitoring and engineering feedback control (*see Engineering Process Control*). That is, it has been presented as a tool that can be appropriately used for either the detection of process changes, or to estimate the mean of an autocorrelated process for the purpose of feedback control. In this article, we will review both applications of the EWMA, and will discuss strengths and weaknesses associated with applying the EWMA in either case. We will also briefly discuss some of the many modifications and enhancements of the EWMA that have appeared in the literature.

Process Monitoring

Detection of Mean Shifts

Most of the statistical quality control literature regarding the EWMA has focused on the problem of detecting a shift in the process mean. It is typically assumed that the process mean μ shifts from a nominal mean μ_0 to a new process average $\mu = \mu_0 + \Delta$, where Δ represents the amount of shift. In the absence of any process upsets, $\Delta = 0$. The purpose of the control chart is to detect the mean shift,

identify the **assignable cause**, and return the process average to nominal.

The EWMA is defined, for a series of observations $y_1, y_2, \dots, y_t, \dots$, by

$$E_0 = \mu_0 \quad (1)$$

$$E_t = (1 - \lambda)E_{t-1} + \lambda y_t, \quad 0 < \lambda \leq 1,$$

$$t = 1, 2, 3, \dots \quad (2)$$

The EWMA has the property that the most recent observation receives the greatest amount of weight, and all previous observations receive weight decreasing exponentially from the most recent to the first observation. The y_t 's can be individual observations or, in many cases, sample averages. If each of the y_t 's has mean μ_0 , E_t also has mean μ_0 . When the process mean shifts to $\mu = \mu_0 + \Delta$, the mean of E_t shifts away from μ_0 and the process is considered "out of control". If the y_t 's are independent with common standard deviation σ , it is easy to show [1] that the standard deviation of E_t is

$$\sigma_{E_t} = \sigma \times \sqrt{\frac{\lambda}{2 - \lambda} [1 - (1 - \lambda)^{2t}]} \quad (3)$$

with limiting value

$$\sigma_E = \sigma \times \sqrt{\frac{\lambda}{2 - \lambda}} \quad (4)$$

The control chart limits are usually defined in terms of this limiting value, set at

$$\mu_0 \pm k\sigma_E \quad (5)$$

The control chart is specified by choice of λ and k . The performance of the EWMA is usually measured in terms of the average number of observations until the control chart produces a signal. This metric is referred to as the average run length (ARL, *see Average Run Lengths and Operating Characteristic Curves*). The ARL depends on the particular choice of λ and k and the shift Δ in the process mean. Various ways of evaluating the ARL and identifying the best choices of λ and k have been proposed in the literature.

Roberts [1] evaluated the ARL of the EWMA by simulating a series of independent observations with a shift in the mean introduced into the process. Crowder [3] showed that the ARL, as well as other

2 Exponentially Weighted Moving Average (EWMA) Control Chart

moments of the run length, could be represented as an integral equation, solvable with simple numerical methods. A computer program to solve the integral equation appeared in [4]. Lucas and Saccucci [5] showed that the ARL, as well as other moments of the run length, could be represented as a Markov chain, also solvable with simple numerical methods. A computer program to solve the **Markov chain** appeared in [5]. These papers assumed a normal process model for the observations; however, the results can be extended to other distributions.

Given the ARL performance of competing EWMA schemes, the optimal chart can be easily determined. Crowder [3] presented a two-step approach to choosing an optimal (λ, k) combination. The first step involved choosing λ to minimize the ARL for a given process shift of interest ($\Delta \neq 0$). The second step involved choosing k to give the desired ARL for the nominal process mean ($\Delta = 0$). Lucas and Saccucci [5] outlined a similar design procedure also based on ARLs.

The properties of the EWMA for the process mean and optimal EWMA design strategies have been thoroughly investigated by numerous authors. The conclusions have consistently been that smaller values of λ are optimal for detecting small shifts in the process mean, and that the EWMA outperforms the standard Shewhart chart for that special case. For very large shifts in the process mean, the Shewhart chart is optimal.

Detection of Variance Shifts

Although most applications of the EWMA have concentrated on the process mean, the more important problem for quality control practitioners is the problem of detecting and removing causes of increased process variation, which typically have greater impact on product quality. Recently, much more attention has been paid to this important problem.

Crowder and Hamilton [6] proposed an EWMA to monitor a process standard deviation based on the observations $y_t = \ln(s_t^2)$, where the s_t^2 values are successive sample variances, each based on subgroup size n . The log transform was recommended for variances of normally distributed data, as the logarithms will be much closer to normally distributed than the sample variances themselves. The EWMA was defined as

$$E_t = \max \left\{ (1 - \lambda)E_{t-1} + \lambda y_t, \ln(\sigma_0^2) \right\}, \quad 0 < \lambda \leq 1,$$

$$t = 1, 2, 3, \dots \quad (6)$$

with $E_0 = \ln(\sigma_0^2)$. The one-sided control chart limit was set at $k\sigma_E$, where σ_E is the limiting standard deviation of E_t , approximated by

$$\sigma_E \approx \left(\frac{\lambda}{2 - \lambda} \right) \left(\frac{2}{n - 1} + \frac{2}{(n - 1)^2} + \frac{4}{3(n - 1)^3} - \frac{16}{15(n - 1)^5} \right) \quad (7)$$

Note from equation (6) that the EWMA is reset to $\ln(\sigma_0^2)$ if its value ever falls below $\ln(\sigma_0^2)$. This is owing to the fact that the primary interest is in increases in process variability. An integral equation was developed to allow easy computation of the ARL for the one-sided EWMA, and a design strategy using those ARLs was presented. A computer program to compute the ARLs appears in [7]. The authors concluded that small λ values are optimal for detecting small percent increases in process standard deviation, while large values of λ are optimal for detecting large percent increases. The authors also concluded that the EWMA outperforms the standard Shewhart range chart for detecting small increases in the process standard deviation. These conclusions were not surprising, as the log transformation tends to “normalize” the data and turn the problem into one of detecting shifts in a process mean.

MacGregor and Harris [8] provided a brief overview of other uses of the EWMA to monitor process variation. They also proposed using “exponentially weighted moving variance” and “exponentially weighted mean squared deviation” charts to monitor process variation. They concluded that the charts are particularly useful when only individual observations are collected and when observations are positively correlated. Gan [9] investigated the properties of “omnibus” EWMA charts designed to detect shifts in both the process mean and the variance with a single chart, but found them to be ineffective. He concluded that the best approach is to design the EWMA mean chart and the one-sided EWMA variance charts separately, as described above.

Strengths and Weaknesses

The major strengths of the EWMA include its simplicity and its ability to quickly detect small shifts in the process mean or variance. These conclusions are reached for idealized models with independent,

identically distributed, normal observations. It is also typically assumed that the parameters of the **normal distribution** are known and are constant until some process upset causes a shift in value.

Numerous modifications to the EWMA chart have been proposed to address failures in the assumptions and specific weaknesses in the approach. One failure in the assumptions is that the process parameters may not be known and must be estimated from the data. Jones and Champ [10] investigated the performance of the EWMA chart with estimated parameters and concluded that **estimation** of the parameters can lead to substantially more frequent false alarms and reduce the sensitivity of the EWMA to detect changes in the process mean. The authors show how the ARLs are affected by estimation and what sample sizes are necessary to apply the standard EWMA procedures without modification.

A weakness of the EWMA involves its inability to quickly detect both small and large shifts in the process mean with a single chart. Lucas and Saccucci [5] proposed a combined Shewhart-EWMA chart, which is constructed by adding Shewhart limits to an EWMA control chart. They also table ARLs of the proposed combined chart and give design recommendations. Capizzi and Masarotto [11] proposed an adaptive EWMA to address the same problem in a single chart. The weighting constant λ is replaced by an adaptive weighting function. The proposed weighting function is constructed to respond to either small or large process shifts. The authors conclude that the adaptive EWMA provides reasonably balanced protection against shifts of different sizes.

Another weakness of the EWMA has been called the *"inertia" problem*. The inertia problem may occur if the process mean is off target when the EWMA is initiated. This could happen because of startup problems or because of ineffective control actions taken after the previous process alarm. When the value of λ is small, it may take the EWMA more time than desirable to react to the process being off target. Lucas and Saccucci [5] proposed a "fast initial response" (FIR) feature to address this problem. The FIR feature is obtained by implementing two one-sided EWMA's, each with an appropriate head start. Instead of initializing a single EWMA with the nominal mean, one of the one-sided EWMA's has a starting value above the nominal mean and the other one-sided EWMA has a starting value below the nominal mean. If the process is off target at startup,

one of the EWMA's will give a signal more quickly than a single two-sided EWMA.

The traditional use of the EWMA in the statistical quality control literature has been to detect shifts in either the mean or the variance of an independent, identically distributed normal process. In the sections that follow, we will address the use of the EWMA for monitoring autocorrelated processes, and how the EWMA can appropriately be used as a feedback control tool.

Feedback Control and Monitoring

Many manufacturing processes have continuous flows of material that produce correlations over consecutive measurements. In these cases, past data are useful for predicting future observations and it is natural to use those predictions for active feedback regulation (*see Engineering Process Control; Feedback Control*). In this context, EWMA's can usefully play three distinct roles: forecasting, compensation, and monitoring. These are discussed in the sections titled "EWMA's for Prediction", "EWMA's for Feedback Control", and "EWMA's for Monitoring Controlled Processes".

Brief Overview of Literature

Kalman [12] revolutionized prediction and control of dynamic systems by introducing a filter to calculate minimum variance predictions through sequential updating. The EWMA update is a special case of the Kalman filter and is the optimal one-step-ahead forecast for an integrated moving average (IMA) process. Trigg [13] used EWMA's to monitor the performance of a forecasting system. His monitoring signal is the ratio of an EWMA of forecast errors to an EWMA of their absolute values. Bather [14] and Box and Jenkins [15] made an interesting connection between EWMA charts and optimal feedback control. They showed that incorporation of a fixed cost for any nonzero adjustment to an IMA process results in an optimal regulation policy that adjusts only when an EWMA statistic falls outside of certain upper and lower dead-band limits. A follow-up paper of Box *et al.* [16] contrasts the function of control charts for monitoring with that of control algorithms for process regulation.

The discussion paper of Box and Kramer [17] explores interrelationships between feedback control

4 Exponentially Weighted Moving Average (EWMA) Control Chart

and process monitoring and the use of EWMA's in both contexts. Vander Wiel *et al.* [18] call the combination of control and monitoring *algorithmic statistical process control* (ASPC). They use EWMA's to regulate viscosity in a batch polymerization reaction and then apply a cumulative sum (CUSUMs) as a monitoring scheme (see **Cumulative Sum (CUSUM) Chart**). Box and Luceño's book [19] includes derivations and examples of EWMA's for optimal prediction and control and for process monitoring.

EWMA's for Prediction

Exponential smoothing has been widely used for forecasting **time series** [20]. Single exponential smoothing uses an EWMA of the time series to predict the next value. This minimizes the variance of forecast errors, if the time series is a so-called IMA and the EWMA parameter matches the time series model, as shown below.

Industrial processes that tend to wander are sometimes modeled as random walks. The IMA is a slightly more general model that is equivalent to noisy observation of a random walk. The usual description of an IMA process, N_t (N stands for noise) is

$$N_t = N_{t-1} + a_t - \rho a_{t-1} \quad (t = 1, 2, \dots) \quad (8)$$

where $-1 < \rho < 1$, the a_t are independent "innovations" with mean zero and constant standard deviation, and the recursion starts with $N_0 = a_0$. If $\rho = 0$, N_t is a random walk. (See **Large Autoregressive Integrated Moving Average (ARIMA) Modeling** for additional methods of time series modeling.)

Taking an EWMA of the process N_t for the purpose of forecasting is known as *single exponential smoothing*:

$$\hat{N}_{t+1|t} = (1 - \lambda)\hat{N}_{t|t-1} + \lambda N_t \quad (9)$$

If the EWMA parameter λ is set equal to the IMA parameter ρ , then a little algebra provides the forecast error

$$e_{t+1} \equiv N_{t+1} - \hat{N}_{t+1|t} = a_{t+1} \quad (10)$$

which has the smallest possible variance of any forecast based on $(N_0, \dots, N_t)$. That is, for an IMA process, the best one-step-ahead forecast is an EWMA of the past observations.

EWMA's for Feedback Control

Box and Kramer [17] note that the IMA model can be thought of as an interpolation between an uncorrelated noise process and random walk. They say it "occupies a central place in representing the behavior of many nonstationary disturbances" that impact industrial processes. If N_t represents the "noise" affecting a linear system, then the forecast $\hat{N}_{t|t-1}$ is natural to use for feedback control. A general linear system with IMA noise is

$$Y_{t+1} = \alpha + \beta X_t + f(X_{t-1}, X_{t-2}, \dots) + N_{t+1} \quad (11)$$

where Y_{t+1} is an observation of the system at time $t + 1$, X_t is the setting of a control variable and $f(X_{t-1}, X_{t-2}, \dots)$ is the sum of the linear effects of previous control settings. To steer the system toward a target T , we set

$$T = \alpha + \beta X_t + f(X_{t-1}, X_{t-2}, \dots) + \hat{N}_{t+1|t} \quad (12)$$

and solve for X_t . Here, $\hat{N}_{t+1|t}$ is defined by equation (9) with N_t determined from equation (11). The resulting feedback adjustment rule is

$$X_t = -\beta^{-1} \left[\alpha + f(X_{t-1}, X_{t-2}, \dots) + \hat{N}_{t+1|t} - T \right] \quad (13)$$

and plugging this rule into equation (11) shows that the regulated process becomes

$$Y_t = T + a_t \quad (14)$$

No other regulation scheme gives as small a mean squared error (MSE) between the process and target value. That is, the feedback rule equation (13) based on an EWMA of the process noise is the minimum MSE controller. The above discussion assumes knowledge of the parameters $\lambda = \rho$, α , β , and $f(\cdot)$. Luceño [21] discusses estimation of the best λ for feedback control in cases where ρ is not known or the noise process is not an IMA.

EWMA's for Monitoring Controlled Processes

The previous sections showed how EWMA's can be used for prediction and regulation of a process that tends to wander. However, a sequence of forecasts or a regulated process also needs to be monitored – and

various control charts, including EWMA charts, have been suggested for this role.

Vander Wiel [2] evaluates the use of an EWMA control chart of forecast errors, $e_t = N_t - \hat{N}_{t|t-1}$ to detect a step change in an IMA process and provides guidance on choosing the EWMA control chart parameter. Montgomery *et al.* [22] compare various control charts, including EWMA, applied to forecast errors for detecting upsets that occur within an autoregressive moving average process.

Various authors recommend applying control charts to the *control inputs* X_t or the *controlled outputs* Y_t or both series when a process is run under feedback control. Runger [23] discusses this issue and concludes that, rather than using either inputs or outputs, a more sensitive monitoring procedure is to generate the one-step-ahead forecast errors, e_t , and use a control chart, such as an EWMA, that is tuned to detect process shifts of specified magnitudes.

Strengths and Weaknesses

EWMA is versatile. They can be used for prediction and feedback regulation as well as detection of sudden changes in autocorrelated and regulated processes. A strength of EWMA for prediction and regulation is their simplicity. They are easy to understand and they can produce good forecasts (or form the basis for good feedback regulation) on a wide range of processes with positive **autocorrelation**. On the other hand, EWMA is only optimal for processes with IMA noise structure. For other autocorrelation structures better forecasts can be achieved using time series modeling (*see Large Autoregressive Integrated Moving Average (ARIMA) Modeling*). Similarly, feedback control based on EWMA can be suboptimal. Sometimes better results are achieved using more general control algorithms such as proportional integral derivative (PID) controllers or optimal controllers based on time series modeling (*see Engineering Process Control; Feedback Control*).

Use of EWMA to monitor forecasts or regulated processes is similar to the conventional use of EWMA for process monitoring; however, determining performance properties such as ARLs is more complicated. If forecasts are not optimal, then the forecast errors being monitored will be autocorrelated and this can have a large impact on ARL properties of EWMA control charts. Furthermore, feedback control will act on an abrupt shift in the level of a

process and bring the resulting system output back to the target level. Often, if a shift is not detected within a short “window of opportunity” it will be missed altogether. This complicates the analysis of control chart performance. See [2] for an example.

Finally, using EWMA for forecasting, regulation, and monitoring of forecasts and regulated processes requires a more sophisticated approach to process improvement than conventional control charts. These applications combine time series modeling and parameter estimation, automatic control theory, and control chart design. Many organizations are not equipped with engineers or statisticians that are proficient in each of these areas.

Summary

An EWMA follows the ups and downs of process data; however, it is much smoother than the sequence of raw measurements because each new data point gets averaged in with the previous EWMA value. EWMA control charts are good at detecting small shifts in a process mean or variance and can be combined with Shewhart charts to enhance the ability to quickly detect large process shifts. They can also be initialized with head starts to produce a signal more quickly if the process mean is off target when the EWMA is initiated.

EWMA is also used for forecasting autocorrelated processes, for feedback regulation and for monitoring forecast errors and regulated processes. These advanced uses of EWMA draw on technologies from control charts, time series modeling, and automatic control.

References

- [1] Roberts, S.W. (1959). Control chart tests based on geometric moving averages, *Technometrics* **1**(3), 239–250.
- [2] Vander Wiel, S.A. (1996). Monitoring processes that wander using integrated moving average models, *Technometrics* **38**(2), 139–151.
- [3] Crowder, S.V. (1989). Design of exponentially weighted moving average schemes, *Journal of Quality Technology* **21**(2), 155–162.
- [4] Crowder, S.V. (1987). A simple method for studying run-length distributions of exponentially weighted moving average charts, *Technometrics* **29**(4), 401–407.
- [5] Lucas, J.M. & Saccucci, M.S. (1990). Exponentially weighted moving average control schemes: properties and enhancements, *Technometrics* **32**(1), 1–12.

- [6] Crowder, S.V. & Hamilton, M.D. (1992). An EWMA for monitoring a process standard deviation, *Journal of Quality Technology* **24**(1), 12–21.
- [7] Hamilton, M.D. & Crowder, S.V. (1992). Average run lengths of EWMA control charts for monitoring a process standard deviation, *Journal of Quality Technology* **24**(1), 44–50.
- [8] MacGregor, J.F. & Harris, T.J. (1993). The exponentially weighted moving variance, *Journal of Quality Technology* **25**(2), 106–118.
- [9] Gan, F.F. (1995). Joint monitoring of process mean and variance using exponentially weighted moving average control charts, *Technometrics* **37**(4), 446–453.
- [10] Jones, L.A. & Champ, C.W. (2001). The performance of exponentially weighted moving average charts with estimated parameters, *Technometrics* **43**(2), 156–167.
- [11] Capizzi, G. & Masarotto, G. (2003). An adaptive exponentially weighted moving average control chart, *Technometrics* **45**(3), 199–207.
- [12] Kalman, R.E. (1960). A new approach to linear filtering and prediction problems, *Journal of Basic Engineering* **82**(1), 35–45.
- [13] Trigg, D.W. (1964). Monitoring a forecasting system, *Operational Research Quarterly* **15**(3), 271–274.
- [14] Bather, J.A. (1963). Control charts and the minimization of costs, *Journal of the Royal Statistical Society. Series B (Methodological)* **25**(1), 49–80.
- [15] Box, G.E.P. & Jenkins, G.M. (1963). Further contributions to adaptive quality control: simultaneous estimation of dynamics: nonzero costs, *Bulletin of the International Statistical Institute* **34**, 943–974.
- [16] Box, G.E.P., Jenkins, G.M. & MacGregor, J.F. (1974). Some recent advances in forecasting and control, *Applied Statistics* **23**(2), 158–179.
- [17] Box, G. & Kramer, T. (1992). Statistical process monitoring and feedback adjustment: a discussion, *Technometrics* **34**(3), 251–267.
- [18] Vander Wiel, S.A., Tucker, W.T., Faltin, F.W. & Doganaksoy, N. (1992). Algorithmic statistical process control: concepts and an application, *Technometrics* **34**(3), 286–297.
- [19] Box, G.E.P. & Luceno, A. (1997). *Statistical Control by Monitoring and Feedback Adjustment*, John Wiley & Sons, New York.
- [20] Brown, R.G. (1962). *Smoothing, Forecasting and Prediction of Discrete Time Series*, Prentice-Hall, Englewood Cliffs.
- [21] Luceno, A. (1995). Choosing the EWMA parameter in engineering process control, *Journal of Quality Technology* **27**(2), 162–168.
- [22] Montgomery, D.C., Keats, J.B., Runger, G.C. & Messina, W.S. (1994). Integrating statistical process control and engineering process control, *Journal of Quality Technology* **26**(2), 79–87.
- [23] Runger, G.C. (2002). Assignable causes and autocorrelation: control charts for observations or residuals? *Journal of Quality Technology* **34**(2), 165–170.

Related Articles

Autocorrelated Data; Average Run Lengths and Operating Characteristic Curves; Change-Point Methods; Control Charts for the Mean; Control Charts for the Standard Deviation; Control Charts, Overview; Engineering Process Control; Feedback Control; Large Autoregressive Integrated Moving Average (ARIMA) Modeling; Multivariate Exponentially Weighted Moving Average (MEWMA) Control Chart; Profile Monitoring.

STEPHEN V. CROWDER AND SCOTT
A. VANDER WIEL

Cumulative Sum (CUSUM) Chart

Introduction

In any manufacturing process, variations in the quality of products always exist and they can never be completely eliminated. In other words, no two products can be completely identical in all aspects. These variations are either inherent in a process or are due to some assignable causes. The inherent variation is due to the cumulative effects of many random causes peculiar to a process. No single cause can account for any major part of the inherent variation. If a process is properly optimized using some experimental design procedures (*see Factorial Experiments; Fractional Factorial Designs; Response Surface Methodology; Assessment of Experimental Designs*), the inherent variation can be kept small. The variations due to assignable causes are relatively larger and these causes include changes in process settings, changes in raw materials, or human errors. A process in the presence of assignable causes is considered to be out of control. A control chart (*see Control Charts, Overview; Control Charts for the Mean*) is a charting procedure designed to detect the presence of assignable causes. One will have a better understanding of the cumulative sum (CUSUM) chart if it is compared to the Shewhart chart.

Consider a process in which there is an in-control measurement of a certain quality characteristic and any deviation from this value is not desirable. In the Shewhart control charting procedure, samples are taken at regular intervals and the sample mean of measurements is plotted against the sample number. A signal is issued for any sample mean falling beyond certain chart limits. The sensitivity of the Shewhart chart is determined by the change in the proportion of sample means falling outside these limits, which in turn depends on the shift in the process mean. This change will be small if the shift is small and thus the Shewhart chart is not sensitive in detecting small shifts. For a chart to be sensitive in detecting a small shift, past samples as well as the most recent sample will have to be utilized in signaling. The CUSUM chart introduced by Page [1] is one such chart. The CUSUM chart can also be used to provide an estimate

of the process mean although this estimate is a poor one.

Let $Y_1, Y_2, \dots$ be the measurements of a certain quality characteristic. These measurements are assumed to be independent and identically distributed normal random variables with mean μ and variance σ^2 (*see Normal Distribution*). Also assume that the in-control mean and variance are given as μ_0 and σ_0 , respectively. The original measurements can be standardized as $X_t = (Y_t - \mu_0)/\sigma_0, t = 1, 2, \dots$ which will have a standard normal distribution. The upper-sided and lower-sided CUSUM charts are obtained by plotting

$$S_t = \max\{0, S_{t-1} + X_t - k\}, \quad t = 1, 2, \dots \quad (1)$$

and

$$T_t = \min\{0, T_{t-1} + X_t + k\}, \quad t = 1, 2, \dots \quad (2)$$

against the sample number t , respectively. The chart parameter k is a positive constant and the initial values $S_0 = u, 0 \leq u < h$ and $T_0 = v, -h < v \leq 0$ are usually chosen to be 0. The upper-sided CUSUM chart, which is intended for detecting an upward shift in the mean will issue a signal when $S_t > h$. Similarly, the lower-sided CUSUM chart which is intended for detecting a downward shift in the mean will issue a signal when $T_t < -h$. The quantities h and $-h$ are usually known as the *chart limits*.

Sensitivity of CUSUM Chart

The sensitivity of a CUSUM chart in detecting a small shift is best illustrated using a simulated data set. In Figures 1–3, we consider a data set in which the first 50 measurements were simulated from a normal distribution with mean 0 and variance 1, and the last 10 with mean 1 and variance 1. The chart parameter k is set at 1 and chart limits of the CUSUM charts are selected such that jointly they have the same in-control average run length (ARL, *see Average Run Lengths and Operating Characteristic Curves*) as the Shewhart chart with 3σ limits. For the first 50 measurements, all the charts show no unusual pattern which corresponds to an in-control process. The shift in the mean occurs after the 50th measurement and Figure 1 shows that the Shewhart chart is unable to detect the presence of such a shift. This shift is clearly reflected as a

2 Cumulative Sum (CUSUM) Chart

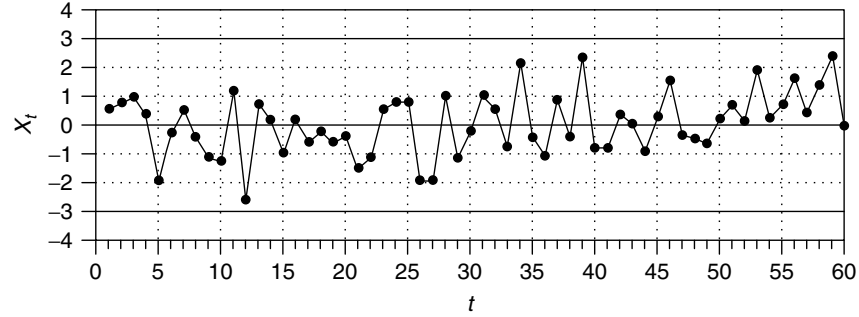


Figure 1 Shewhart chart

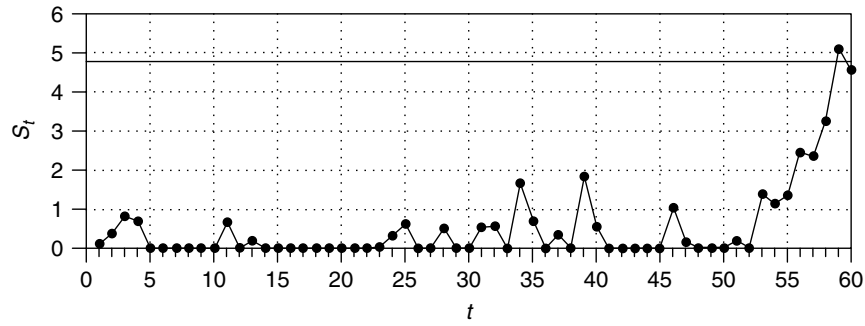


Figure 2 Upper-sided CUSUM chart

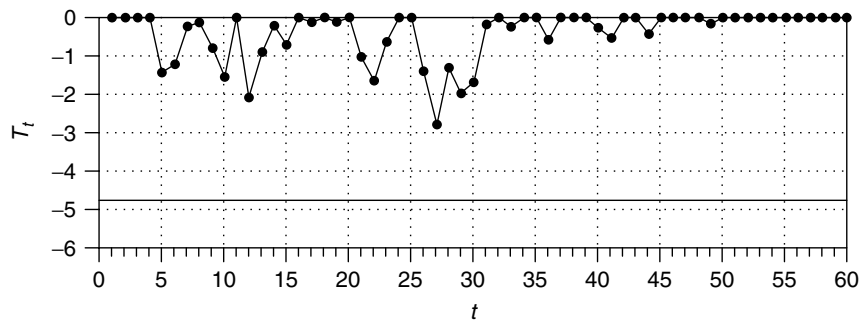


Figure 3 Lower-sided CUSUM chart

slope on the upper-sided CUSUM chart after the 51st sample and a signal is issued at the 59th sample. The lower-sided CUSUM chart shows complete inactivity after the 50th measurement. This example clearly demonstrates the sensitivity of the CUSUM chart in detecting a small shift.

Unfortunately, it is this strength that makes the CUSUM chart less sensitive in detecting a big

shift. This fact is revealed in Figures 4 and 5. In this data set, the first 50 measurements were simulated from a normal distribution with mean 0 and variance 1, and the last four with mean 2.5 and variance 1. The presence of such a big shift is immediately detected by the Shewhart chart and it signals at the 51st sample. This shift is also detected by the CUSUM chart, but it issues a signal only at

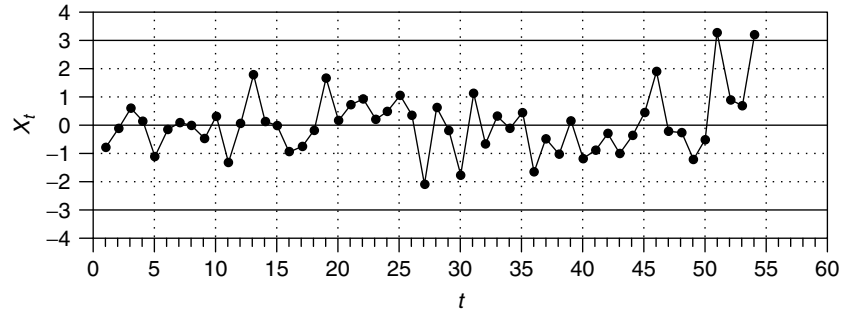


Figure 4 Shewhart chart

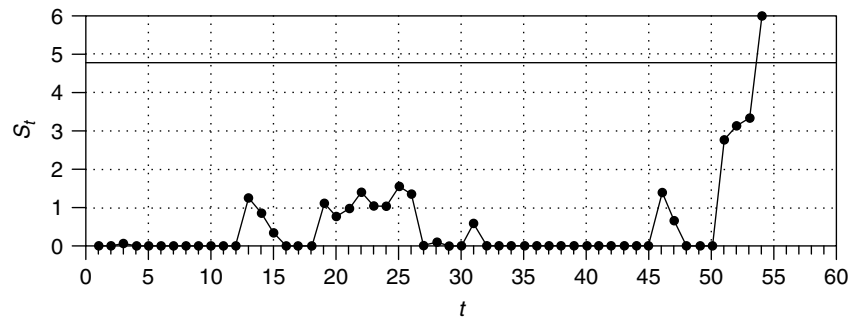


Figure 5 Upper-sided CUSUM chart

the 54th measurement because it only accumulates $X_t - k$ each time. The worst scenario for the CUSUM chart happens when a shift occurs and the CUSUM is at zero. If sampling is done once every hour, three samples slower in signaling would translate into 3 h of nonconforming products.

Other Aspects of CUSUM Chart

An alternative implementation of the CUSUM chart is the use of the V-mask introduced by Barnard [2], but mathematically both implementations are identical in the sense that when one signals, the other one will signal too. The V-mask implementation is not popular because points plotted could be in an increasing or decreasing trend for a long period (which means plotting out of a piece of paper) and yet does not trigger any signal. The CUSUM chart is often used to estimate the process mean (*see Estimation*) after a signal is issued and this can be done using $\hat{\mu} = k + S_t/N_t$ [3, p. 394 or p. 27; 4] where N_t is the number of samples since the chart became last

active. This formula is actually the same as the mean of the sample means since the chart became last active. As this estimator is associated with a signal, it is biased for all values of the process mean except for one value where it is unbiased. The ARL of a CUSUM chart can be studied using the integral equation approach [1] or the Markov chain approach [5]. In terms of weighting of samples for signaling, the CUSUM chart is perhaps the most intelligent when compared to the exponentially weighted moving average (EWMA, *see Exponentially Weighted Moving Average (EWMA); Exponentially Weighted Moving Average (EWMA) Control Chart*) and the straight moving average (SMA, *see Moving Averages*) charts. Hunter [6] showed that while the EWMA and SMA charts always weight the past samples in a deterministic fashion, the CUSUM chart will only include those when it first senses ($X_t > k$) that the process could be out of control. When it is convinced ($S_t = 0$) that the process is in control, it discards all past samples and starts a new.

When a process is first started and one suspects that the process might not start in an in-control state, Lucas and Crosier [7] suggested using S_0 and T_0 that are nonzero for a faster detection. Combined Shewhart–CUSUM scheme [8] and simultaneous CUSUM schemes [9] could be used to improve the performance of the CUSUM chart. The CUSUM chart is similar to the sequential probability ratio test (SPRT) Moustakides [10] proved that the CUSUM chart is optimal in its run-length performance. Gan [11] provided a simple procedure for the design of an optimal CUSUM chart. A more comprehensive treatment of the CUSUM charts can be found in Hawkins and Olwell [4]. A multivariate CUSUM scheme (see **Multivariate Control Charts Overview–Multivariate Control Charts, Interpretation of**) will have to be employed in processes where the quality is characterized by multiple quality measurements. One approach is to implement a chart using a single monitoring statistic summarized from the multivariate data in each sample. Another approach is to compute a CUSUM statistic for each quality variable. These statistics can either be monitored individually or be combined into a single monitoring statistic on a control chart. Woodall and Ncube [12], Crosier [13], Healy [14], Lowry *et al.* [15], and Mason *et al.* [16] have more detailed discussions of various multivariate charting procedures.

References

- [1] Page, E.S. (1954). Continuous inspection schemes, *Biometrika* **41**, 100–115.
- [2] Barnard, G.A. (1959). Control charts and stochastic processes, *Journal of the Royal Statistical Society* **B21**, 239–271.
- [3] Montgomery, D.C. (2005). *Introduction to Statistical Quality Control*, 5th Edition, John Wiley & Sons, New York, p. 394.
- [4] Hawkins, D.M. & Olwell, D.H. (1997). *Cumulative Sum Charts and Charting for Quality Improvement*, Springer-Verlag, New York.
- [5] Brook, D. & Evans, D.A. (1972). An approach to the probability distribution of CUSUM run length, *Biometrika* **59**, 539–550.
- [6] Hunter, S.J. (1986). The exponentially weighted moving average, *Journal of Quality Technology* **18**, 203–210.
- [7] Lucas, J.M. & Crosier, R.B. (1982). Fast initial response for CUSUM quality-control schemes: give your CUSUM a head start, *Technometrics* **24**, 199–205.
- [8] Lucas, J.M. (1982). Combined Shewhart-CUSUM quality-control schemes, *Journal of Quality Technology* **14**, 51–59.
- [9] Sparks, R.S. (2000). CUSUM charts for signaling varying location shifts, *Journal of Quality Technology* **32**, 157–171.
- [10] Mustakides, G. V. (1986). Optimal Stopping Times for Detecting Changes in Distributions, *The Annals of Statistics* **14**, 1379–1387.
- [11] Gan, F.F. (1991). An optimal design of CUSUM quality-control charts, *Journal of Quality Technology* **23**, 279–286.
- [12] Woodall, W.H. & Ncube, M.M. (1985). Multivariate CUSUM quality-control procedures, *Technometrics* **27**, 285–292.
- [13] Crosier, R.B. (1988). Multivariate generalizations of cumulative sum quality-control schemes, *Technometrics* **30**, 291–303.
- [14] Healy, J.D. (1987). A note on multivariate CUSUM procedures, *Technometrics* **29**, 409–412.
- [15] Lowry, C.A., Woodall, W.H., Champ, C.W. & Rigdon, S.E. (1992). A multivariate exponentially weighted moving average control chart, *Technometrics* **34**, 46–53.
- [16] Mason, R.L., Champ, C.W., Tracy, N.D., Wierda, S.J. & Young, J.C. (1997). Assessment of multivariate process control techniques, *Journal of Quality Technology* **29**, 140–143.

Related Articles

Average Run Lengths and Operating Characteristic Curves; Control Charts for the Mean; Control Charts, Overview; Exponentially Weighted Moving Average (EWMA) Control Chart; Multivariate Cumulative Sum (CUSUM) Chart.

FAH FATT GAN

Moving Range and R Charts

Introduction

The quality of a manufactured item is often measured by one or more quality measurements. This measurement(s) varies from item to item following a (joint) distribution. The parameter or parameters of the (joint) distribution of the quality measurement(s) are measures of the quality of the process. As an example consider a production process that fills bottles with a soft drink that are advertised to contain 16 ounces. The fill X is a measure of the quality of a container. The mean fill μ and the standard deviation σ of the distribution of fills X are typical quality measures of the process. Loss of customers can be associated with μ being too low and unnecessary cost of production when μ is too high. Increases in σ are associated with increases in the proportion of unacceptable over and/or under filled bottles and decreases with quality improvement. It is therefore desirable to control both μ and σ .

The meaning of the process being “in control” is most commonly based on what Shewhart [1] referred to as *natural causes of variability*. We assume that these causes of variability are inherent and can only be removed by redesigning the process. Here we use the word *in control* to mean the state of the process when there are no assignable causes of variability present. When our filling process is in control, we will assume that μ and σ have the values μ_0 and σ_0 , respectively. If assignable causes of variability are present, then $\mu \neq \mu_0$ or $\sigma \neq \sigma_0$ and the process is said to be “out of control.”

Initially, the practitioner must bring the process into a state of statistical control. The Phase I **control chart** is often recommended as an aid to the practitioner for this purpose. The Shewhart Phase I $\bar{X}$ and $\bar{R}$ charts are the most commonly recommended charts to aid the practitioner in bringing the mean μ into a state of statistical control (see Ryan [2] for a discussion of control charts for the mean). The Phase I Shewhart moving range (MR), the sample range (R), and the sample standard deviation (S) charts are recommended for σ . These charts are designed to look at the data collected from the process in retrospect to answer the question “were these data collected from

a process that is in control?” When the process is believed to be in control, it is desirable to monitor the process for a possible change from in control to out of control. The question the practitioner wishes to answer is “has the process changed?” The charts used in the monitoring phase are often referred to as Phase II charts. The Phase II Shewhart MR , R , and S charts are most often used to monitor σ . The R and S (with runs rules) (see Burroughs *et al.* [3], Mitra [4]), the run sum R and S (see Champ and Rigdon [5]), and cumulative sum (**CUSUM**) R and S (see Page [6], Gan [7], Acosta-Mejia [8], and Amin and Wilde [9]), have also been developed to monitor for a change in σ . For a general overview of control charts, see Faltin [10].

Typically, it is recommended that the process standard deviation σ also be controlled when controlling the process mean μ . This is due to the fact that under the normal model the distribution of the plotted statistic for the Shewhart $\bar{X}$ and $\bar{R}$, CUSUM, and **exponentially weighted moving average** (EWMA) charts depend on σ . Thus, if the chart for monitoring μ indicates a possible change in the process this could be due to a change in σ . It is then argued that a chart for controlling for a change σ be examined first. If there is no indication that σ has changed, then any change in the process indicated by the chart for controlling the mean is interpreted as being due to a change in μ . It can be shown that the MR chart is affected not only by changes in σ but also by changes in μ whereas the R and S charts are only affected by changes in σ .

In this article, we will examine the design and use Phase I and II MR and R charts. The charts are designed under the assumption that the quality measurements are independent each having a **normal distribution**. The relative performance of these charts can be examined under the assumption that both μ_0 and σ_0 are known, one is to be estimated, or both are to be estimated.

Moving Range and R Charts

The Shewhart MR and R charts are designed under the assumption that the observed values of X are independent and each has a normal distribution. Also, it is assumed that periodically a random sample of size $n \geq 1$ is taken from the output of the production process. For convenience of discussion, we let $X_{t,j}$

2 Moving Range and R Charts

represent the X measurement on the j th item in sample t . The MR chart is often recommended for the case in which $n = 1$. The common scenario in Phase I is that the practitioner will have available X measurements on the items in m independent random samples each of size n that have been taken periodically from the output of the production process. In our example, the practitioner will have available the measurements $\{X_{t,1}, \dots, X_{t,n}\}$ for $t = 1, \dots, m$.

When $n = 1$, we define the sequence of moving ranges $MR_2, \dots, MR_m$ by

$$MR_t = |X_{t-1,1} - X_{t,1}| \quad (1)$$

It can be shown under our model assumptions that if the process is in control the mean and standard deviation of the distribution MR_t are given, respectively, by $(2/\sqrt{\pi})\sigma_0$ and $\sqrt{2(\pi-2)}/\pi\sigma_0$. If σ_0 were known, the Shewhart 3σ limits upper control limits (UCL) and lower control limits (LCL) for the Phase I Shewhart MR chart are

$$\begin{aligned} LCL &= 0 \quad \text{and} \\ UCL &= \frac{2 + 3\sqrt{2(\pi-2)}}{\sqrt{\pi}}\sigma_0 \approx 3.686\sigma_0 \end{aligned} \quad (2)$$

since $(2 - 3\sqrt{2(\pi-2)})/\sqrt{\pi} \approx -1.429 < 0$. Generally, it is not the case that σ_0 is known. The value of σ_0 is then estimated by the average $\overline{MR}$ of the $m-1$ MR s times an unbiased constant $\sqrt{\pi}/2$, that is, $\overline{MR} = (m-1)^{-1} \sum_{t=2}^m MR_t$. The Shewhart 3σ control limits are then given by

$$\begin{aligned} LCL &= 0 \quad \text{and} \\ UCL &= \frac{2 + 3\sqrt{2(\pi-2)}}{\sqrt{\pi}} \left(\frac{\sqrt{\pi}}{2} \overline{MR} \right) \doteq 3.2665 \overline{MR} \end{aligned} \quad (3)$$

The chart “signals” a potential out-of-control process if $MR_t > UCL$. If a positive LCL is desired or if it is desired to have an overall probability α of at least one of the MR_t ’s indicating a potential out-of-control process when the process is in control, the control limits could be adjusted by selecting the control limits as

$$\begin{aligned} LCL &= (1 - k_L \sqrt{\pi-1}) \overline{MR} = h_L \overline{MR} \quad \text{and} \\ UCL &= (1 + k_U \sqrt{\pi-1}) \overline{MR} = h_U \overline{MR} \end{aligned} \quad (4)$$

where k_L (h_L) and k_U (h_U) are appropriately selected. It is not difficult to show that as the number of samples m increases the probability α the chart signals a potential out-of-control process for fixed chart parameters k_L (h_L) and k_U (h_U) also increases. Some criteria for selecting these chart parameters are such that (a) it is just as likely to observe a value for the plotted statistic less than the LCL as greater than the UCL when the process is in control and (b) it is least likely to observe at least one of the plotted statistics outside the control limits when the process is in control. Using one or both of these criteria for selecting the control limits of the MR chart has not been studied in the literature.

The control limits for the Phase II MR chart are defined in a similar way to those for the Phase I MR chart. The difference is the method for choosing the chart parameters $k_{MR,L}$ ($h_{MR,L}$) and $k_{MR,U}$ ($h_{MR,U}$) and the value of $\overline{MR}$ is based on m independent samples each of size $n = 1$ from a process that is believed to be in control. Typically, the chart parameters of the Phase II MR chart are chosen such that the **average run length** (ARL) of the chart is relatively large when the process is in control and as small as possible when the process is out of control. The run length is the number of samples taken in Phase II until the plotted statistic plots either below LCL or above UCL.

For the case of $n > 1$, often the Phase I Shewhart R chart is recommended as an aid to the practitioner in bringing the process into a state of statistical control. The sample range R is the difference between the maximum and minimum values of the sample. For a discussion of the sample range R see Montgomery [11] and Alwan [12]. Under the normal model, the Phase I Shewhart R has LCL and UCL are defined by

$$\begin{aligned} LCL &= \left(1 - k_{R,L} \frac{d_3}{d_2}\right) \overline{R} \quad \text{and} \\ UCL &= \left(1 + k_{R,U} \frac{d_3}{d_2}\right) \overline{R} \end{aligned} \quad (5)$$

where $\mu_R = d_2\sigma_0$ and $\sigma_R = d_3\sigma_0$. Values of d_2 and d_3 are functions of the sample size n and must be determined numerically. Table 1 gives values for d_2 and d_3 for $n = 2, \dots, 12$. Harter [13], Montgomery [11], and Alwan [12] provide tables of these values.

Often it is recommended that $k_{R,L}$ and $k_{R,U}$ be set to 3 as suggested by Shewhart [1]. As with the MR

Table 1 Constants d_2 and d_3^2

n	d_2	d_3^2
2	1.1283791671	0.7267604553
3	1.6925687506	0.7891977107
4	2.0587507460	0.7740624738
5	2.3259289473	0.7466376009
6	2.5344127212	0.7191713092
7	2.7043567512	0.6942311313
8	2.8472006121	0.6721236717
9	2.9700263244	0.6525962151
10	3.0775054617	0.6352897762
11	3.1728727038	0.6198643117
12	3.2584552798	0.6060285277

chart, the chart parameters $k_{R,L}$ and $k_{R,U}$ are functions of the number of samples m for a given value of α . Similar to the Phase II Shewhart MR chart, the control limits for the Phase II Shewhart R charts are defined using the same formulas as in Phase I. The differences are that the value of $\bar{R}$ is based on m independent random samples each of size n that are believed to have been taken from a process that is in control. Further, the chart parameters for these charts are usually selected so that the run length distribution when the process is in control has certain properties such as a given in control ARL .

Example

Periodically, from our bottle filling process, we have selected $m = 20$ independent samples each of size $n = 5$. The X fills of these bottles along with the MR for the first sampled value and the range of each sample are recorded in the Table 2.

To illustrate the MR chart, we will consider that we have available only the measurement $x_{t,1}$ taken on the first item in each sample. For these data, the MR chart has $LCL = 0$ and $UCL = 3.2670 \times 0.1974 = 0.6449$ using equation (3). The Phase I MR is displayed in Figure 1.

The control limits for the Phase I R chart are $LCL = 0$ and $= 2.1145 \times 0.1159 = 0.2451$ using equation (5). There is no indication by the Phase I MR chart that these data are from a process that is out of control.

The Phase I R chart is displayed in Figure 2. We see from Figure 2 that the range of sample taken at sampling stage 5 signals a potential out-of-control process. It was also noted that the range of the 17th sample is relative large. Upon further investigation, the practitioner decides that the relatively large values of the range were due to natural causes of variability.

Table 2 Phase I Data

t	$x_{t,1}$	$x_{t,2}$	$x_{t,3}$	$x_{t,4}$	$x_{t,5}$	$MR_{t,1}$	R_t
1	32.0175	31.9947	31.9916	31.9927	32.0093	—	0.0258
2	32.1399	32.0075	31.9786	31.9925	32.0108	0.1224	0.1612
3	31.9159	31.9926	32.0039	31.9993	31.9898	0.2240	0.0880
4	32.2082	31.9892	32.0071	31.9961	31.9990	0.2923	0.2190
5	31.7144	31.9853	31.9933	31.9894	31.9932	0.4938	0.2788
6	32.0526	31.9971	31.9959	32.0160	31.9789	0.3381	0.0737
7	32.0037	31.9977	31.9943	32.0123	32.0011	0.0488	0.0180
8	31.9426	31.9937	31.9842	31.9874	31.9947	0.0611	0.0521
9	32.0792	32.0035	32.0056	32.0029	32.0019	0.1366	0.0774
10	31.8532	31.9883	32.0092	31.9857	31.9892	0.2261	0.1561
11	32.0017	31.9988	31.9972	31.9969	31.9901	0.1485	0.0116
12	31.9017	32.0129	31.9970	31.9925	32.0040	0.1000	0.1112
13	32.1456	31.9828	31.9991	31.9779	31.9872	0.2439	0.1677
14	32.1496	32.0012	31.9983	31.9890	32.0035	0.0039	0.1605
15	32.0326	32.0084	32.0037	32.0028	31.9980	0.1169	0.0346
16	31.8605	32.0242	31.9850	31.9952	31.9961	0.1722	0.1637
17	32.2333	32.0059	31.9956	32.0065	32.0038	0.3728	0.2376
18	31.9382	32.0078	31.9914	31.9930	31.9951	0.2950	0.0696
19	32.0970	32.0186	32.0006	32.0050	32.0087	0.1587	0.0964
20	31.9029	32.0065	32.0165	31.9959	32.0182	0.1940	0.1153

$$\bar{x} = 32.0006, \bar{MR} = 0.1974, \bar{R} = 0.1159, \text{ and } \bar{S} = 0.0485.$$

4 Moving Range and R Charts

Table 3 Phase II Data

t	$x_{t,1}$	$x_{t,2}$	$x_{t,3}$	$x_{t,4}$	$x_{t,5}$	$\bar{x}_t$	$MR_{t,1}$	R_t
1	31.7649	31.9616	31.9763	31.9691	31.9981	31.9340	—	0.2332
2	32.2237	32.0310	32.0029	32.0056	32.0406	32.0608	0.4588	0.2208
3	31.9838	32.0089	32.0017	31.9927	31.9965	31.9967	0.2399	0.0251
4	31.8351	31.9757	32.0009	31.9754	31.9646	31.9503	0.1487	0.1658
5	31.9435	32.0098	31.9967	32.0081	32.0057	31.9928	0.1084	0.0663
6	32.3005	31.9935	31.9983	31.9894	32.0040	32.0572	0.3570	0.3111

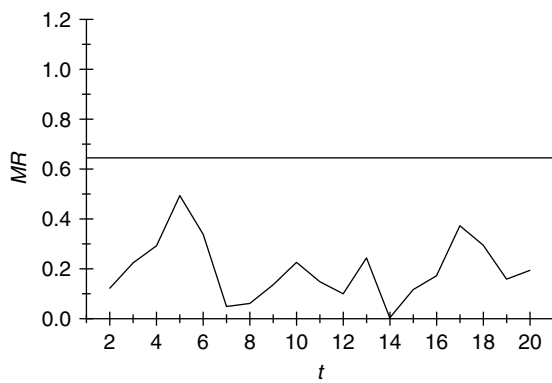


Figure 1 Phase I MR chart with $m = 20$

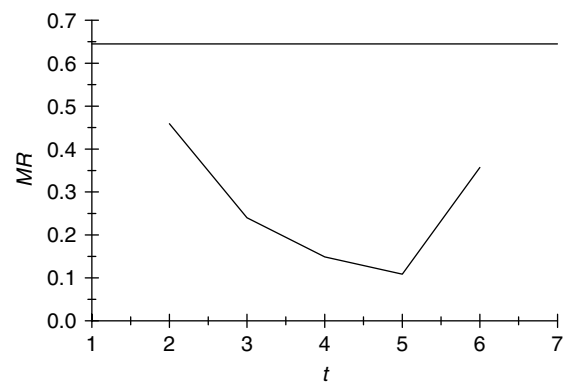


Figure 3 Phase II MR chart

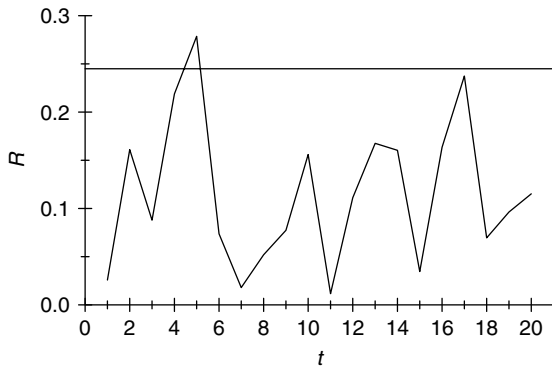


Figure 2 Phase I R chart with $m = 20, n = 5$

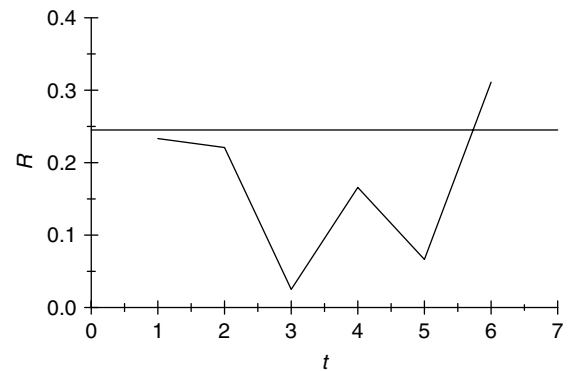


Figure 4 Phase II R chart

Since it is believed that the data collected in Phase I is in control, then the same $3\text{-}\sigma$ limits are used for the Phase II chart. Table 3 contains the data collected in Phase II for the first six samples.

The Phase II MR and R charts are given respectively in Figures 3 and 4.

The Phase II MR chart provides no evidence

that the process has changed. The sample range R_6 of the sixth sample in Phase II signals a potential out-of-control process suggesting to the practitioner that the process should be examined to see it had changed.

Conclusion

In this article, we have discussed the *MR* and *R* charts. These charts are simple to use and interpret. The *MR* chart unlike the *R* chart is affected by changes in both the mean and the standard deviation of the distribution of the quality measurement. Also, it has been shown by Rigdon *et al.* [14] that the Phase II *X* chart performs about as well as the combined Phase II *X* and *MR* chart. We do not recommend using the *MR* chart.

References

- [1] Shewhart, W.A. (1931). *Economic Control of Quality of Manufactured Product*, D. Van Nostrand, New York.
- [2] Ryan, T. (2007). Control charts for the mean, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons, New York.
- [3] Burroughs, T.E., Rigdon, S.E. & Champ, C.W. (1995). Analysis of the Shewhart *S*-chart with runs rules when no standards are given, in *Proceedings of the Twenty-Sixth Annual Meeting of the Midwest Decision Sciences Institute*, St. Louis, pp. 268–270, 4–6 May 1995.
- [4] Mitra, A. (2007). Standard deviation charts, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons, New York.
- [5] Champ, C.W. & Rigdon, S.E. (1998). Monitoring variability using the run sum chart, in *Proceedings of the Quality and Productivity Section of the American Statistical Association*, Dallas, pp. 24–25, 9–13 August 1998.
- [6] Page, E.S. (1963). Controlling the standard deviation by CUSUMs and warning lines, *Technometrics* **5**, 307–315.
- [7] Gan, F.F. (2007). CUSUM charts, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons, New York.
- [8] Acosta-Mejia, C.A. (1998). Monitoring reduction in variability with the range, *IIE Transactions* **30**, 515–523.
- [9] Amin, R.W. and Wilde, M. (2000). Two-Sided CUSUM Control Charts for Variability, *Allgemeines Statistisches Archiv* **84**, 295–313.
- [10] Faltin, F.W. (2007). Control charts, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons, New York.
- [11] Montgomery, D.M. (2005). *Introduction to Statistical Quality Control*, 5th Edition, John Wiley & Sons, New York.
- [12] Alwan, L.C. (2001). *Statistical Process Analysis*, Irwin McGraw-Hill, Boston.
- [13] Harter, H.L. (1960). Tables of range and studentized range, *Annals of Mathematical Statistics* **31**, 1122–1147.
- [14] Rigdon, S.E., Cruthis, E.N. & Champ, C.W. (1994). Design strategies for individual and moving range control charts, *Journal of Quality Technology* **26**, 274–287.

Further Reading

- Amin, R.W. & Matthias, W. (2000). Two-sided CUSUM control charts for variability, *Allgemeines Statistisches Archiv* **84**, 295–313.
- Vander Wiel, S. (2007). Exponentially weighted moving average control charts, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons, New York.

CHARLES W. CHAMP AND DEBORAH K.
SHEPHERD

Control Charts for the Standard Deviation

Introduction

All processes, regardless of the degree of sophistication, exhibit variation. This variation may be caused by factors over which we have some degree of control, for example, equipment, materials, methods, personnel, or policies. Alternatively, environmental factors over which we have no control may also contribute to variability. A control chart is a graphical tool that monitors an ongoing process (*see Control Charts, Overview*). The selected quality characteristic is plotted along the vertical axis, whereas the sample or subgroups, selected in order of time, are plotted along the horizontal axis.

A control chart for the standard deviation monitors process variability by plotting the sample standard deviation through an *s-control chart*. The sample standard deviation (s) of a subgroup (or sample) of observations ($X_i, i = 1, 2, \dots, n$) of size n is a measure of their variability around the sample mean and is given by

$$s = \sqrt{\frac{\sum_{i=1}^n (X_i - \bar{X})^2}{n-1}} = \sqrt{\frac{\sum_{i=1}^n X_i^2 - \left(\sum_{i=1}^n X_i\right)^2 / n}{n-1}} \quad (1)$$

where $\bar{X}$ represents the sample mean. Typically, monitoring a process for changes in location and variation is done through simultaneous operation of a chart for the sample mean ($\bar{X}$) and a chart for the standard deviation (s). An alternative is to monitor control charts for the mean ($\bar{X}$) (*see Control Charts for the Mean*) and range (R) (*see Moving Range and R Charts*). However, when the subgroup size is large ($n > 10$), the s chart is preferable to the R chart [1].

Preliminary Considerations

There are two assumptions made in the construction of control charts for the standard deviation. First, the observations are assumed to be independent and identically distributed (i.i.d.). Second, the

observations are assumed to be normally distributed. When the first assumption is not valid, control charts for autocorrelated data could be pursued [2]. These charts will not be discussed in this paper. When the second assumption is not valid, techniques for transforming the observations such that the transformed data meet the assumption of normality, is a possible option. In this context, a common transformation is the **Box–Cox transformation** [3], discussed briefly in a later section. Some preliminary considerations on the manner in which samples are to be selected, the size and frequency of sampling, and the placement of the control limits based on a chosen false alarm rate are necessary prior to the construction of control charts.

Selection of Rational Subgroups

Subgroups are chosen such that the variation within them is only due to common causes. Hence, if special causes are present, subgroups are selected such that they will occur between samples. So, using these concepts, differences *within* samples will be minimized while differences *between* samples will be maximized. Sampling is usually conducted by time order, where observations within a subgroup are homogeneous. Usually, at least 25–30 subgroups of observations must be obtained prior to the construction of control charts.

Sample Size

The sample or subgroup size (n) represents the number of observations in each sample. It is usually greater than or equal to 4. The degree of shift in the process parameter that is expected to take place will influence the choice of the sample size. The larger the sample size, the better the chance of detecting small shifts. Other factors such as the cost of obtaining the information or the cost of shipping a nonconforming item to the customer also influence the choice of sample size.

Frequency of Sampling

Choosing large samples very frequently is the sampling scheme that provides the most information. However, resource constraints may not allow us to do so. Other options include choosing small sample

2 Control Charts for the Standard Deviation

sizes at frequent intervals or large sample sizes at infrequent intervals. The former method is usually chosen. Other factors also influence the frequency of sampling and the sample size. The current state of the process is one factor. For a stable process that is in control, we could get by with sampling at infrequent intervals. The cost of sampling and the loss incurred due to a nonconforming item being passed on to the customer are other factors.

Errors in Decision Making from Control Charts

Decision making from a control chart is analogous to testing a hypothesis. Since sample information is used to make a decision on a population parameter, two types of errors may happen.

False Alarm. A false alarm (known as a *Type I error*) occurs from inferring that a process is out of control when it is actually in control. The probability of a false alarm is denoted by α . Let us suppose that we are using the decision rule of concluding a process to be out of control if the sample statistic plots outside the control limits. Suppose the process is in control. Since the control limits are at a finite distance (usually three standard deviations of the plotted quality characteristic) from the centerline (CL), there is a small chance (given by α) of the sample statistic falling outside the control limits. In such instances, inferring that the process is out of control is erroneous and is known as a *false alarm*. The false alarm can be minimized by having control limits that are wider apart (e.g., four standard deviations of the plotted statistic from the CL). Usually, the control limits are constructed based on a chosen level of significance (α). For a monitored statistic that is normally distributed, a common choice of placement of the control limits is three standard deviations from the CL, yielding a false alarm rate of about 0.0026.

Failure to Signal. A failure to signal (known as a *Type II error*) occurs from inferring that a process is in control when it is really out of control. Suppose our decision rule for concluding that a process is in control is for the observed sample statistic that is being monitored to fall within the control limits. Now suppose that, owing to a change in the vendor, the variation in the quality of the incoming raw material has increased. However, it is possible that

a random sample, chosen after the occurrence of this change, could yield a sample standard deviation that lies within the previously established control limits. Thus, we would erroneously conclude that the process is in control – a failure to signal, the probability of which is denoted by β . The magnitude of the failure to signal error is, however, dependent on the degree of change; in this case the increase in variability. The larger the change, the less the chances of committing a Type II error. The failure to signal error can be minimized by moving the control limits closer to each other (such as two standard deviations of the plotted statistic from the CL), thereby increasing the chance of detection of an out-of-control process. But, doing so will increase the chance of a false alarm. Hence, we note the inverse relationship between the false alarm and failure to signal errors, for a fixed sample size, as the location of the control limits are changed. A measure of performance of control charts is the ability to quickly detect a process change, when the process is out of control. If we denote P_d to be the probability of detection of an out-of-control condition, $P_d = 1 - \beta$. Accordingly, a performance measure is the number of subgroups, on average, to detect an out-of-control signal given by the *average run length* (ARL) where

$$ARL = \frac{1}{P_d} = \frac{1}{(1 - \beta)} \quad (2)$$

Location of Control Limits

Control limits are placed at k standard deviations of the plotted statistic from the CL. Selecting a value for k is influenced by the chosen false alarm rate and by the probability of detecting an out-of-control condition when a process change of a certain magnitude takes place. This involves a trade-off between the false alarm and failure to signal errors. Usually, k is selected based on a chosen level of the false alarm rate (α), where importance is placed on limiting the probability of a false alarm. Assuming normality of distribution of the characteristic, k is typically chosen to be 3, justifying the name “ 3σ ” limits, for which the probability of a false alarm (α) is limited to 0.0026. In this chapter, we will discuss and present only “ 3σ ” control limits, even though, in general, k could take on other values.

Rules for Identifying Out-of-Control Conditions

A major objective of using control charts is to determine when a process is out of control so that necessary remedial actions may be taken. A systematic or nonrandom pattern of the plotted values suggests an out-of-control condition. Some commonly used rules are as follows:

1. A single-point plots outside the control limits.
2. Two out of three consecutive points fall outside the “ 2σ ” warning limits on the same side of the CL.
3. Four out of five consecutive points fall outside the “ 1σ ” limit on the same side of the CL.
4. Nine or more consecutive points fall to one side of the CL.

The first rule is the one that is routinely used. As more rules are used simultaneously, the chance of a false alarm increases. Suppose that the number of independent rules used is p , and let α_i be the probability of a false alarm of rule i . The *overall false alarm rate* is given by

$$\alpha = 1 - \prod_{i=1}^p (1 - \alpha_i) \quad (3)$$

Construction of Standard Deviation Control Charts

For a normally distributed quality characteristic with a *population standard deviation* denoted by σ , the mean and the standard deviation of the sample standard deviation [2] are given by

$$E(s) = c_4 \sigma \quad (4)$$

$$\sigma_s = \sigma \sqrt{1 - c_4^2} \quad (5)$$

respectively, where c_4 is a factor that depends on the sample size and is given by

$$c_4 = \left[\frac{2}{(n-1)} \right]^{1/2} \frac{\Gamma(n/2)}{\Gamma[(n-1)/2]} \quad (6)$$

The expression $\Gamma(\cdot)$ represents the gamma function.

No Given Standards

The CL of a standard deviation chart is

$$CL_s = \bar{s} = \frac{\sum_{i=1}^g s_i}{g} \quad (7)$$

where g is the number of samples and s_i is the standard deviation of the i th sample. The upper control limit is

$$UCL_s = \bar{s} + 3\sigma_s = \bar{s} + 3\sigma \sqrt{1 - c_4^2} \quad (8)$$

In accordance with equation (4), an estimate of the population standard deviation σ is

$$\hat{\sigma} = \frac{\bar{s}}{c_4} \quad (9)$$

Substituting this estimate of $\hat{\sigma}$ in the preceding expression yields

$$UCL_s = \bar{s} + \frac{3\bar{s}\sqrt{1 - c_4^2}}{c_4} = B_4 \bar{s} \quad (10)$$

where $B_4 = 1 + 3 \frac{\sqrt{1 - c_4^2}}{c_4}$. Similarly,

$$LCL_s = \bar{s} - \frac{3\bar{s}\sqrt{1 - c_4^2}}{c_4} = B_3 \bar{s} \quad (11)$$

where $B_3 = \max \left[0, 1 - 3\sqrt{(1 - c_4^2)}/c_4 \right]$. Thus, the 3σ control limits are

$$UCL_s = B_4 \bar{s} \quad (12)$$

$$LCL_s = B_3 \bar{s} \quad (13)$$

The CL of the chart for the mean $\bar{X}$ is given by

$$CL_{\bar{X}} = \bar{\bar{X}} = \frac{\sum_{i=1}^g \bar{X}_i}{g} \quad (14)$$

The control limits on the $\bar{X}$ chart are

$$\bar{\bar{X}} \pm 3\sigma_{\bar{X}} = \bar{\bar{X}} \pm \frac{3\sigma}{\sqrt{n}} \quad (15)$$

4 Control Charts for the Standard Deviation

where $\bar{\bar{X}}$ is the CL and represents the grand mean (or mean of the sample means). Using equation (9) to obtain $\hat{\sigma}$, we find the control limits to be

$$UCL_{\bar{X}} = \bar{\bar{X}} + \frac{3\bar{s}}{c_4\sqrt{n}} = \bar{\bar{X}} + A_3\bar{s} \quad (16)$$

$$LCL_{\bar{X}} = \bar{\bar{X}} - \frac{3\bar{s}}{c_4\sqrt{n}} = \bar{\bar{X}} + A_3\bar{s} \quad (17)$$

where $A_3 = 3/(c_4\sqrt{n})$.

The above-constructed limits are known as *trial control limits*. The values of the sample mean ($\bar{X}_i$) and sample standard deviation (s_i) are plotted on the respective control charts along with the CL and the trial control limits. On the basis of the chosen rules for out-of-control conditions, we determine if any points are out of control on the s chart or the $\bar{X}$ chart. The s chart is looked at first. Only if it is in control should the $\bar{X}$ chart be developed, since the standard deviation of $\bar{X}$ is dependent on $\bar{s}$. If the s chart is not in control, any estimate of the standard deviation of $\bar{X}$ will be unreliable, which will create unreliable control limits for $\bar{X}$ [1.]

When there are out-of-control points, the investigator tries to identify special causes associated with them, if any. The pattern of the plot may indicate possible special causes and, further, the corresponding remedial actions. For example, a sequence of six or more consecutive points that are increasing on an s chart may indicate a progressive increase in the variability of the process. On investigation, a possible cause could be the gradual deterioration in the raw material quality due to the cumulative effects of humidity over time. After identification of remedial actions to eliminate the special causes, the out-of-control points are deleted, and the revised CL and control limits are calculated. This cycle of obtaining process information, determining trial limits, finding out-of-control points, identifying and correcting special causes, and determining revised control limits continues. The revised control limits serve as trial control limits for the immediate future until the limits are revised again. Such an ongoing process is a critical component of continuous improvement.

Given Standard

If a target standard deviation is specified as σ_0 , the CL of the s chart is found by using equation (4) as

$$CL_s = c_4\sigma_0 \quad (18)$$

The upper control limit (UCL) for the s chart is found by using equation (5)

$$\begin{aligned} UCL_s &= c_4\sigma_0 + 3\sigma_s = c_4\sigma_0 + 3\sigma_0\sqrt{1 - c_4^2} \\ &= \left(c_4 + 3\sqrt{1 - c_4^2}\right)\sigma_0 = B_6\sigma_0 \end{aligned} \quad (19)$$

where $B_6 = c_4 + 3\sqrt{1 - c_4^2}$. Similarly, the lower control limit (LCL) for the s chart is

$$LCL_s = \left(c_4 - 3\sqrt{1 - c_4^2}\right)\sigma_0 = B_5\sigma_0 \quad (20)$$

where $B_5 = \max\left[0, c_4 - 3\sqrt{1 - c_4^2}\right]$. Thus, the control limits for the s chart are

$$UCL_s = B_6\sigma_0 \quad (21)$$

$$LCL_s = B_5\sigma_0 \quad (22)$$

If a target value for the mean is specified as $\bar{X}_0$, then the CL on the $\bar{X}$ chart is given by

$$CL_{\bar{X}} = \bar{X}_0 \quad (23)$$

Equations for the control limits on the $\bar{X}$ chart are given by

$$UCL_{\bar{X}} = \bar{X}_0 + A\sigma_0 \quad (24)$$

$$LCL_{\bar{X}} = \bar{X}_0 - A\sigma_0 \quad (25)$$

where $A = 3/\sqrt{n}$.

Interpretations of the s Chart

Note that the LCL on the standard deviation control chart cannot be negative. Hence, using “3 σ ” control limits, if the LCL is calculated to be negative, it is converted to zero.

When the LCL on the s chart is positive and an observation plots below it, the process is considered to be out of control. However, in this case it is a desirable condition that indicates that the variability in the process is unusually small. When the corresponding special cause is identified, management should make an effort to move process conditions toward this situation so that the reduced process variability may be sustained in the future.

Related Topics

Estimating Process Standard Deviation

A point of interest from the standard deviation control chart is to obtain a suitable estimate of the process standard deviation (σ). There are three common estimates of σ .

1. One approach is to base the estimation on the average range of the samples. Let R_i denote the range of observations for sample i , $i = 1, 2, \dots, g$. The average range is given by

$$\bar{R} = \frac{\sum_{i=1}^g R_i}{g} \quad (26)$$

When sampling from a normal population, the distribution of the relative range (W), given by R/σ , is dependent on the sample size, n . The mean of W is represented by a factor d_2 [1]. Hence, an estimate of σ is obtained from

$$\hat{\sigma}_1 = \frac{\bar{R}}{d_2} \quad (27)$$

2. A second estimate of σ is found from the average of the standard deviations of the individual samples, where the sample size of each is n . This estimator is given by equation (9):
3. A third estimate of σ is based on the pooled estimate of the sample variances. If the sample size is constant for each sample and denoted by n , the pooled estimate is given by

$$s_p = \left[\frac{1}{g} \sum_{i=1}^g s_i^2 \right]^{1/2} \quad (28)$$

An unbiased estimator of the process standard deviation is obtained as

$$\hat{\sigma}_p = \frac{s_p}{c_4} \quad (29)$$

When the subgroup sizes vary from subgroup to subgroup, a pooled estimate of the process standard deviation is found as

$$s_p^* = \left[\frac{\sum_{i=1}^g (n_i - 1) s_i^2}{\sum_{i=1}^g n_i - g} \right]^{1/2} \quad (30)$$

Correspondingly, an unbiased estimator of the process standard deviation is given by

$$\hat{\sigma}_p^* = \frac{s_p^*}{c_4} \quad (31)$$

where c_4 is calculated using equation (6) with an equivalent sample size of $(\sum n_i - g + 1)$. This pooled estimate, s_p^* , may be used to construct the CL and control limits for the s chart when the subgroup size varies. The procedure is similar to that outlined before (where $\bar{s}$ is replaced by s_p^*). However, the control limits are not constant, but change from subgroup to subgroup.

4. One major objective of the control chart is to determine out-of-control conditions quickly, when the process is out of control. On the basis of simulation results that determine the average and the standard deviation of the run length, for detection of out-of-control conditions [4], the pooled estimate (s_p) is found to be desirable, relative to the other estimates.

Estimating Process Capability

The capability of a process is its ability to meet customer specifications. Capability can be estimated only when a process is in control, since only then is the output from a process stable and predictable. Hence, control charts serve a purpose in **estimating process capability**.

Typical capability measures involve computations that compare the specification spread to the process spread. Since the process spread is influenced by the estimated process standard deviation, the type of estimate used has an impact on the estimated capability. For short-term capability measures, estimates of the process standard deviation given by equations (9), (27), (29), or (31), may be used. These estimates are measures of the within-subgroup variation.

For long-term capability measures, a measure of the overall process variability is given by

$$s_{LT} = \frac{\sum_{i=1}^g \sum_{j=1}^{n_i} (X_{ij} - \bar{\bar{X}})^2}{(\sum n_i) - 1} \quad (32)$$

where $\bar{\bar{X}}$ represents the grand average of all the observations, once the process is in control.

6 Control Charts for the Standard Deviation

Table 1 Copper deposition thickness (in microns) in printed circuit boards

Subgroup number	Observations					Subgroup number	Observations				
1	43.47	38.43	44.41	36.48	31.42	14	29.58	39.86	34.84	21.77	36.63
2	37.25	40.65	46.40	22.73	30.74	15	23.88	47.92	22.71	21.98	28.76
3	19.78	19.13	20.31	26.48	33.69	16	21.76	35.77	34.98	33.36	33.20
4	22.39	31.54	24.40	20.63	31.64	17	17.86	26.15	56.48	15.83	38.37
5	36.61	23.31	27.93	25.45	19.65	18	35.18	30.73	30.43	30.31	30.47
6	22.48	32.74	31.69	38.88	37.09	19	20.67	34.47	35.56	19.35	27.21
7	29.78	34.41	26.07	19.56	10.01	20	28.56	41.36	28.67	33.61	39.63
8	43.38	22.64	38.32	31.76	26.13	21	35.63	25.18	39.60	42.70	9.53
9	24.86	27.48	31.30	25.21	37.95	22	35.82	23.28	27.27	27.38	30.26
10	29.19	26.51	35.05	36.67	36.53	23	23.49	29.29	30.49	21.35	29.14
11	34.61	21.21	24.66	36.58	35.70	24	38.66	49.59	24.46	32.21	21.05
12	55.44	38.68	34.74	42.53	32.18	25	46.43	26.89	41.78	25.87	29.23
13	24.54	40.76	43.19	34.67	26.02						

An unbiased estimate of the overall process standard deviation is obtained from

$$\hat{\sigma}_{LT} = \frac{s_{LT}}{c_4} \quad (33)$$

where c_4 is calculated using equation (6) with an equivalent sample size of $(\sum n_i - g + 1)$.

Nonnormality of Distribution

When the distribution of the quality characteristic does not follow a **normal distribution**, one approach is to transform the data. A form of power transformation, where the transformed variable follows normality, is the Box–Cox transformation [3, 5] (*see Box and Cox Transformation*). It is given by

$$\begin{aligned} X_T &= X^\lambda, \lambda \neq 0 \\ &= \ell n(\lambda), \text{ when } \lambda = 0 \end{aligned} \quad (34)$$

The Box–Cox method estimates a value for λ which minimizes the standard deviation of the standardized transformed variable [3, 5].

Another approach is to use a distribution-free or nonparametric procedure for constructing control charts, when the quality characteristic could be from any general distribution (*see Control Charts, Nonparametric*). While this approach vastly expands the scope of application of control charts, without any restriction on the distributional form of the characteristic, it may suffer from having a smaller power to detect process changes relative to parametric control charts.

Example

In the manufacture of printed circuit boards (PCBs), an electroplating process is utilized where the circuitry is plated with copper. The copper deposition rates are found to be varying quite a bit and are influenced by factors such as plating time, bath temperature, current density, and bath concentration. Table 1 shows the copper deposition thickness (in microns) for 25 randomly chosen groups from the process. Each subgroup consists of five observations, randomly selected in close proximity in time. The specifications of thickness are 25–35 μm .

The control chart constants for construction of the 3σ control charts for the s chart may be calculated using previously defined equations. For example, using equation (6), we get

$$c_4 = \left[\frac{2}{(5-1)} \right]^{1/2} \frac{\Gamma(5/2)}{\Gamma[(5-1)/2]} = 0.9400 \quad (35)$$

Consequently, the control chart factor for obtaining the UCL for the s chart is obtained as

$$B_4 = 1 + 3 \frac{\sqrt{1 - c_4^2}}{c_4} = 1 + 3 \frac{\sqrt{1 - 0.94^2}}{0.94} = 2.089 \quad (36)$$

Similarly, the control chart factor for obtaining the LCL for the s chart is found to be

$$\begin{aligned} B_3 &= \max \left[0, 1 - 3 \frac{\sqrt{1 - c_4^2}}{c_4} \right] \\ &= \max[0, -0.089] = 0 \end{aligned} \quad (37)$$

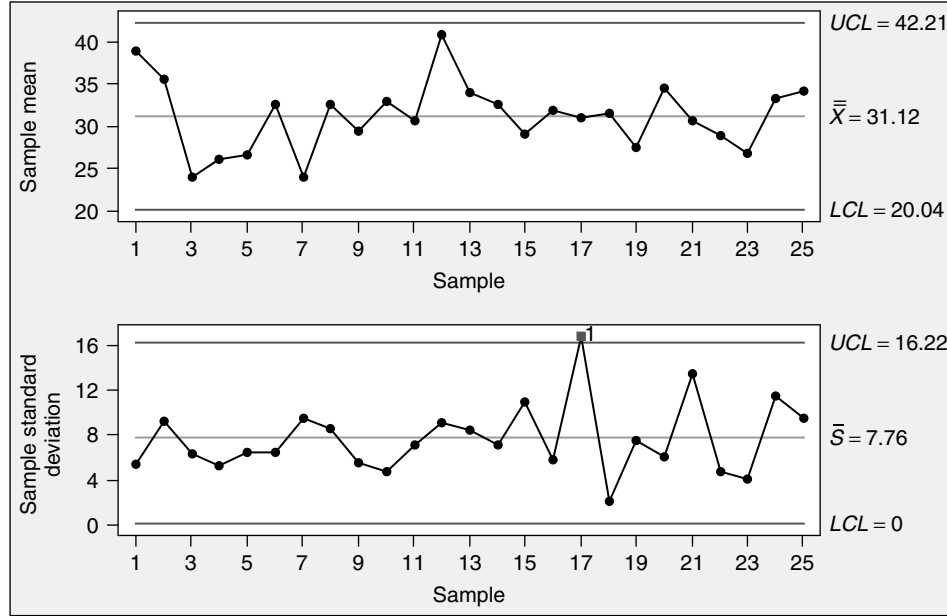


Figure 1 Control charts for the mean and standard deviation using trial control limits

For the $\bar{X}$ chart, the control chart factor for the control limits, as given by equation (16), is obtained as

$$A_3 = \frac{3}{(c_4\sqrt{n})} = \frac{3}{(0.9400\sqrt{5})} = 1.427 \quad (38)$$

We demonstrate construction of control charts for the mean ($\bar{X}$) and standard deviation (s) using a statistical software package, Minitab [6], that is widely used in practice. In order to use Minitab, the data must be input as a worksheet, where the columns are treated as variables and the rows represent the selected samples or subgroups. Here, each subgroup is input in a single row, with the five observations being input in columns 1 through 5 (C1–C5).

Since a prior estimate of the process mean or standard deviation is not specified, we use the sample data to estimate these parameters and accordingly construct control charts. The following commands are executed: *Stat > Control Charts > Variable Charts for Subgroups > X bar – S*. Indicate that the observations for a subgroup are in one row of columns. Input the column numbers, here C1–C5, since the worksheet contains the observations in the first five columns. Click on *X bar – S Options*, select *Estimate*, and under *Method for estimating standard*

deviation, choose either *S bar* or *Pooled standard deviation* (the default option). The *S bar* option uses the version given by equation (7), while the latter uses the versions given by equations (28) or (30). In each case, Minitab allows us to select “Use unbiasing constant”, which is the default option. Thus, the process standard deviation is estimated either by equation (9), or by equations (29) or (31), depending on the selection. Under the *Tests for special causes* option, we select the first rule, which is a single point outside the “ 3σ ” control limits. Click *OK*.

Figure 1 shows the control charts for the mean and standard deviation using trial control limits. On the $\bar{X}$ chart, the CL is 31.12, while the control limits are (42.21, 20.04). On the s chart, the CL is 7.76, while the control limits are (16.22, 0). Using the specified rule for detecting out-of-control conditions, we find that on the s chart, the sample standard deviation for subgroup number 17 is unusually high and falls above the UCL. This calls for investigation of possible special causes associated with this subgroup. On looking into the process, it is deemed that major changes in the bath concentration occurred. Accordingly, remedial measures were put into place that would indicate whenever the bath concentration is outside certain specified bounds.

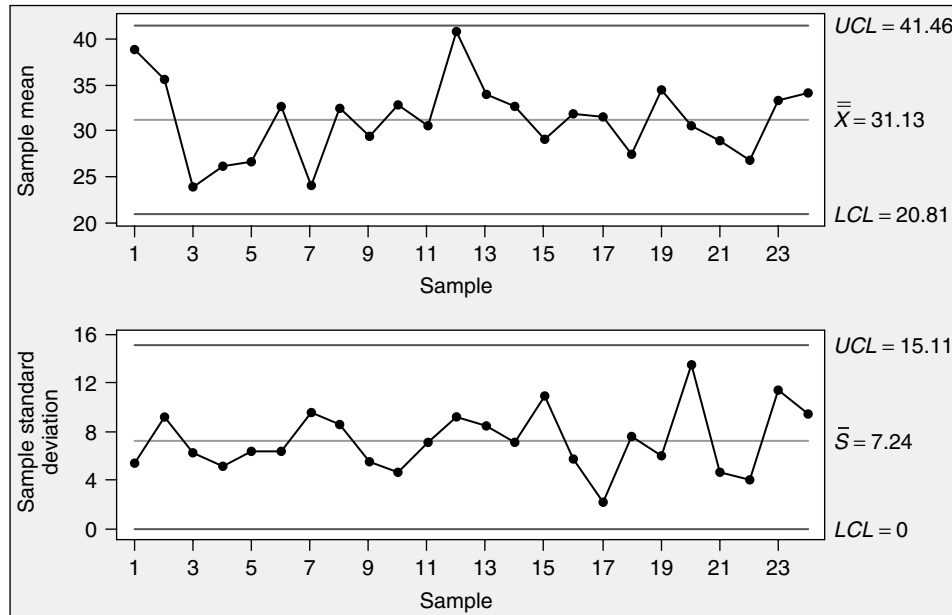


Figure 2 Control charts for the mean and standard deviation using revised control limits

Subgroup number 17 is now deleted, and revised control charts are constructed using the remaining subgroups. Figure 2 shows the revised control charts for the mean and standard deviation. Using the specified rule for detection of out-of-control conditions, none is found to be out of control. Now that we have a process that is in control and stable, we can estimate the process standard deviation. Since the revised CL on the s chart is 7.24, and the constant (c_4) from equation (6) using a subgroup size of 5 is 0.9400, an estimate of the process standard deviation using equation (29) is $7.24/0.94 = 7.702$. An estimate of the process average, from the $\bar{X}$ chart, is observed to be 31.13. These values may be used to determine process capability, the ability of the process, to meet the required specifications of 25–35 μm .

Conclusions

The standard deviation control chart is a very handy tool to monitor process variability as measured by its standard deviation. While a control chart for the range is another tool for monitoring process variation, the standard deviation chart is preferred to the range chart for subgroups of large size.

Process changes may take place either in the location of a process, variation of a process, or both. Control charts for the mean are used to monitor changes in process location. Hence, to monitor the stability of a process, either an $\bar{X}$ and s chart, or an $\bar{X}$ and R chart are used in conjunction with each other.

Estimation of the process standard deviation is a desired outcome of process control through an s chart. Such an estimate helps us assess process capability, which is the ability of the process to meet set specifications. Further, computation of proportion nonconforming of the product is possible once an estimate of the process mean and process standard deviation are available.

References

- [1] Mitra, A. (2006). *Fundamentals of Quality Control and Improvement*, 2nd Edition-updated, Thomson Publishing, Mason.
- [2] Montgomery, D.C. (2004). *Introduction to Statistical Quality Control*, 5th Edition, John Wiley & Sons, New York.
- [3] Chou, Y., Polansky, A.M. & Mason, R.L. (1998). Transforming nonnormal data to normality in statistical process control, *Journal of Quality Technology* **30**, 133–141.

- [4] Chen, G. (1997). The mean and standard deviation of the run length distribution of $\bar{X}$ charts when control limits are estimated. *Statistica Sinica* **7**, 789–798.
- [5] Pearn, W.L., Kotz, S. & Johnson, N.L. (1992). Distributional and inferential properties of process capability indices, *Journal of Quality Technology* **24**, 216–231.
- [6] Minitab Inc. 2007. Statistical software, Release 15, State College, Pennsylvania.

Related Articles

Control Charts, Overview; Control Charts for the Mean; Control Charts, Nonparametric Box and Cox Transformation; Moving Range and R Charts

AMITAVA MITRA

Control Charts for Attributes

Introduction

Control charts for attributes are useful when the quality characteristics are classifications rather than measures of a continuous trait. For example, when an item is to be classified as conforming to quality standards or nonconforming, attribute control charts are often used. The p -chart is an attribute chart based on the binomial model and is used to monitor the proportion of nonconforming items in a sample. Another type of attribute chart, the c -chart, is based on the Poisson model and is useful when a nonconforming item has multiple nonconformities or defects. A c -chart can be used to monitor the number of occurrences of defects over some interval of time or area rather than the proportion of nonconforming items in a sample. For example, one might monitor the number of surface blemishes per 20 sq ft of surface area of a finished product.

p -Chart

The proportion of nonconforming items in the t^{th} sample or subgroup is defined by

$$\hat{p}_t = \frac{X_t}{n} \quad (1)$$

where X_t is the number of nonconforming items in the subgroup and n is the total number of items in the sample. To establish the control limits, m subgroups of size n are gathered when the process is in-control and the number of nonconforming items is recorded for each subgroup. The average of the proportion of nonconforming items in the initial subgroups is defined as

$$\bar{p} = \frac{\sum_{i=1}^m X_i}{mn} \quad (2)$$

The lower control limit (LCL), center line (CL), and upper control limit (UCL) of the p -chart are given by

$$LCL = \bar{p} - 3\sqrt{\frac{\bar{p}(1-\bar{p})}{n}} \quad (3)$$

$$CL = \bar{p} \quad (4)$$

$$UCL = \bar{p} + 3\sqrt{\frac{\bar{p}(1-\bar{p})}{n}} \quad (5)$$

These limits are considered trial control limits. The subgroup proportions from the initial sample, $\hat{p}_1, \hat{p}_2, \dots, \hat{p}_m$, should be compared to the control limits, and if the values fall outside the limits, those subgroups should be investigated for assignable causes to an out-of-control situation. If assignable causes are found, these subgroups should be excluded from the calculation of $\bar{p}$, and new limits should be determined. It is often recommended that $m > 30$ subgroups be used to establish the chart. Once the chart is established, samples of size n are continually observed, and the proportion of nonconforming items in the subgroup, $\hat{p}_t$, is plotted.

Example of p -Chart

A manufacturer of cordless telephones tracks the proportion of nonconforming units produced. At the end of each of the three daily shifts, $n = 50$ phones are sampled from the production run. After the first 11 days, data is gathered from $m = 33$ subgroups of size $n = 50$. This information is given in Table 1.

The CL is given by $\bar{p} = 0.1085$; thus the trial control limits are

$$\begin{aligned} LCL &= \bar{p} - 3\sqrt{\frac{\bar{p}(1-\bar{p})}{n}} = 0.1085 - 0.1320 \\ &= -0.0235 \end{aligned} \quad (6)$$

and

$$\begin{aligned} UCL &= \bar{p} + 3\sqrt{\frac{\bar{p}(1-\bar{p})}{n}} = 0.1085 + 0.1320 \\ &= 0.2405 \end{aligned} \quad (7)$$

Because the LCL is computed to be a negative value, it will be set to zero on the chart.

Three consecutive subgroups (13, 14, and 15) are above the UCL suggesting excessive nonconforming items. It is discovered that these values occurred during a holiday when the production line was staffed with inexperienced workers. The three subgroups are removed from the calculation of the control limits,

2 Control Charts for Attributes

Table 1 Example data

Subgroup number, i	Number	
	Nonconforming units, X_i	Fraction nonconforming, p_i
1	8	0.16
2	3	0.06
3	1	0.02
4	5	0.10
5	3	0.06
6	3	0.06
7	6	0.12
8	3	0.06
9	4	0.08
10	5	0.10
11	3	0.06
12	5	0.10
13	13	0.26
14	15	0.30
15	17	0.34
16	2	0.04
17	7	0.14
18	4	0.08
19	1	0.02
20	6	0.12
21	1	0.02
22	5	0.10
23	6	0.12
24	6	0.12
25	8	0.16
26	3	0.06
27	5	0.10
28	7	0.14
29	3	0.06
30	3	0.06
31	7	0.14
32	5	0.10
33	6	0.12
179		0.1085

and the revised CL is found to be $\bar{p} = 0.0893$. The revised control limits are

$$LCL = \bar{p} - 3\sqrt{\frac{\bar{p}(1-\bar{p})}{n}} = 0.0893 - 0.1210 = -0.0316 \quad (8)$$

and

$$UCL = \bar{p} + 3\sqrt{\frac{\bar{p}(1-\bar{p})}{n}} = 0.0893 + 0.1210 = 0.2103 \quad (9)$$

Once again, the LCL is set to zero. The revised control chart is pictured in Figure 1. This chart contains 21 additional plotted subgroup values.

The limits of the chart in Figure 1 are based on 30 of the 33 initial subgroups of data. Following subgroup 33, a workforce retraining program was instituted. The final 21 subgroups are plotted to monitor the process improvement following the retraining. There appears to be a downward shift in the fraction of nonconforming items as a result of the training.

***c*-Chart**

The statistic plotted on a *c*-chart is the count of the number of nonconformities or defects over a specified unit. The specified unit may be a single item, a sample of multiple items, or a fixed time interval. The LCL, CL, and UCL of a *c*-chart are given by

$$LCL = \bar{c} - 3\sqrt{\bar{c}} \quad (10)$$

$$CL = \bar{c} \quad (11)$$

$$UCL = \bar{c} + 3\sqrt{\bar{c}} \quad (12)$$

where $\bar{c}$ is the average number of nonconformities observed in an initial in-control sample of m inspection units. Again, the limits are considered to be trial limits. If any of the initial counts plot outside of the limits and assignable causes are found, these units should be eliminated from the calculation of the control limits. Once the limits are established, future counts are plotted to monitor the process or product quality.

Other Attribute Charts

Other common attribute charts include the *np*-chart, which is used to monitor the number of nonconforming items and, like the *p*-chart, is based on the binomial model. The *u*-chart, similar to the *c*-chart, is based on the Poisson model, and can be used to monitor for the rate of occurrence of nonconformities. The *u*-chart is useful when the specified unit or window of opportunity for nonconformities varies over time. Additionally the *p*-chart can be adapted

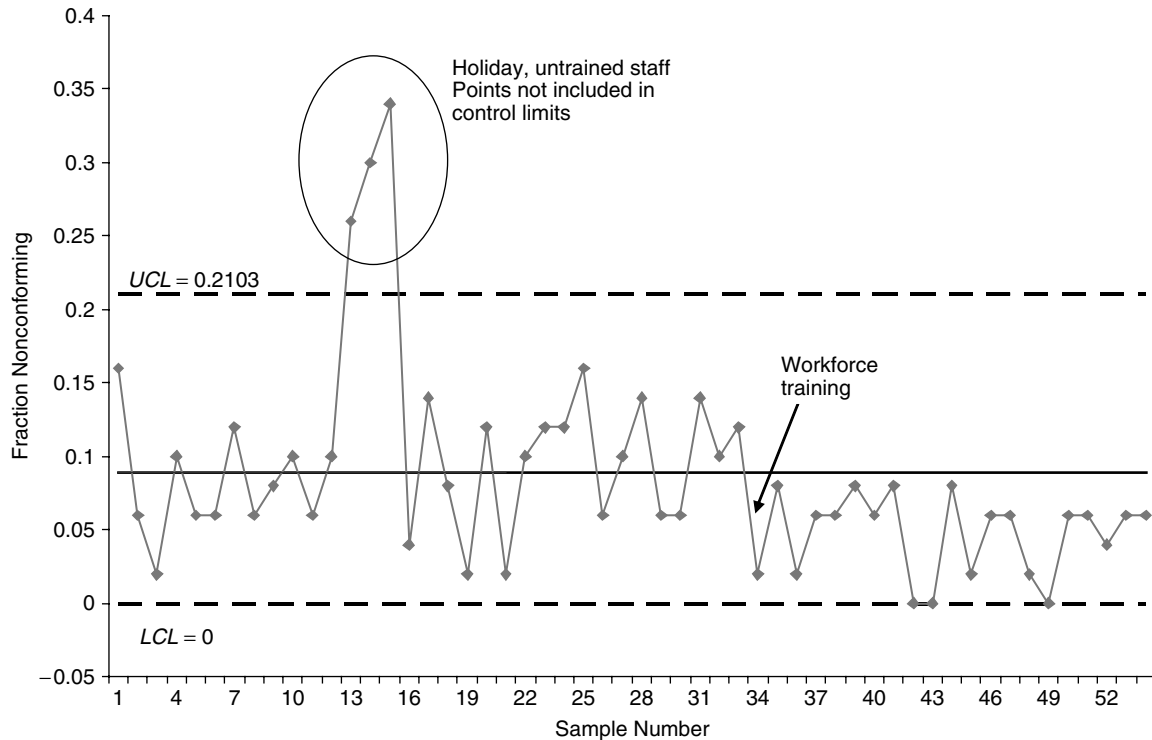


Figure 1 Revised control limits

for use when the subgroup size is not constant over time.

Adaptive sampling methods (*see Variable Sampling Rate Control Charts*), cumulative sum (CUSUM) (*see Cumulative Sum (CUSUM) Chart*), and exponentially weighted moving average (EWMA) (*see Exponentially Weighted Moving Average (EWMA) Control Chart*) charts have been recommended to improve the sensitivity of control charts for attributes. Ryan and Schwertman [1] discuss methods for finding optimal control limits for attribute charts. Woodall [2] provides a thorough bibliography of articles related to attribute control charts.

Merits and Limitations of Attribute Charts

It is commonly believed that variable control charts such as the $\bar{X}/R$ chart (*see Control Charts for the Mean; Moving Range and R Charts*) are preferred over attribute control charts. The benefits of variable control charts over attribute charts include the ability to separately monitor the process mean and variance,

the ability to monitor specific quality characteristics rather than defective products or nonconformities, and the ability to indicate process problems prior to a result of defective products. Variable control charts also tend to provide more information regarding potential out-of-control conditions within the process. However, attribute control charts are useful when the process under study is a complex assembly or service operation where several types of problems might occur simultaneously. Several characteristics can be summarized into a single easy-to-understand chart, which makes attribute charts particularly useful for historical studies and presentations to management. Additionally, when measurement data cannot be obtained, statistical process control may still be possible with attribute charts.

References

- [1] Ryan, T.P. & Schwertman, N.C. (1997). Optimal limits for attributes control charts, *Journal of Quality Technology* **29**, 86–98.

- [2] Woodall, W.H. (1997). Control charts based on attribute data: bibliography and review *Journal of Quality Technology* **29**, 172–183.

Further Reading

Montgomery, D.C. (2005). *Introduction to Statistical Quality Control*, 5th Edition, John Wiley & Sons, New York.

Related Articles

Control Charts, Overview; Control Charts for the Mean; Cumulative Sum (CUSUM) Chart; Exponentially Weighted Moving Average (EWMA) Control Chart; Precontrol.

L. ALLISON JONES-FARMER

Multivariate Control Charts Overview

Introduction

Multivariate statistical process control (MSPC) is used to simultaneously monitor multiple measurements from a process. Similar to the case for a single variable, the objective is to detect assignable causes of variation. Multiple measurements are quite common in modern processes. Multiple process variables might be measured such as temperatures, pressures, material properties, concentrations, flow rates, valve settings, revolutions, voltages, and so on. Multiple measurements might also be derived from the product (such as numerous physical dimensions, thickness, surface finish, gloss, uniformity, and so on). In general cases there may be very different units of measurement. In special cases of MSPC, there might be multiple measurements in the same units, such as thickness measurements at several sites on a surface, or lengths of multiple vanes on a single part, or an absorption spectrum for a range of wavelengths of light. Common units often occur in **multiple stream processes**, and [1] provides a summary. Profile measurements often provide additional examples (*see Profile Monitoring*). Multiple measurements allow one to consider a number of interesting strategies to monitor the process.

In addition to the cases mentioned previously, multiple measurements arise in other ways. A summary statistic might be used to combine data. For example, the mean thickness from several sites on a semiconductor wafer can be used to summarize multiple measurements to one. Consequently, the summary can be monitored with a **control chart** for a single variable. If some sites were near the center of the wafer and some were near the edge, one might also consider additional summaries such as the mean thickness at the edge minus the mean at the center. This leads to another summary, and another variable to be monitored. Consequently, multivariate data occurs again, although the number of metrics might be reduced from every site to fewer summary scores. Still, with more than one summary score to monitor, the application is again within the area of MSPC.

In a batch process one might monitor the temperature over time as the batch starts and then ends. This becomes more interesting when one considers that batches should also be monitored over time. From the view of the process, each batch provides potentially a large number of temperature measurements. Furthermore, a single temperature is probably supplemented with other temperature sensors, and other variable measurements. There may be several variables measured over time for every batch, and then the batches are monitored in sequence. Consequently, a number of alternative designs are feasible for a process monitor, but clearly multiple measurements are involved (*see Control Charts for Batch Processes*).

In yet another application for MSPC, the (rpm) revolutions per minute of a pump is often monitored with the pressure change across the pump. One expects an increase in rpm to correspond to an increased pressure differential. Although there are two variables, (the rpm might be considered an input variable and the pressure change an output), the variables do not have the same symmetric relationship that occurs with measurements from different sites on a wafer. Consequently, the nature of a strategy to monitor this data might incorporate the input–output structure. Still, multiple measurements are involved. A regression-type approach for MSPC has been applied, with and without, a direct cause and effect relationship among the variables (*see Regression Control Charts*).

As a transaction progresses through a business process, intermediate cycle times become available. The set of cycle times provides multiple measurements for each transaction.

Advantages of MSPC

With multiple measurements, each can be monitored in its own control chart. However, this has two disadvantages. One is that it is difficult to control the number of false alarms. An out-of-control action plan (OCAP) needs to begin whenever any control chart for a process signals. With multiple measurements and a correspondingly large number of control charts, there is much greater opportunity for false alarms. If the measurements plotted on each chart are independent, the increase in false alarms can be computed from basic probability. However, multiple

2 Multivariate Control Charts Overview

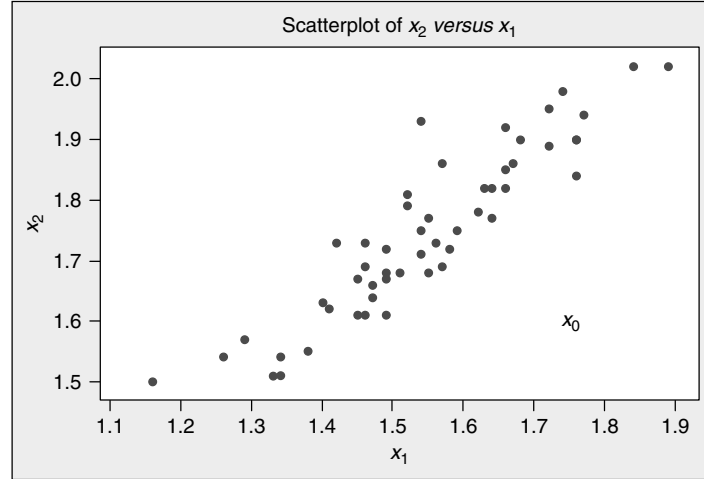


Figure 1 A scatter plot of two correlated measurements x_1 and x_2 displays an elliptical shape, but one point in the lower right does not conform to the pattern

measurements are often correlated and this makes the calculation more difficult. In any case, the control limits on each chart need to be widened to control the false alarms and this may degrade the detection of assignable causes.

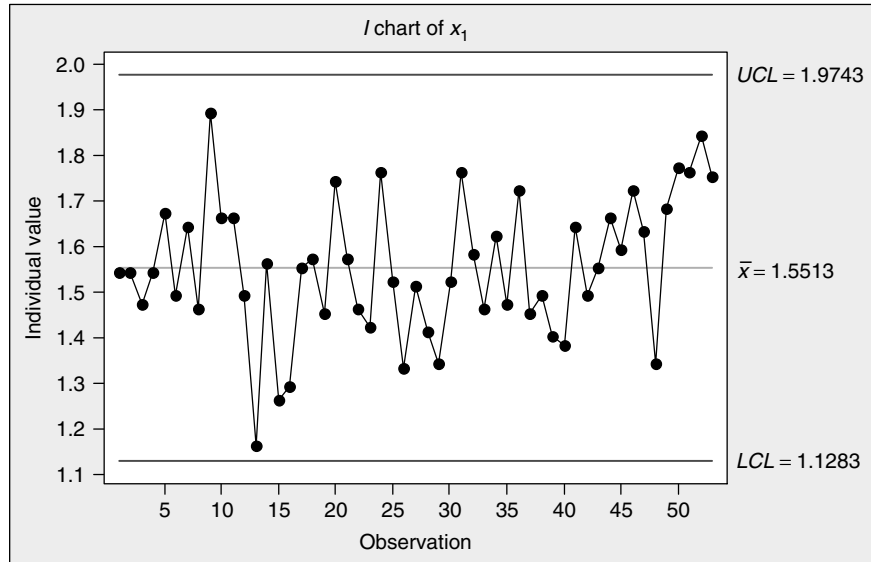
The second disadvantage of separate control charts is more important. There are often important relationships between variables that should be considered for MSPC. Multiple measurements from a process are often redundant, and correlated, and reflect data from many fewer true process states. Figure 1 provides a scatter plot of two measurements from a process. Notice the data is not randomly dispersed in the plot. Instead, the data clusters into an ellipse shape (except for a single point denoted as x_0). This is expected from measurements that are correlated. In Figure 1, the single point x_0 is clearly seen as unusual. Figure 2 shows that neither of the separate control charts for x_1 and x_2 indicate any problems. If x_1 and x_2 are measurements from two sites on a wafer, there is an unusual relationship observed for the wafer corresponding to point x_0 . As another example, if x_1 is the rpm of a pump and x_2 is the pressure differential, point x_0 suggests a broken blade (increased rpm without a corresponding pressure change). An unusual data point potentially signifies that an **assignable cause** is present. Therefore, MSPC has detected unusual process results that neither separate control chart can identify.

Similar relationships occur between variables in higher dimensions. Although simple graphics cannot be viewed, the shape can be confirmed from numerical calculations. Consequently, MSPC can be used to monitor for unusual points in multiple dimensions and this provides a powerful methodology. Data points that are distant from the typical, in-control data in multiple dimensions signal potential assignable causes.

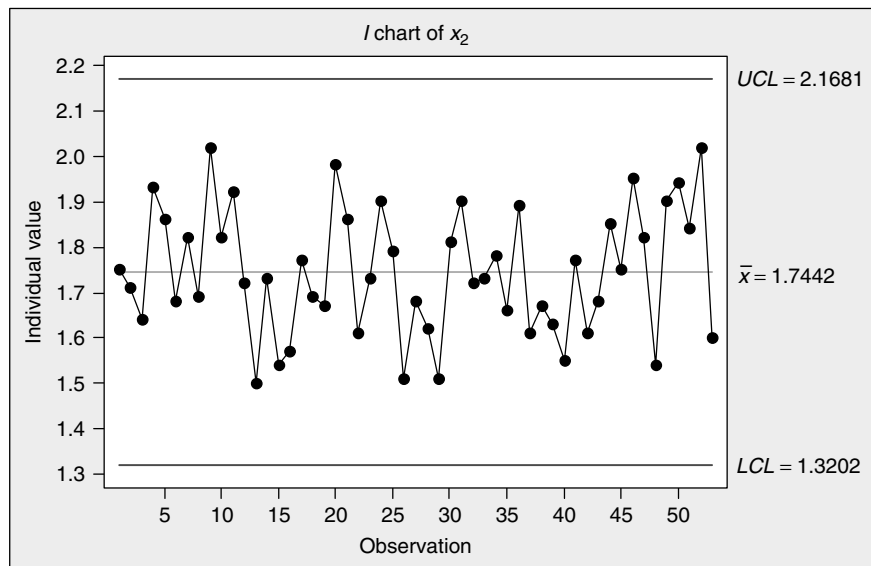
The data points tend to form an elliptical shape, and the major axis of the ellipse can be described by a line. The data points cluster along this line due to the **correlation**. In higher dimensions, data often clusters into lower-dimensional subspaces such as a plane in three dimensions (or a hyperplane in higher dimensions). An MSPC strategy can identify these subspaces and build monitors that focus on the subspace or deviations from the subspace.

Structure of the Data

MSPC has traditionally assumed that the measurements are available in the form of a table of numbers. Each row represents a time instant, and the numbers in the row are the measurements obtained from the process at that time (temperatures, pressures, etc.). Consequently, categorical measurements such as *{poor, average, good}* need to be quantified. However, additional work might be needed to acquire



(a)



(b)

Figure 2 Individual control charts for each variable x_1 and x_2 do not show any assignable causes. Can you identify the control chart points that correspond to the point in the lower right of Figure 1?

this convenient table. For example, the cycle times in the business process described previously will require some work to fit this table, and several alternatives might be useful (based on the objectives of the analysis). Similarly, the final quality measurement of a

product produced now is not dependent on the particle size distribution of the raw material now, but on an earlier, tracked lot of raw material that was used for the current product. Furthermore, many products might have been produced from the same lot of raw

material. Laboratory measurements might be taken from only a sample of products. This leads to different frequencies and phases of measurements that probably influence how to construct the table. Also, a particle size distribution may need to be summarized to a few numbers. Consequently, the traditional table might require substantial work as material flows from a start to an end of a process and measurements are derived in between. The batch process described previously presents additional challenges with time of batch, measurement time within the batch, and multiple variables. The design alternatives are driven by the objectives of the analysis.

Also, the numbers in the table are assumed to not be missing. If so, an interpolation or estimation might be needed. The statistical literature has considered several options for measurement **imputation**. The alternative is to eliminate rows that are not complete. However, with different frequencies of measurement this is not an attractive option. Large numbers of rows might be forced to be discarded. More modern, computationally intensive alternatives exist, but these are beyond the scope of this article.

Assignable Causes

MSPC has focused on control procedures to monitor process parameters. The parameters are the natural extensions from the case of a single variable. There is a mean μ_j and variance σ_j^2 for every variable X_j , say for $j = 1, 2, \dots, p$. Furthermore, MSPC considers the collection of covariances σ_{jk} between variables X_j and X_k . Consequently, there are control charts to monitor means the μ_j 's and those for the variances and covariances.

Usually the variances and covariances are organized into a table called the **covariance matrix** as

$$\begin{pmatrix} \sigma_1^2 & \sigma_{12} & \dots & \sigma_{1p} \\ \sigma_{21} & \sigma_2^2 & \dots & \sigma_{2p} \\ \dots & \dots & \dots & \dots \\ \sigma_{p1} & \sigma_{p2} & \dots & \sigma_p^2 \end{pmatrix}$$

and the objective is to monitor for an assignable cause that changes any one of these parameters. The multiple parameters dramatically increase the challenge from the case of a simple range chart. The article **Multivariate Charts for Variability** presents MSPC for such a covariance matrix.

The means of the variables are also assembled into a mean vector,

$$\mu = (\mu_1, \mu_2, \dots, \mu_p)$$

and monitors for a change to *any* element of this vector are of substantial interest. As shown in the previous figures an integrated strategy that uses the data from all variables can detect assignable causes not visible to separate control charts.

MSPC has focused on a change to any or all elements of the mean vector. Such an omnibus strategy can simultaneously monitor the means of all process variables. However, the chapters here for MSPC show that the omnibus strategy suffers as the number of process variables (denoted p) increases. Consequently, there is also interest in monitors for more restricted assignable causes.

As mentioned previously, high-dimensional data often cluster into lower-dimensional subspaces because of the correlations between the variables. Consequently, assignable causes that preserve the correlation are expected to shift the mean vector within this subspace. Alternatively, an assignable cause that specifically changes the relationships between the variables, might be expected to shift the mean vector from the subspace. When a pump blade breaks the relation between rpm and pressure differential is no longer valid. In either case, these are restrictions on the mean vector under an assignable cause, and such knowledge can be used to improve the performance of MSPC. Principal components analysis (PCA) has been widely used in this context. Much of the PCA work was summarized by Jackson [2], and further development of dimensional reduction currently continues. A more general subspace was considered by Runger [3]. For related examples, *see* **Control Charts for Batch Processes**.

One-sided control charts are common for a single variable. Therefore, the same motivation applies to multivariate data. Interest might be only in an increase in the mean of one or more variables. It is surprisingly more difficult to accommodate such simple information into a multivariate control chart, but solutions have been developed [4].

Nonparametric approaches have also been applied to the MSPC problem. For an introduction, *see* **Control Charts, Nonparametric**. Nonparametric density estimates and monitors to detect changes in these have been considered. Computationally intensive alternatives such as neural networks (*see* **Neural**

Networks: Construction and Evaluation) and other methods have also been considered, but these extensions are beyond the scope of this discussion.

Multivariate Control Charts

Common control charts for MSPC parallel the charts for a single variable. Consequently, there are multivariate extensions to Shewhart, **exponentially weighted moving average** (EWMA), and cumulative sum (CUSUM) control charts. The chapters **Hotelling's T^2 Chart**, **Multivariate Exponentially Weighted Moving Average (MEWMA) Control Chart**, and **Multivariate Cumulative Sum (CUSUM) Chart** describe each of these charts in more detail. The charts apply to a data row, or vector, of arbitrary length p . Consequently, the charts can be applied either before or after a dimensional reduction is used with the data. Suitably modified charts can monitor the mean vector or covariance matrix of the process. The charts are developed with a phase I and phase II strategy identical to the one used for a single variable (see **Control Charts, Overview**). Trial control limits might need to be modified on the basis of assignable causes. Rational subgrouping principles apply, although individual observations are typically used.

MSPC charts are designed for numerical data, and for processes that are characterized in terms of the mean vector and covariance matrix (called *second-order moments*). Although the charts are robust to nonnormality, the process measurements are expected to form elliptical regions, with greater density at the centers of the ellipse. The distance from an observed data vector x_i (or a appropriate sum or average) to the historical process mean vector μ is an important component of these charts. Rather than Euclidean distance, Mahalanobis distance is used

$$(x_i - \mu)' \Sigma^{-1} (x_i - \mu)$$

where Σ denoted the covariance matrix of the variables. Of course, the unknown parameters in the distance calculation need to be estimated in the phase I analysis. From basic algebra it can be shown that points with equal Mahalanobis distance from μ fall on the surface of an ellipse centered at μ . Consequently, a point with large Mahalanobis distance does not follow the expected elliptical pattern as described in Figure 1.

These charts have been effective for measurements from numerous processes in diverse industries. A further assumption is that the measurement vectors from different sample times are independent. If violated, this leads to the case of autocorrelated data, and suitably modified charts may be needed (see **Autocorrelated Data; Large Autoregressive Integrated Moving Average (ARIMA) Modeling**).

Signal Interpretation

One might presume that a signal from a control chart for a single variable is easier to analyze than MSPC chart. However, it is the process, not a variable, that must be interrogated as part of an OCAP. For example, if the length of an injection-molded part exceeds its control limit, there is not a single dial for part length. Instead, there is a combination of settings for temperatures, pressures, holding times, and so on that might need to be adjusted. The same is true for a multivariate control chart signal – interpretation of the signal and corrective action is not any different. Furthermore, in the previous pump example, a signal is not “caused” by either rpm or pressure differential. It might be caused by a broken blade, or a clog in a pipe. Consequently, the cause is in the process, not from a variable.

Still, in a multivariate control chart, one might eliminate some variables and focus on others that contribute to the signal. This can help corrective actions. If rpm and pressure differential were only two of several variable monitored by MSPC, then the change in the relationship between these two would no doubt benefit diagnosis and OCAP. Methods for identifying contributors to an MSPC signal have received substantial attention (see **Multivariate Control Charts, Interpretation of**).

Conclusions

With modern measurement automation, the importance of MSPC increases. It is now unusual to be concerned with a single measurement. And the ubiquitous computational resources have eliminated the traditional burden of onerous calculations. MSPC is an important tool for modern processes and a rich area for further innovations. Other chapters provide details on many of the topics mentioned here.

Acknowledgments

This material is based upon work supported by the National Science Foundation under Grant No. 0355575.

References

- [1] Mortell, R. & Runger, G.C. (1995). Process control of multiple stream processes, *Journal of Quality Technology* **27**(1), 1–12.
- [2] Jackson, J.E. (1991). *A User's Guide to Principal Components*, John Wiley & Sons, New York.
- [3] Runger, G.C. (1996). Projections and the U^2 control chart for multivariate statistical process control, *Journal of Quality Technology* **28**(3), 313–319.
- [4] Testik, M. & Runger, G.C. (2005). One-sided multivariate control charts, *IIE Transactions* **38**(8), 635–645.

GEORGE C. RUNGER

Hotelling's T^2 Chart

Introduction

Statistical process control (monitoring) has a purpose of quickly detecting instances when a process is *out of control* (OC), as indicated by a shift in the statistical distribution of one (univariate) or multiple (multivariate) measures of quality. The Hotelling's T^2 **control chart** is a tool for simultaneously monitoring the mean of several quality characteristics with a single control chart. This method is particularly well suited for monitoring several correlated quality characteristics (see **Multivariate Control Charts Overview**).

One classification of the various applications depends on the knowledge of the *in-control* (IC) parameters, with *phase II* denoting use of known parameters and two categories of *phase I* denoting imperfectly known parameters [1]. *Stage I* of phase I is a retrospective analysis of historical observations with the simultaneous objectives of estimating the IC distribution and detecting observations likely to have been from a different statistical distribution. *Stage II* of phase I is a prospective, or real time, detection of OC conditions as new observations are obtained, based on the distribution estimated in stage I. Logically the parameter estimates from stage I would be updated at some point after accumulating additional IC history in stage II, but when to do this and what its statistical effects are have not been published. Some authors use the terminology “phase I” (or phase 1) for the retrospective analysis and “phase II” (or phase 2) for the prospective, real-time monitoring with estimated parameters, including the case of perfect estimation of phase II. There are alternatives to the dichotomy of retrospective and prospective analysis in the form of *self-starting* control charts that can begin the real-time monitoring process with very few observations, then simultaneously signal the detection of sustained shifts and smoothly update the parameter estimates [2]. With the *change-point* (see **Change-Point Methods**) paradigm, the entire history of observations (or a large number of the most recent ones to limit computational requirements) is evaluated with each new observation to detect a boundary separating the history into two groups with statistically significantly different mean vectors. An OC signal is given when the test statistic exceeds a

control limit, which may vary with the number of observations [3, 4].

Observations are classified into *rational subgroups* of n observations when the n observations are acquired at nearly the same instant of time so that it is likely that they all have the same statistical distribution. The term *individual observations* applies when the subgroup size is $n = 1$, a situation for which the task of accurately estimating the IC **covariance** matrix in stage I is greatly complicated. With individual observations, the subgroup mean is the single observation.

Hotelling's T^2 Statistic and its Distribution

A basic measure of the statistical (Mahalanobis) distance between two multivariate vectors $\mathbf{x}_1$ and $\mathbf{x}_2$ is $\sqrt{(\mathbf{x}_1 - \mathbf{x}_2)^T \mathbf{\Sigma}^{-1} (\mathbf{x}_1 - \mathbf{x}_2)}$, where the matrix $\mathbf{\Sigma}$ is positive definite, ensuring that the radicand is never negative and is zero only when the two vectors are the same. As the name indicates, the Hotelling's T^2 chart statistic is the square of such a distance, based on the mean vector of a subgroup and the IC mean vector. When specialized to univariate observations with the variance estimated, Hotelling's T^2 is the square of a Student's t statistic. Hotelling's T^2 chart is called a *multivariate Shewhart chart*. Although a univariate Shewhart chart has upper and lower control limits, they are symmetrical about the process mean, so that a control chart signal can be uniquely related to the distance of a subgroup mean from the IC process mean, as with Hotelling's T^2 chart. Professor Harold Hotelling introduced the multivariate generalization and gave its distribution [5], also noting its invariance to any linear transformation of the measurements $\mathbf{y} = \mathbf{A}\mathbf{x} + \mathbf{b}$, where $\mathbf{A}$ is a nonsingular matrix. A common early reference for Professor Hotelling's work is an article that appeared later and gave a specialized bivariate application [6], but that work lacks a matrix formulation and generalization beyond two dimensions.

The classic form for Hotelling's T^2 statistic [5] applies in stage II, with the formula $T_s^2 = n(\bar{\mathbf{x}}_s - \bar{\bar{\mathbf{x}}})^T \bar{\mathbf{S}}^{-1} (\bar{\mathbf{x}}_s - \bar{\bar{\mathbf{x}}})$, which is also given in [7, p. 216] and [8, p. 366], where $\bar{\mathbf{x}}_i = n^{-1} \sum_{j=1}^n \mathbf{x}_{i,j}$ is the mean vector of subgroup i and $\bar{\bar{\mathbf{x}}}$ denotes the mean vector of a subgroup “subsequent” to (and not included in) the estimation sample of m subgroups. The grand mean vector is $\bar{\bar{\mathbf{x}}} = m^{-1} \sum_{i=1}^m \bar{\mathbf{x}}_i$, and

$\bar{\mathbf{S}} = m^{-1} \sum_{i=1}^m \bar{\mathbf{S}}_i$ is the mean of the subgroup covariance matrices, $\bar{\mathbf{S}}_i = (n-1)^{-1} \sum_{j=1}^n (\mathbf{x}_{i,j} - \bar{\mathbf{x}}_i)(\mathbf{x}_{i,j} - \bar{\mathbf{x}}_i)^T$. Wierda [9, p. 151] gives the distribution of T_s^2 as $(p(m+1)(n-1)/m(n-1) - p + 1)F_{p, m(n-1)-p+1; \tau^2}$, where p is the number of elements (quality measurements) in each measurement vector $\mathbf{x}_{i,j}$, and $\tau^2 = (mn/m+1)(\boldsymbol{\mu}_s - \boldsymbol{\mu}_0)^T \boldsymbol{\Sigma}^{-1}(\boldsymbol{\mu}_s - \boldsymbol{\mu}_0)$ is the noncentrality parameter based on the true mean vector of the subsequent subgroup $\boldsymbol{\mu}_s$, the true IC mean vector $\boldsymbol{\mu}_0$, and the true covariance matrix for single observations $\boldsymbol{\Sigma}$. The OC properties depend on a shift in the mean vector only through the noncentrality parameter, so that shifts having the same value of the noncentrality parameter result in T^2 statistics having the same statistical distributions.

The *average run length* (ARL) (see **Average Run Lengths and Operating Characteristic Curves**) is the average number of observations until a signal, with performance typically specified by an IC ARL₀ and an OC ARL₁ that is based on some specific shift in the mean vector. The *upper control limit* (UCL) for a specified ARL₀ is $(p(m+1)(n-1)/m(n-1) - p + 1)F_{\alpha; p, m(n-1)-p+1; 0}$, where $\alpha = 1/\text{ARL}_0$ and $F_{\alpha; p, m(n-1)-p+1; 0}$ is the $(1-\alpha)$ quantile (which is the **100(1- α) percentile**) of the (central) F distribution with numerator **degrees of freedom** p and denominator degrees of freedom $m(n-1) - p + 1$.

Numerical Example

As an example of the calculations, suppose that each subgroup has $n = 5$ observations consisting of measurement vectors of $p = 2$ elements and ARL₀ is 370.4, so that α is 0.0027. With an IC estimation sample of $m = 50$ observations, the UCL is $(2(50+1)(5-1)/50(5-1) - 2 + 1)F_{0.0027; 2, 50(5-1)-2+1; 0} = (392/199)6.094 = 12$. Suppose the estimated parameters are means of 55.5 and 33.3, variances of 4.68 and 2.76, and a covariance of 2.52. With a new subgroup mean vector of 53.45 and 33.6, the deviations from the overall means are -2.05 and 0.3. The inverse of the covariance matrix is $\begin{bmatrix} 0.420 & -0.384 \\ -0.384 & 0.713 \end{bmatrix}$. The first product of the quadratic form is $[-2.05 \quad 0.30] \begin{bmatrix} 0.420 & -0.384 \\ -0.384 & 0.713 \end{bmatrix} = [-0.977 \quad 1.001]$, so that the T^2 statistic for this subgroup is $5[-0.977 \quad 1.001] \begin{bmatrix} -2.05 \\ 0.30 \end{bmatrix} = 11.51$.

Since this statistic is less than the UCL, no signal is given for this subgroup. If the OC distribution had a shift in the mean vector of, say $[1 \quad 0.5]$, and the true covariance matrix was the same as the estimate from the sample, then the noncentrality parameter would be $\tau^2 = (50(5)/50+1)[1 \quad 0.5] \begin{bmatrix} 0.420 & -0.384 \\ -0.384 & 0.713 \end{bmatrix} \begin{bmatrix} 1 \\ 0.5 \end{bmatrix} = 4.902[0.228 \quad -0.027] \begin{bmatrix} 1 \\ 0.5 \end{bmatrix} = 1.053$. The probability a chart statistic would exceed the UCL is the same as that of a random draw from the F distribution exceeding $(m(n-1) - p + 1/p(m+1)(n-1))\text{UCL} = (50(5-1) - 2 + 1/2(50+1)(5-1))12 = 6.094$, which is 0.01522, so the OC ARL₁ would be 65.7. Additional examples are given in [10–12].

General Form and Form for Phase II

The general case of the T^2 distribution is given by [13, p. 74] as the distribution of the statistic $Q = \mathbf{y}^T(\mathbf{W}/f)^{-1}\mathbf{y} \sim T_{p,f}^2$, where $\mathbf{y}$ has a multivariate **normal distribution** with zero mean vector, $\mathbf{W}$ is independently distributed as a Wishart with f degrees of freedom, and a common covariance matrix is shared by $\mathbf{y}$ and $\mathbf{W}$, of which $\mathbf{W}/f$ is an unbiased estimator. Equivalently, $(f - p + 1/fp)Q \sim F_{p, f-p+1; 0}$. Since $\bar{\mathbf{x}}_s$ requires division of a sum of random vectors by n and $\bar{\bar{\mathbf{x}}}$ requires division of an independent sum by mn , the expression $(\bar{\mathbf{x}}_s - \bar{\bar{\mathbf{x}}})\sqrt{mn/(m+1)}$ has zero mean and the same covariance matrix as a single observation. This gives the general form $T_s^2 = n(\bar{\mathbf{x}}_s - \bar{\bar{\mathbf{x}}})^T(\mathbf{W}/f)^{-1}(\bar{\mathbf{x}}_s - \bar{\bar{\mathbf{x}}}) \sim (m+1/m)(fp/f - p + 1)F_{p, f-p+1}$.

Instead of using $\bar{\mathbf{S}}$, which has $m \times (n-1)$ degrees of freedom, the overall sample covariance matrix $\mathbf{S} = (1/mn - 1) \sum_{i=1}^m \sum_{j=1}^n (\mathbf{x}_{i,j} - \bar{\bar{\mathbf{x}}})(\mathbf{x}_{i,j} - \bar{\bar{\mathbf{x}}})^T$, which ignores the subgrouping and has degrees of freedom $mn - 1$, can be used. With this variation the chart statistic is $T_{s,p}^2 = n(\bar{\mathbf{x}}_s - \bar{\bar{\mathbf{x}}})^T \mathbf{S}^{-1}(\bar{\mathbf{x}}_s - \bar{\bar{\mathbf{x}}})$, with UCL given by $(m+1/m)((mn-1)p/(mn-p))F_{\alpha; p, mn-p; 0}$ in [9, p. 151].

In phase II, when the true parameters are known, Hotelling's T^2 chart is called a χ^2 chart with chart statistic given by $\chi_s^2 = n(\bar{\mathbf{x}}_s - \boldsymbol{\mu}_0)^T \boldsymbol{\Sigma}^{-1}(\bar{\mathbf{x}}_s - \boldsymbol{\mu}_0)$ and UCL = $\chi_{\alpha; p}^2$. As the size of the estimation sample increases, the stage II UCL smoothly decreases and approaches the UCL for phase II. For example, with $n = 5$, $p = 2$, and $\alpha = 0.0027$, the UCL is 12.49 with

$m = 50$ observations, decreasing to 11.89 with 500 observations and decreasing to the χ^2 UCL of 11.83 with “known” parameters.

Statistical Formula for Stage I

In a stage I chart, a subgroup mean $\bar{\mathbf{x}}_i$ from the estimation sample is used to calculate the chart statistic $T_i^2 = n(\bar{\mathbf{x}}_i - \bar{\mathbf{x}})^T \bar{\mathbf{S}}^{-1}(\bar{\mathbf{x}}_i - \bar{\mathbf{x}})$, with UCL given by $(p(m-1)(n-1)/m(n-1) - p + 1)F_{\alpha; p, m(n-1)-p+1; 0}$. All m observations are analyzed at once, so that the concept of specifying an ARL is not meaningful. Instead, performance is described in terms of signal probabilities. The probability of a signal when all observations are from an IC process is the *false signal probability (fsp)* or false alarm probability or similar terminology. Thus, α is the *fsp* for a single observation, often taken to be 0.0027.

Stage I objectives include estimation of the IC parameters for future monitoring in stage II, and OC observations are typically excluded from the estimation set, reducing m accordingly. Special care and attention is required in identifying *all* the OC observations when there is a sustained shift in the mean vector, as all of the affected subgroups *may not be associated with chart statistics that exceed the UCL*. This is so partially because the IC mean vector is not known and is estimated from the observations being analyzed. Another difficulty is that only distance is presented by the chart statistics, without indication of direction. Consider the case of, say, 25 IC observations followed by a set of 25 OC observations with a common mean vector that is shifted from the IC mean vector. The grand mean vector will be the midpoint between the mean vectors of each set. Thus, the group means would be equally distant from the grand mean and would have the same T^2 distance from it (although these distances would not appear on the control chart). The chart statistics for the 50 subgroups tend to have similar distances as well, so that in the plot of the distances (squared), a visually discernable break between the two sets may be lacking. Useful directional information can be obtained by projecting the subgroup mean vectors onto the line joining the mean vectors of the sets and applying a univariate Shewhart chart to the result. This will make it visually evident that one set clusters in a negative direction while the other clusters in a positive direction,

although equally distant from the grand mean. Simple visual interpretation of a T^2 control chart can easily lead to misinterpretation when there is a sustained shift in the mean vector. Especially when only one or a few statistics exceed the UCL it is easy to classify incorrectly the associated subgroups as **outliers** and remove only those subgroups from the estimation sample, because of the visual similarity to the chart pattern when only outlying observations are present. For more discussion on this point (see [14]; **Multivariate Control Charts, Interpretation of**).

Special Issues with Individual Observations in Stage I

In stage I with individual observations ($n = 1$), a covariance matrix cannot be calculated from the single observation in each subgroup, so $\bar{\mathbf{S}}$ cannot be used. Although there are compelling reasons for not doing so, as will be discussed later, there are recommendations to use the sample covariance matrix $\mathbf{S} = (1/(m-1)) \sum_{i=1}^m (\mathbf{x}_i - \bar{\mathbf{x}})(\mathbf{x}_i - \bar{\mathbf{x}})^T$, which is unbiased *only* in the IC condition when all observations are from a common distribution. Tracy *et al.* [15] show that when the process is IC, $T_{i,1}^2 = (\mathbf{x}_i - \bar{\mathbf{x}})^T \mathbf{S}^{-1}(\mathbf{x}_i - \bar{\mathbf{x}})$ has the distribution $((m-1)^2/m)\beta_{p/2, (m-p-1)/2}$, so the UCL would be $((m-1)^2/m)\beta_{\alpha; p/2, (m-p-1)/2}$.

Since the observations are correlated slightly the overall *fsp* is not exactly $1 - (1 - \alpha)^m$ but is close to that for reasonably large values of m . In stage II with $n = 1$ it is reasonable to use $\mathbf{S}$, and the corresponding chart statistic is $T_{s,1}^2 = (\mathbf{x}_s - \bar{\mathbf{x}})^T \mathbf{S}^{-1}(\mathbf{x}_s - \bar{\mathbf{x}})$ with UCL $= (p(m+1)(m-1)/m(m-p))F_{\alpha; p, m-p; 0}$, as given in [7, p. 225].

With larger subgroup size n , the distribution of a subgroup average vector approximates the multivariate normal distribution, although the covariance estimator may be affected by nonnormality. Other multivariate control charts, such as multivariate **exponentially weighted moving average (MEWMA)** (see **Multivariate Exponentially Weighted Moving Average (MEWMA) Control Chart**) and multivariate cumulative sum (**CUSUM**) (see **Multivariate Cumulative Sum (CUSUM) Chart**), can be designed so as to preserve the control chart performance with respect to nonnormal distributions even with individual observations or small subgroup size [16].

The T^2 statistic takes into account the IC correlations among the elements of the measurement vectors

in contrast to an alternative monitoring technique of using univariate control charts applied to each measurement element individually. Thus, a subgroup may have a large T^2 statistic, indicating its mean vector is unusually distant from the IC mean vector with respect to the IC correlations, although none of the individual means is unusually distant from its IC mean, which is generally considered an advantage. It is mathematically possible to standardize the subgroup mean vectors as $\bar{\mathbf{z}}_s = \Sigma^{-1/2}(\bar{\mathbf{x}}_s - \boldsymbol{\mu}_0)$ so that the T^2 statistic for $\bar{\mathbf{z}}_s$ is the simple sum of squares of its elements and each element can be monitored separately with the same performance characteristics maintained.

When a subgroup leads to a signal, it may be helpful to diagnose which element or combination of elements most contributed to the signal, which can be estimated using a *step-down* method as described in [17, 18] (see **Multivariate Control Charts, Interpretation of**). Alternatively, *principal components* can be used [19], but users may have difficulty interpreting principal components and relating a signal for a principal component with a specific special cause of variation. A criticism of monitoring only a few principal components is that they are determined by the IC variation among the elements of a measurement, which may not be the same subspace as shifts in the mean vector with the presence of special causes.

In stage I with individual observations, the sample covariance matrix $\mathbf{S}$ is not unbiased unless all observations are from the same distribution. With numerous IC and OC observations present, the true covariance matrix may be estimated so badly that the T^2 chart fails to signal, a problem referred to as *masking*. For example [20] shows that when $\mathbf{S}$ is used the signal probability actually *decreases* with an increasing size of a step shift in the mean vector when the shift is near the middle of the observations. A simple estimator that is robust to a small number of shifts in the mean vector is based on the vector difference between successive observations. Sullivan and Woodall [20] gave some properties of the T^2 chart based on this estimator, which is not effective in detecting a large number of outliers, and [21] gives formulas for calculation of the control limits. Other estimators such as the *minimum volume ellipsoid* [22] and an estimator given in [20] are useful in detecting isolated outliers but less effective in detecting multiple sustained shifts in the mean vector. Another proposal for detecting multiple shifts and/or multiple

outliers in the mean vector and/or covariance matrix is given by Sullivan and Woodall [23].

References

- [1] Alt, F.B. & Smith, N.D. (1988). Multivariate process control, in *Handbook of Statistics*, P.R. Krishnaiah & C.R. Rao, eds, Elsevier Science Publishers B.V., Amsterdam, Vol. 7, pp. 333–351.
- [2] Sullivan, J.H. & Jones, L.A. (2002). A self-starting control chart for multivariate individual observations, *Technometrics* **44**, 24–33.
- [3] Srivastava, M.S. & Worsley, K.J. (1986). Likelihood ratio tests for a change in the multivariate normal mean, *Journal of the American Statistical Association* **81**, 199–204.
- [4] Young, S.S. & Kim, S.W. (2005). Bayesian single change point detection in a sequence of multivariate normal observations, *Statistics* **39**, 373–387.
- [5] Hotelling, H.H. (1931). The generalization of student's ratio, *Annals of Mathematical Statistics* **2**, 360–378.
- [6] Hotelling, H.H. (1947). Multivariate quality control illustrated by the air testing of sample bombsights, in *Techniques of Statistical Analysis*, C. Eisenhart, M.W. Hastay & W.A. Wallis, eds, McGraw-Hill, New York, pp. 111–184.
- [7] Ryan, T.P. (1989). *Statistical Methods for Quality Improvement*, John Wiley & Sons, New York.
- [8] Montgomery, D.C. (1996). *Introduction to Statistical Quality Control*, 3rd Edition, John Wiley & Sons, New York.
- [9] Wierda, S.J. (1994). Multivariate statistical process control – recent results and directions for future research, *Statistica Neerlandica* **48**, 147–168.
- [10] Tracy, N.D., Young, J.C. & Mason, R.L. (1997). Some aspects of hotelling's T^2 statistic for multivariate quality control, in *Statistics of Quality*, S. Ghosh, W.R. Schucany & W.B. Smith, eds, Marcel Dekker, New York, pp. 77–100.
- [11] Mason, R.L. & Young, J.C. (2001). *Multivariate Statistical Process Control with Industrial Application*, Society for Industrial and Applied Mathematics, Philadelphia.
- [12] Fuchs, C. & Kenett, R.S. (1998). *Multivariate Quality Control Theory and Applications*, Marcel Dekker, New York.
- [13] Mardia, K.V., Kent, J.T. & Bibby, J.M. (1979). *Multivariate Analysis*, Academic Press, London.
- [14] Sullivan, J.H., Woodall, W.H. & Barrett, J.D. (1995). Interpreting the retrospective T^2 control chart, IAPQR transactions, *Journal of the Indian Association for Productivity, Quality, and Reliability* **20**, 183–195.
- [15] Tracy, N.D., Young, J.C. & Mason, R.L. (1992). Multivariate control charts for individual observations, *Journal of Quality Technology* **24**, 88–95.
- [16] Stoumbos, Z.G. & Sullivan, J.H. (2002). Robustness to non-normality of the multivariate EWMA control chart, *Journal of Quality Technology* **34**, 260–276.

-
- [17] Murphy, B.J. (1987). Selecting out of control variables with the T^2 multivariate quality control procedure, *The Statistician* **36**, 571–583.
 - [18] Mason, R.L., Tracy, N.D. & Young, J.C. (1995). Decomposition of T^2 for multivariate control chart interpretation, *Journal of Quality Technology* **27**, 99–108.
 - [19] Jackson, J.E. (1991). *A User's Guide to Principal Components*, Wiley-Interscience, New York.
 - [20] Sullivan, J.H. & Woodall, W.H. (1996). A comparison of multivariate control charts for individual observations, *Journal of Quality Technology* **28**, 398–408.
 - [21] Williams, J.D., Woodall, W.H., Birch, J.B. & Sullivan, J.H. (2006). On the distribution of Hotelling's T^2 statistic based on the successive differences covariance matrix estimator, *Journal of Quality Technology* **38**, 217–229.
 - [22] Vargas, N.J.A. (2003). Robust estimation in multivariate control charts for individual observations, *Journal of Quality Technology* **35**, 367–376.
 - [23] Sullivan, J.H. & Woodall, W.H. (2000). Change-point detection of mean vector or covariance matrix shifts using multivariate individual observations, *IIE Transactions* **32**, 537–549.

Related Articles

Control Charts, Selection of; Multivariate Control Charts Overview; Multivariate Exponentially Weighted Moving Average (MEWMA) Control Chart; Multivariate Cumulative Sum (CUSUM) Chart; Multivariate Control Charts, Interpretation of.

JOE H. SULLIVAN

Multivariate Exponentially Weighted Moving Average (MEWMA) Control Chart

Introduction

The multivariate **exponentially weighted moving average (MEWMA) control chart** is a tool for simultaneously monitoring several correlated process variables. Historically, univariate charts have been used to determine if a process is in a state of statistical control (*see Control Charts, Overview; Control Charts for the Mean*). If multiple variables need to be considered, then separate charts have been used for monitoring each individual variable. However, many variables that affect the quality of a process are often correlated. Using separate univariate charts assumes that the variables are independent and ignores the **correlation** among them. Multivariate control charts are designed to incorporate the correlation structure of the variables into the chart when determining the state of statistical control of a process (*see Multivariate Control Charts Overview*). The exponentially weighted moving average (EWMA) statistic computes a weighted average over time in an attempt to detect mean changes more quickly.

MEWMA Statistic

To detect small changes in a process mean vector, a commonly used chart is the MEWMA control chart. In the univariate case, the EWMA chart (*see Exponentially Weighted Moving Average (EWMA) Control Chart*) has been shown to provide quicker detection of small changes in the process mean than the traditional $\bar{x}$ chart (*see Control Charts for the Mean*). Extending this to the multivariate situation was proposed by Lowry *et al.* [1]. Given a set of vectors $X_1, X_2, \dots, X_n$ that follow a multivariate **normal distribution** of dimension p with mean vector μ and variance–**covariance** matrix Σ , the MEWMA statistic can be defined as in equation (1)

$$Z_i = R[X_i - \mu_0] + (I - R)Z_{i-1} \quad (1)$$

for $i = 1, 2, \dots, n$ where R is a diagonal matrix, $R = \text{diag}(r_1, \dots, r_p)$ with $0 < r_i \leq 1$, and Z_0 is usually initialized to the target mean vector, μ_0 . The elements of the matrix R may be chosen to reflect the relative importance of the recorded data. If all variables are equally important, then $r_1 = r_2 = \dots = r_p$. Using this choice for the matrix R yields $Z_i \sim N_p(0, \Sigma_z)$ where the variance–covariance matrix is given in equation (2).

$$\Sigma_{z_i} = \frac{r[1 - (1 - r)^{2i}]}{2 - r} \Sigma \quad (2)$$

Since the variance depends on the weights of the previous time periods it increases over time to its limiting value of $\frac{r}{2-r}\Sigma$. For values of $r \geq 0.5$, this limiting value is achieved rapidly. However, for values of $0 < r < 0.5$, it may take more than 25 time periods to get close to this limiting value. Using the exact variances in the series ensures that the probability of a false alarm is constant for all plotted statistics.

Plotted Statistic

The form of the statistic plotted on the control chart depends on whether Σ is known or not. If Σ is known then a χ^2 statistic may be computed as $Q_i^2 = Z_i' \Sigma_{Z_i}^{-1} Z_i$ with an upper control limit based on the χ_p^2 distribution with a false alarm rate of α . If Σ is not known, one can use the sample variance–covariance matrix (S) as an estimate of Σ (*see Multivariate Analysis*) and compute the Hotelling's T^2 value for each successive element of the series, Z_i . The Hotelling's T^2 value is computed as follows: $T_i^2 = Z_i' S_{Z_i}^{-1} Z_i$, where $S_{Z_i} = \frac{r[1 - (1 - r)^{2i}]}{2 - r} S$. The T^2 value can be plotted on a control chart with the upper control limit defined by $\frac{(n-1)p}{n-p} F_{(\alpha, p, n-p)}^{-1}$ where $F_{(\alpha, p, n-p)}^{-1}$ denotes the upper $100(1 - \alpha)$ th percentile of the F distribution with p and $n - p$ **degrees of freedom**. This is similar to the T^2 control chart (*see Hotelling's T^2 Chart*), but replaces observations at time t with the EWMA value computed for time t . In practice, one would prefer to use the χ^2 statistic if Σ is known, since it will result in a smaller **average run length** (ARL). For situations where there is not a great deal of historical data, the Hotelling's T^2 form should be used.

2 Multivariate EWMA Control Chart

Example

A simulated example is constructed with five variables ($p = 5$) and 35 observations. The simulated data represent measurements for five variables ($X_1 - X_5$). The data set contains 36 observations where the last observation, which is out of control, yields a signal. The first 35 data points were generated from an in-control process and used to estimate the sample variance-covariance matrix. Fixing the target mean to be $\mu'_0 = (5, 5, 5, 5, 5)$ the MEWMA vectors were calculated using $r = 0.2$ for all variables. Table 1 contains the original data ($X_1 - X_5$), the MEWMA values ($Z_1 - Z_5$) and the resulting

T^2 statistic for each vector. These T^2 statistics are also plotted in Figure 1. The control limit is computed as $\frac{(n-1) \times p}{n-p} F_{(0.05, p, n-p)}^{-1} = \frac{34 \times 5}{30} 2.534 = 14.357$. In this example the first 35 observations are inside the control limit, but the 36th observation falls above the upper control limit. For this demonstration the false alarm rate was 0.05, which may yield too many false alarms. It is customary to use false alarm rates of 0.01 or less.

Related Topics

When a value is plotted above the upper control limit (UCL), it is called a *signal*. While a signal indicates a

Table 1 Data used for example MEWMA control chart

Index	X_1	X_2	X_3	X_4	X_5	Z_1	Z_2	Z_3	Z_4	Z_5	T^2
1	5.358	4.924	4.431	3.093	4.487	0.0716	-0.015	-0.114	-0.381	-0.1026	4.34248
2	2.148	4.156	4.001	5.865	4.371	-0.513	-0.181	-0.291	-0.132	-0.208	6.389396
3	4.161	3.618	3.836	4.505	6.543	-0.578	-0.421	-0.465	-0.205	0.142	8.366432
4	4.879	6.672	5.742	4.68	5.275	-0.487	-0.003	-0.224	-0.228	0.169	3.286738
5	2.612	7.032	4.242	5.184	6.012	-0.867	0.404	-0.331	-0.145	0.337	9.219914
6	5.208	5.251	5.575	5.477	5.847	-0.652	0.374	-0.150	-0.021	0.439	5.762943
7	5.875	5.32	6.257	3.191	3.045	-0.347	0.363	0.132	-0.379	-0.039	4.006224
8	3.884	6.257	4.894	3.95	5.082	-0.501	0.542	0.084	-0.513	-0.015	7.274752
9	5.753	5.027	5.112	6.295	5.67	-0.250	0.439	0.090	-0.151	0.122	2.379272
10	6.138	5.312	6.101	5.402	3.569	0.028	0.413	0.292	-0.041	-0.189	2.694534
11	4.728	4.547	5.035	5.266	6.13	-0.032	0.240	0.241	0.021	0.075	1.153995
12	6.336	4.227	6.07	4.775	5.962	0.241	0.038	0.406	-0.028	0.252	2.406579
13	4.035	4.256	5.876	5.387	5.51	0.000	-0.119	0.500	0.055	0.304	3.436191
14	4.707	5.394	5.214	5.427	3.586	-0.058	-0.016	0.443	0.129	-0.040	2.698012
15	4.78	6.347	5.311	5.011	5.448	-0.091	0.256	0.417	0.105	0.058	2.817496
16	3.1	5.306	3.938	4.102	4.949	-0.453	0.266	0.121	-0.095	0.036	2.737835
17	3.512	5.559	5.688	3.657	4.912	-0.660	0.325	0.234	-0.345	0.011	7.085626
18	3.636	3.778	3.853	5.002	3.468	-0.801	0.015	-0.042	-0.275	-0.297	8.250673
19	4.532	3.322	6.813	3.897	7.156	-0.734	-0.323	0.329	-0.441	0.193	9.862655
20	5.59	4.301	3.737	5.855	4.56	-0.469	-0.398	0.011	-0.182	0.067	3.65062
21	6.596	4.47	5.165	7.168	4.18	-0.056	-0.425	0.042	0.288	-0.111	2.063864
22	5.162	5.873	5.245	3.02	6.678	-0.013	-0.165	0.082	-0.165	0.247	0.897306
23	5.666	3.901	2.89	4.596	4.194	0.123	-0.352	-0.356	-0.213	0.036	3.229715
24	4.105	4.256	6.375	4.951	6.004	-0.080	-0.430	-0.010	-0.180	0.230	2.192427
25	5.486	4.584	4.564	2.917	4.481	0.033	-0.427	-0.095	-0.561	0.080	4.631125
26	4.12	4.309	5.689	5.217	5.058	-0.150	-0.480	0.062	-0.405	0.076	3.970696
27	5.257	6.448	4.698	3.678	6.807	-0.068	-0.095	-0.011	-0.589	0.422	3.986599
28	3.976	5.837	3.75	4.88	5.332	-0.260	0.092	-0.259	-0.495	0.404	4.263655
29	4.705	5.076	5.318	4.384	5.471	-0.267	0.089	-0.143	-0.519	0.417	3.960487
30	4.928	3.942	4.491	4.616	4.719	-0.228	-0.141	-0.217	-0.492	0.278	3.586314
31	3.667	2.771	5.078	5.735	4.877	-0.449	-0.558	-0.158	-0.247	0.198	5.02378
32	5.895	5.502	4.25	3.95	6.476	-0.180	-0.346	-0.276	-0.407	0.453	4.912568
33	4.839	7.25	4.816	6.68	3.481	-0.176	0.173	-0.258	0.010	0.059	1.176238
34	6.808	5.034	4.715	3.373	3.204	0.221	0.145	-0.263	-0.317	-0.312	2.815775
35	5.610	6.081	6.005	3.898	4.175	0.299	0.332	-0.010	-0.474	-0.415	4.350111
36	4.777	5.939	5.243	11.000	11.000	0.194	0.454	0.041	0.821	0.868	17.82526

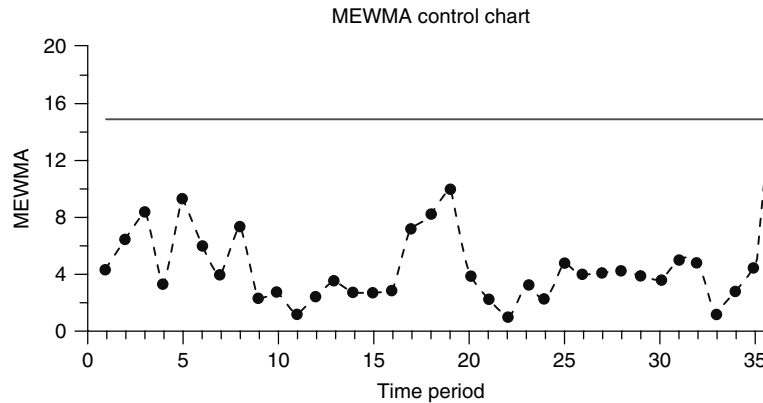


Figure 1 MEWMA control chart is based on the computed T^2 statistics

potential problem, it does not indicate which variable or set of variables are problematic. See **Multivariate Control Charts, Interpretation of** for techniques to determine which variable(s) are problematic in such a signal.

MEWMA charts are useful for detecting small shifts in the overall mean. Typically one uses values such as $0 < r \leq 0.2$ for these situations. In practice, some authors have suggested using two simultaneous MEWMA charts, one with a small value for r and another with $r \geq 0.5$. This latter chart will be able to detect large shifts better. Controlling the overall false alarm rate for these combined charts is somewhat problematic and has not been thoroughly addressed in the literature.

There are some difficulties with the MEWMA chart that also need to be addressed. First, there are some computational issues to deal with. However, all of the calculations in the example above were done using Microsoft Excel. Secondly, there is an “inertia” problem that arises when a process that has been in-control shifts out of control. It may take several time periods before the control chart signals. This problem can be diminished to some extent with the combination chart mentioned above. Calculating the false alarm rate for the MEWMA chart can be challenging. Rigdon [2] and [3] has developed an algorithm for computing false alarm rates. Molnau *et al.* [4] have developed a **Markov chain** based algorithm for calculating the false alarm rates. Linderman and Love [5] and Prabhu and Runger [6] discuss the design of MEWMA control charts.

Despite the computation and interpretation challenges of the MEWMA control chart, it is a good

choice when monitoring the mean vector of several correlated process variables. This chart is especially useful when one is interested in detecting relatively small shifts in the mean vector. Competing techniques include the Hotelling’s T^2 control chart (*see Hotelling’s T^2 Chart*) and the multivariate CUSUM control chart (*see Multivariate Cumulative Sum (CUSUM) Chart*). A general discussion of multivariate control chart may be found in **Multivariate Control Charts Overview**.

References

- [1] Lowry, C.A., Woodall, W.H., Champ, C.W. & Rigdon, S.E. (1992). A multivariate exponentially weighted moving average control chart, *Technometrics* **34**, 46–53.
- [2] Rigdon, S.E. (1995). An integral equation for the in-control average run length of a multivariate exponentially weighted moving average control chart, *Journal of Statistical Computation and Simulation* **52**, 351–365.
- [3] Rigdon, S.E. (1995). A double-integral equation for the average run length of a multivariate exponentially weighted moving average control chart, *Statistics and Probability Letters* **24**, 365–373.
- [4] Molnau, W.E., Runger, G.C., Montgomery, D.C., Skinner, K.R., Lored, E.N. & Prabhu, S.S. (2001). A program for ARL calculation for multivariate EWMA charts, *Journal of Quality Technology* **33**, 515–521.
- [5] Linderman, K. & Love, T.E. (2000). Economic and economic statistical designs for MEWMA control charts, *Journal of Quality Technology* **32**, 410–417.
- [6] Prabhu, S.S. & Runger, G.C. (1997). Designing a multivariate EWMA control chart, *Journal of Quality Technology* **29**, 8–15.

MICHAEL D. CONERLY

Multivariate Cumulative Sum (CUSUM) Chart

Introduction

Slight deviations of process variables from their target values are inevitable even when the production process is running normally. Under normal process operating conditions, these deviations in the variables are the interplay of multiple small causes, which are practically unavoidable and generally called the *common causes* or *chance causes of variation*. A process operating under common causes of variation is often considered to be statistically *in control* or *stable*. By statistically analyzing the common variation or the in-control process, process monitoring can be established for early detection of shifts to an out-of-control process. By assigning causes of these shifts, corrective actions can be taken and quality can be improved.

Multivariate control charts simultaneously utilize several related measures of process properties for collectively monitoring the statistical state of a process. The intention is rapid detection of abnormal situations and occurrences of faults, which are assumed to shift the process to an out-of-control state and are referred to as *assignable causes of variation*. Although routine practice of multivariate procedures may require computer computations, over the past decade multivariate control charts have been widely utilized in industry and new multivariate control charting techniques have been developed. This is due, at least in part, to the increased use of computers on the production floor and the amplified ability of computers and sensors to capture data on several process properties simultaneously. Various multivariate control charts (see **Multivariate Control Charts Overview**) exist in the literature. Among these multivariate control charts, the Hotelling's T^2 (see **Hotelling's T^2 Chart**; [1]), the multivariate **exponentially weighted moving average** (MEWMA) (see **Multivariate Exponentially Weighted Moving Average (MEWMA) Control Chart**; [2]), and the multivariate cumulative sum (MCUSUM) are the familiar ones.

In this article, the state of the art for the MCUSUM control charts is reviewed. The univariate cumulative sum (CUSUM) (see **Cumulative Sum (CUSUM) Chart**) procedure is described first since it is the

basis of the multivariate extensions. Next, some multivariate extensions to the univariate CUSUM control chart are explained with their descriptions and basic formulas. References are provided for the other MCUSUM procedures. Advantages and disadvantages of each multivariate procedure are discussed.

The Cumulative Sum Procedure

CUSUM control charts are important tools for the real-time, on-line monitoring of *univariate* (with one variable of interest) processes. These charts are especially effective in detecting small to moderate sized shifts in a process mean.

The classical one-sided CUSUM control chart can be derived from a sequential probability ratio test (SPRT) for the mean of a **normal distribution** or from a likelihood-ratio test for a change point in a normal distributed mean. Page [3] used SPRT to develop a one-sided CUSUM control for detecting shifts in the mean of a normal distributed variable. The null and the alternative hypotheses are

$$H_0 : \mu = \mu_0 \text{ and } H_1 : \mu = \mu_1 \quad (1)$$

respectively. Here, μ represents a process mean whose true value is unknown, μ_0 is the known on-target (in-control process) mean value and μ_1 is a specified off-target (out-of-control process) mean value to be detected quickly.

To generate a procedure for monitoring two-sided mean shifts, i.e., the mean μ increases to some value $\mu_1 > \mu_0$ or decreases to some value $\mu_1 < \mu_0$, Page proposed to simultaneously operate two one-sided CUSUMs. These two one-sided CUSUMs are designed typically with the same shift magnitude value δ for the increasing and decreasing off-target means. That is, $\mu_1 = \mu_0 \pm \delta$ and $\delta = |\mu_1 - \mu_0|$. The monitoring problem is easily accomplished with the recursive calculations that signal if either one-sided CUSUM

$$s_t^+ = \max \left(s_{t-1}^+ + \sigma^{-1} \left(y_t - \mu_0 - \frac{\delta}{2} \right), 0 \right) \text{ or} \\ s_t^- = \max \left(s_{t-1}^- + \sigma^{-1} \left(\mu_0 - y_t - \frac{\delta}{2} \right), 0 \right) \quad (2)$$

exceeds a predetermined decision interval, $h(s_t^+ > h$ or $s_t^- > h)$. Here, t denotes the time or sequence

2 Multivariate Cumulative Sum (CUSUM) Chart

index, y denotes the observation on the variable of interest distributed as normal with mean μ and constant variance σ^2 , and the statistics s_i^+ and s_i^- (generally with, $s_0^+ = s_0^- = 0$) are the one-sided upper and lower CUSUMs, respectively. It can be seen from equation (2) that sums are accumulated as the name implies, but an observation is accumulated only if it differs from the on-target mean value by more than one-half the specified shift magnitude δ . Such accumulation of information across successive observations is what makes the CUSUM procedure powerful in detecting small to medium sized shifts in the mean. In fact, this chart is the optimal diagnostic for a shift of the mean from the on-target mean value μ_0 to the off-target mean value μ_1 . Since the development of CUSUM by Page, a large number of studies have evaluated their theoretical and practical properties. A good recent review can be found in Hawkins and Olwell [4].

Multivariate Cumulative Sum Procedures

Several authors have proposed multivariate (multiple variables of interest) extensions to the univariate CUSUM procedure for monitoring processes with multiple related variables. Many of the key steps in the development of these MCUSUM procedures also appear in the construction of the CUSUM control charts for two-sided shifts (Runger and Testik [5]).

There are mainly two general approaches for the multivariate process monitoring problems. One approach comes from the univariate perspective where multiple univariate control charts are used simultaneously as a collection. Each of the univariate charts in the collection aims in some direction and these charts may be based on measured variables, principal component variables, regression adjusted variables, etc. A detailed discussion is provided in Runger and Montgomery [6]. An out-of-control alarm can be triggered individually by any of the univariate control charts of the collection. This approach may be ineffective in detecting some types of assignable causes since the multivariate characteristic of the data is no longer present directly. The other approach comes from the multivariate perspective, which utilizes the relationships between the variables. Multivariate observation vectors or a statistic of multivariate observation vectors are reduced to a scalar statistic by utilizing the relationships among the variables or

the process knowledge before being plotted on a single control chart. Then this single control chart is used for monitoring the variables of the process jointly. In the literature, MCUSUM procedures have been developed using one of these two approaches. However, before going into detail there is one more important characteristic of multivariate control charts that should be explained. Multivariate control charts are one of two types; directionally invariant and direction specific.

Direction-Specific versus Directionally Invariant Control Charts

Let us consider a univariate monitoring problem first. Univariate control charts for monitoring a process mean are specifically designed for the purpose of detecting shifts in the mean along its axis. Since the mean is a scalar value, it does not have a direction but has magnitude. If an **assignable cause** results in a shift of the mean, the shifted mean value relative to the on-target mean value would be either an increase or a decrease with some magnitude on the monitored variable's axis. Therefore, a two-sided control chart is used when both increases and decreases in the mean are important to detect. Consider the two-sided CUSUM chart in equation (2) where the δ term used in the upper CUSUM s_i^+ and the lower CUSUM s_i^- is the same. Because the observations are assumed to be distributed as normal, which is symmetric, performance of the two-sided CUSUM chart would be the same, on average, for detecting a mean shift having some magnitude λ , either an increase to $\mu = \mu_0 + \lambda$ or a decrease to $\mu = \mu_0 - \lambda$. Generally, the metric used for performance evaluation of control charts is the **average run length** (ARL), which is the average number of points plotted on a control chart before the control chart signals an alarm the first time. Consequently, ARL performance of the two-sided CUSUM chart is invariant to an increase or decrease of the process mean relative to the on-target mean but depends on the magnitude or equivalently the distance $\lambda(\mu) = |\mu - \mu_0|$ of the shifted mean from the on-target mean.

Now consider multivariate monitoring problems. Suppose that the variables are arranged in a vector. Unlike the scalar process mean values in a univariate setting, multivariate process mean vectors have both magnitude and direction. Due to correlations among the monitored variables, an assignable cause would

often result in a shift of the mean vector with some or all of the on-target mean values being changed. The shifted mean vector may have infinite number of possible shift directions that need to be monitored jointly, if not restricted. One approach to multivariate process mean monitoring is to sensitize the control chart to some known or anticipated shift directions when such process knowledge is available. Therefore, the control chart is expected to perform well in detecting mean shifts in these directions but not necessarily be as effective in other directions. The other approach is to design the control chart to simultaneously look at all shift directions such that the ARL performance is the same, on average, in detecting all mean vectors μ which are equidistant from μ_0 . The distance metric used here as a measure of shift magnitude is the statistical (Mahalanobis) distance, which is the square root of the noncentrality parameter

$$\lambda^2(\mu) = (\mu - \mu_0)' \Sigma^{-1} (\mu - \mu_0) \quad (3)$$

Consequently, a multivariate control chart is called *directionally invariant* if its ARL performance in detecting a shift of the mean from the on-target mean is determined solely by the statistical distance of the shifted mean from the on-target mean. Note that the univariate symmetric two-sided CUSUM in equation (2) is also directionally invariant. On the other hand, a multivariate control chart is *direction specific* if its ARL performance in detecting a shift of the mean from the on-target mean is a function of both the statistical distance and the particular direction of the shifted mean relative to the on-target mean. Multiple univariate control charts aim in specified directions and hence are direction specific.

When selecting a multivariate control chart it is important to consider the types of assignable causes that may affect the process. MCUSUM procedures in the literature are either directionally invariant or direction specific. In the following, considered MCUSUM procedures are classified based on their directional characteristic. Advantages and disadvantages of each are discussed.

The Process Model

Suppose that we have observations on p variables $y_1, y_2, \dots, y_p$ and these are arranged in a $p \times 1$

vector $\mathbf{y}$. Assume a sequence $\mathbf{y}_1, \mathbf{y}_2, \dots$ of the p -dimensional observation vectors and let these successive observation vectors be independent and identically distributed as p -variate normal with known on-target mean vector μ and known and constant $p \times p$ **covariance** matrix Σ , i.e., $\mathbf{y}_t \sim N_p(\mu, \Sigma)$ at time t . In the following, the vector μ will represent a multivariate normal process mean whose true value is unknown. Several cases when testing a shift in the mean vector from a known on-target value μ_0 will be considered by using the characteristics of the off-target mean vector.

Direction-Specific MCUSUM Procedures

Let us first consider the simplest case where the off-target mean vector is known in terms of both magnitude and direction. Denote this anticipated off-target mean vector by μ_1 . Note that the on-target and off-target process distributions are completely specified as $\mathbf{y}_t \sim N_p(\mu_0, \Sigma)$ and $\mathbf{y}_t \sim N_p(\mu_1, \Sigma)$, respectively. The null and alternative hypotheses are

$$H_0 : \mu = \mu_0 \text{ and } H_1 : \mu = \mu_1 \quad (4)$$

respectively. To detect a mean shift in one particular direction, a univariate CUSUM chart aimed in that one direction will give the best ARL performance. The solution introduced by Healy [7] is to apply a univariate CUSUM based on a linear combination of the observed variables. For two-sided shifts along the direction of the off-target mean vector, the chart signals if either

$$\begin{aligned} s_t^+ &= \max \left(s_{t-1}^+ + \mathbf{a}'(\mathbf{y}_t - \mu_0) - \frac{\delta}{2}, 0 \right) > h \text{ or} \\ s_t^- &= \max \left(s_{t-1}^- + \mathbf{a}'(\mu_0 - \mathbf{y}_t) - \frac{\delta}{2}, 0 \right) > h \end{aligned} \quad (5)$$

Here, δ is equal to the Mahalanobis distance of μ_1 to μ_0 , i.e., $\lambda(\mu_1)$ and the vector of the linear combination coefficients is $\mathbf{a} = \Sigma^{-1}(\mu_1 - \mu_0)/\delta$. The terms $\mathbf{a}'(\mathbf{y}_t - \mu_0)$ and $\mathbf{a}'(\mu_0 - \mathbf{y}_t)$ both have a standard univariate normal distribution when $\mu = \mu_0$. All of the available theory for designing univariate CUSUM charts and ARL performance analysis can be used for this MCUSUM chart.

Nevertheless, the solution is not very general since the shifted mean needs to be known exactly in terms of magnitude and direction. If the mean shifts some

4 Multivariate Cumulative Sum (CUSUM) Chart

other direction, this procedure may be anything from fairly effective to ineffective.

The problem is more interesting when it is only assumed that the process mean vector $\boldsymbol{\mu}$ has a known shift magnitude δ (in terms of the Mahalanobis distance of $\boldsymbol{\mu}$ to $\boldsymbol{\mu}_0$) but unknown direction when the process experiences a mean shift. Consider the following null and alternative hypothesis, respectively

$$H_0 : \boldsymbol{\mu} = \boldsymbol{\mu}_0 \text{ and } H_1 : \{\boldsymbol{\mu} | \lambda(\boldsymbol{\mu}) = \delta\} \quad (6)$$

It is convenient to standardize the observation vectors as

$$\mathbf{z}_t = \boldsymbol{\Sigma}^{-1/2}(\mathbf{y}_t - \boldsymbol{\mu}_0) \quad (7)$$

where $\boldsymbol{\Sigma}^{-1/2}$ is a square root of $\boldsymbol{\Sigma}$. In terms of $\mathbf{z}$, we can write

$$H_0 : \boldsymbol{\mu} = \mathbf{0} \text{ and } H_1 : \{\boldsymbol{\mu} | \|\boldsymbol{\mu}\| = \delta > 0\} \quad (8)$$

where $\|\cdot\|$ denotes the norm of a vector. One obvious solution to this problem is to apply a set of p two-sided CUSUM charts, one for each element of the vector $\mathbf{y}$ or one for each element of the vector $\mathbf{z}$. Woodall and Ncube [8] essentially considered this approach.

Charts based on $\mathbf{y}$ have the drawback that they do not directly use the **correlation** between the variables. Charts based on $\mathbf{z}$ have the advantages accrued from a principal components transformation and this is a rather natural choice. This is well known to be easily accomplished with the recursive calculations that signal if either

$$\begin{aligned} s_t^+(j) &= \max \left[s_{t-1}^+(j) + \mathbf{z}_t(j) - \frac{\delta(j)}{2}, 0 \right] > h \text{ or} \\ s_t^-(j) &= \max \left[s_{t-1}^-(j) - \mathbf{z}_t(j) - \frac{\delta(j)}{2}, 0 \right] > h \end{aligned} \quad (9)$$

Here, the index $j(j = 1, 2, \dots, p)$ is used to indicate the parameters of the CUSUM chart for the j th component $\mathbf{z}(j)$ of the standardized observation vector $\mathbf{z}$. Runger and Testik [5] proposed to name this MCUSUM procedure as MPYRAMID due to its geometric characteristic. Because the collection of CUSUM charts is considered as a single monitoring procedure, it is necessary to adjust the decision interval h of the charts to keep the total false alarm probability low. A method for selecting h and determining the total false alarm probability is given in Woodall and Ncube [8].

The ARL performance of the MPYRAMID chart depends on the direction of the shift vector and this dependency is lessened, but not removed, by the use of principal component variables $\mathbf{z}(j)$ rather than original variables. It is important to consider the type of mean shift that is important to detect. Charts based on $\mathbf{y}$ would aim along the axes of the original variables and would generally do a good job in detecting shifts along these axes. However, they may not be as effective in detecting shifts in a different direction such as a principal components axis. On the other hand, the MPYRAMID chart based on $\mathbf{z}$ would aim along the axes of the principal component variables and would generally do a good job in detecting shifts along these axes. Similarly, these charts may not be as effective in detecting shifts in different directions such as the axes of original variables.

Directionally Invariant MCUSUM Procedures

For a known off-target magnitude δ but unknown direction another approach is to apply a generalized likelihood-ratio (GLR) procedure. This procedure signals if

$$g_t = \max_{1 \leq \tau \leq t} \left[\delta \left\| \sum_{i=\tau}^t \mathbf{z}_i \right\| - \frac{\delta^2}{2}(t - \tau + 1) \right] > c \quad (10)$$

Here, g is the control statistic, τ is an unknown shift time ($1 \leq \tau \leq t$), and c is a decision interval. The derivation was provided by Basseville and Nikiforov [9]. This MCUSUM procedure is referred to as *MCONE* by Runger and Testik [5] because of its geometric characteristic.

Unfortunately, simplification of the procedure by recursive calculations is not available. However, an advantage of *MCONE* is that it is equally sensitive to shifts of the same magnitude (Mahalanobis distance) in all directions from the on-target mean. Consequently, it can be thought of as a collection of two-sided CUSUM charts operated in all directions simultaneously.

An example *MCONE* chart is shown in Figure 1 for a multivariate normal process monitoring problem with seven variables. Here the mean shifts to an off-target value at time 15 and the chart triggers an alarm at time 18. To illustrate how the calculations are performed, the control statistic g plotted on the control chart is calculated for the

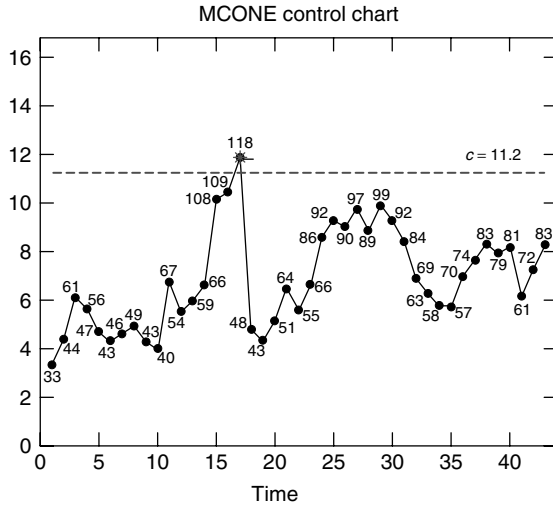


Figure 1 An example of the MCONE control chart

Table 1 An example of the MCONE calculations ($\delta = 1$)

t	τ	$\left\ \sum_{i=\tau}^t \mathbf{z}_i \right\ - \frac{1}{2}(t - \tau + 1)$	g_t
1	1	$3.83 - 0.5 = 3.33$	3.33
2	2	$2.37 - 0.5 = 1.87$	$\max(1.87, 4.40) = 4.40$
2	1	$5.40 - 1.0 = 4.40$	
3	3	$2.96 - 0.5 = 2.46$	$\max(2.46, 3.46, 6.09) = 6.09$
3	2	$4.46 - 1.0 = 3.46$	
3	1	$7.59 - 1.5 = 6.09$	

first three observation vectors as given in Table 1. Note that the value for δ is selected to be 1 in the calculations and the first three observation vectors are

$$\begin{aligned} \mathbf{z}_1 &= [-2.28 \quad 0.06 \quad 1.86 \quad -1.45 \quad 1.46 \\ &\quad -0.52 \quad 1.23]' \\ \mathbf{z}_2 &= [0.04 \quad -1.42 \quad 0.63 \quad -1.03 \quad 1.44 \\ &\quad -0.19 \quad -0.22]' \\ \mathbf{z}_3 &= [-2.26 \quad 0.08 \quad 0.03 \quad -0.56 \quad 1.43 \\ &\quad -0.45 \quad -1.04]' \end{aligned}$$

Other MCUSUM Procedures

Another direction-specific procedure for multivariate observations is the CUSUM of regression adjusted variables proposed by Hawkins [10]. Hawkins extended the approach of a known shift direction to the case in which several shift directions of interest are specified. Some other directionally invariant MCUSUM procedures in the literature include the COT and MCUSUM of Crosier [11] and MC1 and MC2 of Pignatiello and Runger [12]. Of these procedures, MCUSUM and MC1 are the ones with better performance. Although these two MCUSUM procedures have no known optimality properties, they have good practical performance. However, for these two procedures worst-case scenarios exist such that their steady-state performance is impacted by an accumulation of data in one direction before the process shifts another direction, namely the inertia problem. This is because one distinctive characteristic of univariate CUSUM procedure, that is that the steady-state performance is superior to the initial performance, does not appear in these charts. A detailed discussion of the MCUSUM procedures considered in this article as well as some others with performance comparisons can be found in Runger and Testik [5].

Summary

There are several multivariate extensions to the univariate CUSUM control chart. This work reviewed some of these multivariate extensions and compared their advantages and disadvantages. Basic formulas were provided and references were given to allow the interested reader to investigate further.

References

- [1] Hotelling, H.H. (1947). Multivariate quality control—illustrated by the air testing of sample bombsights, in *Techniques of Statistical Analysis*, C. Eisenhart, M.W. Hastay & W.A. Wallis, eds, McGraw-Hill, New York, pp. 111–184.
- [2] Lowry, C.A., Woodall, W.H., Champ, C.W. & Rigdon, S.E. (1992). A multivariate exponentially weighted moving average chart, *Technometrics* **34**, 46–53.
- [3] Page, E.S. (1954). Continuous inspection schemes, *Biometrika* **41**, 100–114.
- [4] Hawkins, D.M. & Olwell, D.H. (1998). *Cumulative Sum Charts and Charting for Quality Improvement*, Springer-Verlag, New York.

6 Multivariate Cumulative Sum (CUSUM) Chart

- [5] Runger, G.C. & Testik, M.C. (2004). Multivariate extensions to cumulative sum control charts, *Quality and Reliability Engineering International* **20**, 587–606.
- [6] Runger, G.C. & Montgomery, D.C. (1997). Multivariate and univariate process control: geometry and shift directions, *Quality and Reliability Engineering International* **13**, 153–158.
- [7] Healy, J.D. (1987). A note on multivariate CUSUM procedures, *Technometrics* **29**, 409–412.
- [8] Woodall, W.H. & Ncube, M.M. (1985). Multivariate CUSUM quality control procedures, *Technometrics* **27**, 285–292.
- [9] Basseville, M. & Nikiforov, I.V. (1993). *Detection of Abrupt Changes: Theory and Application*, Prentice Hall, Englewood Cliffs.
- [10] Hawkins, D.M. (1991). Multivariate quality control based on regression-adjusted variables, *Technometrics* **33**, 61–75.
- [11] Crosier, R.B. (1988). Multivariate generalizations of cumulative sum quality control schemes, *Technometrics* **30**, 291–303.
- [12] Pignatiello, J.J. & Runger, G.C. (1990). Comparisons of multivariate *CUSUM* charts, *Journal of Quality Technology* **22**, 173–186.

MURAT C. TESTIK AND GEORGE C. RUNGER

Multivariate Charts for Variability

Introduction

Multivariate control charts (*see Multivariate Control Charts Overview*) which are specifically designed to detect changes in a **covariance** matrix are discussed here. The covariance matrix represents the variability of a multivariate process. The scope of the discussion is limited to phase II **control charts** for multivariate normal processes. For a more detailed account of the topic being discussed, readers are referred to the works by Weirda [1], Lowry and Montgomery [2], Montgomery [3], and Yeh *et al.* [4], and the references therein (*see Hotelling's T^2 Chart; Multivariate Exponentially Weighted Moving Average (MEWMA) Control Chart; Multivariate Cumulative Sum (CUSUM) Chart*).

Let $X = (X_1, X_2, \dots, X_p)'$ denote the random variable that represents the p correlated quality characteristics derived from a process whose quality is to be monitored. When the process is in control, it is assumed that X follows a p -dimensional **normal distribution**, denoted by $N_p(\mu_0, \Sigma_0)$, where μ_0 is the in-control process mean and Σ_0 is the in-control process covariance matrix. On the other hand, when the process is out of control, it is assumed that X follows $N_p(\mu, \Sigma)$, where either $\mu \neq \mu_0$ or $\Sigma \neq \Sigma_0$ or both. The in-control parameters μ_0 and Σ_0 are assumed to be either known or that they can be sufficiently estimated at the end of phase I control such that their values can be assumed known prior to the beginning of the phase II monitoring. Here we discuss two phase II multivariate control charting techniques for detecting changes in the covariance matrix, each dependent on how the samples of observations are collected. In the first case, independent samples of $n(n > p)$ observations each are repeatedly drawn. In the second case, independent individual observations ($n = 1$) are being drawn from the process.

The $|S|$ Chart

In the case when independent samples each with n observations are available for constructing the control chart, a commonly used Shewhart control

chart is known as the $|S|$ chart (*see Multivariate Control Charts Overview*), or the generalized variance chart. For the t th sample, $t = 1, 2, \dots$, the n observations are denoted by $X_{t1}, X_{t2}, \dots, X_{tn}$, and the corresponding sample mean and covariance matrix are $\bar{X}_t = \sum_{j=1}^n X_{tj}/n$ and $S_t = \sum_{j=1}^n (X_{tj} - \bar{X}_t)(X_{tj} - \bar{X}_t)'/(n-1)$, respectively. To construct the $|S|$ chart, one calculates $|S_t|$, $t = 1, 2, \dots$, the determinant of the sample covariance matrix of the t th sample and plots $|S_t|$ against the sampling sequence t . For 3σ control limits, the upper control limit (UCL), center line (CL), and lower control limit (LCL) of the $|S|$ chart are defined as

$$UCL = \left(1 + 3\frac{\sqrt{b_2}}{b_1}\right) |\Sigma_0| \quad (1)$$

$$CL = |\Sigma_0| \quad (2)$$

$$LCL = \max\left(0, \left(1 - 3\frac{\sqrt{b_2}}{b_1}\right) |\Sigma_0|\right) \quad (3)$$

where $b_1 = (n-1)^{-p} \prod_{i=1}^p (n-i)$ and $b_2 = (n-1)^{-2p} \left[\prod_{i=1}^p (n-i)\right] \times \left[\prod_{j=1}^p (n-j+2) - \prod_{j=1}^p (n-j)\right]$. An out-of-control signal is detected if $|S_t|$ plots outside of the control limits.

We next present the application of the $|S|$ chart using a real-life example. The example from Mitra [5] is related to a component used in the assembly of a transmission mechanism. Two correlated quality characteristics, tensile strength and diameter, are to be monitored. The data set consists of 20 independent samples each with four observations. In this case, $p = 2$ and $n = 4$, which result in $b_1 = \frac{2}{3}$ and $b_2 = \frac{84}{81}$. The centerline $|\Sigma_0|$ was estimated by the average of the sample covariance matrices of all 20 samples which turned out to be 48.4. Therefore the UCL , CL , and LCL were calculated to be 270.1, 48.4, and 0, respectively. The $|S|$ chart, which was produced using Minitab, is shown in Figure 1. The $|S|$ chart indicated that the covariance matrix was in control.

The MEWMS Chart

In the case when only independent individual observations are available, the $|S|$ chart is no longer applicable since the sample covariance matrix does not exist. The alternative is to find other ways to pool observations together and derive control charts based on the pooled observations. Let $Z_t = \Sigma_0^{-1/2}(X_t -$

2 Multivariate Charts for Variability

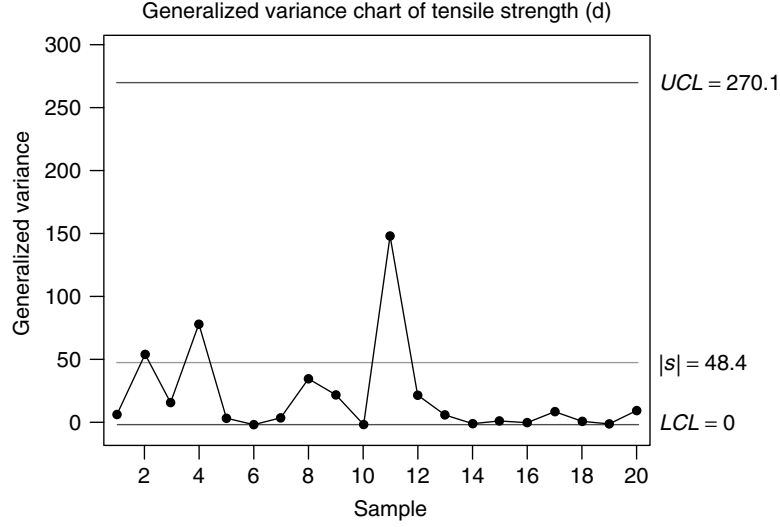


Figure 1 The $|S|$ chart of the transmission assembly example

μ_0), $t = 1, 2, \dots$, be the transformed variable of the t th observation X_t such that when the process is in control Z_t follows $N_p(0, I_p)$, where I_p is a $p \times p$ identity matrix. Let $V_t = \lambda Z_t Z_t' + (1 - \lambda)V_{t-1}$, where $0 < \lambda < 1$ is a smoothing constant and $V_0 = Z_1 Z_1'$. For each sampling sequence t , one calculates $\text{trace}(V_t)$, the trace of the matrix V_t . The values of $\text{trace}(V_t)$ are then plotted on the control chart against the sampling frequencies. For 3σ control limits, the UCL, CL, and LCL are defined as

$$UCL = p + L \sqrt{2p \sum_{i=1}^t c_i^2} \quad (4)$$

$$CL = p \quad (5)$$

$$LCL = p - L \sqrt{2p \sum_{i=1}^t c_i^2} \quad (6)$$

where L is a multiplier dependent on p, λ and the desired in-control chart performance, and $\sum_{i=1}^t c_i^2 = \lambda(2 - \lambda)^{-1} + 2(2 - \lambda)^{-1}(1 - \lambda)^{2t-1}$ which converges to $\lambda/(2 - \lambda)$ as $t \rightarrow \infty$. Using simulations, Huwang *et al.* [6] provided the values of L for λ ranging from 0.1 to 0.9 and $p = 2$ and 3, such that when the process is in control the control chart will have an in-control **average run length** approximately equal to 370, which is equivalent to a 3σ limits Shewhart chart. Part of the tabulated

values of L is provided in Table 1. Note that the chart, called the *multivariate exponentially weighted moving squared-deviation* (MEWMS) chart, is a multivariate extension of the univariate EWMS control chart first proposed by MacGregor and Harris [7]. In general, a smaller value of λ , say $0.1 \leq \lambda \leq 0.3$, gives a better performance of the MEWMS chart – in term of faster detection of out-of-control processes.

Next we present the application of the MEWMS chart. The data were taken from Joshi and Sprague [8] for an application in the semiconductor industry. The observations were related to the quality of wafer manufacturing during photolithography process. There were 74 lots of wafers and one wafer was randomly taken from each lot. Three correlated critical dimension measurements of die were taken on each drawn wafer. Using these samples, except lot 40 which was classified to be out of control, estimates of μ_0 and Σ_0 were obtained and used to establish the control limits for phase II monitoring. The estimates

Table 1 The values of L of the MEWMS chart for $p = 2$ and 3

λ	$L(p = 2)$	$L(p = 3)$
0.1	2.9013	2.8212
0.2	3.4870	3.3281
0.3	3.8713	3.6621

μ_0 and Σ_0 are equal to $\hat{\mu}_0 = \begin{pmatrix} 3.315 \\ 3.108 \\ 3.118 \end{pmatrix}$ and $\hat{\Sigma}_0 = \begin{pmatrix} 0.0093 & 0.0036 & 0.0052 \\ 0.0036 & 0.0085 & 0.0034 \\ 0.0052 & 0.0034 & 0.0088 \end{pmatrix}$, respectively. In phase II, an additional 200 observations were obtained. For every new observation generated, it was transformed using $Z_t = \hat{\Sigma}_0^{-1/2}(X_t - \hat{\mu}_0)$, $t = 1, 2, \dots, 200$. The matrix $V_t = \lambda Z_t Z_t' + (1 - \lambda)V_{t-1}$ was updated and the value of $\text{trace}(V_t)$ was calculated and plotted against the sampling sequence. The process was in control for the first 49 observations. Starting from observation 50, a mean shift occurred in the process and the observations starting from sample 50 were generated from a process having a new mean vector equal to $(3.231, 3.200, 3.118)$ and the same covariance matrix as $\hat{\Sigma}_0$. Beginning at observation 100, an additional change in the covariance matrix also took place. The new covariance matrix starting from sample 100 is equal to $\Sigma = \begin{pmatrix} 0.0186 & 0.0036 & 0.0052 \\ 0.0036 & 0.0170 & 0.0034 \\ 0.0052 & 0.0034 & 0.0088 \end{pmatrix}$. The corresponding phase II MEWMS chart with $p = 3$, $\lambda = 0.2$, and $L = 3.3281$ was produced using S-plus and shown in Figure 2. The MEWMS chart detected out-of-control signals at as early as observation 54, which was primarily due to a mean shift in the process. The chart also showed an out-of-control signal at observation 100 when the change in the covariance matrix actually occurred.

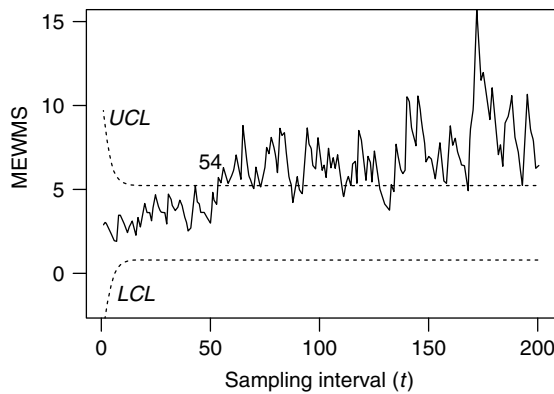


Figure 2 The MEWMS chart of the wafer manufacturing example

Discussion

The $|S|$ chart is relatively easy to construct and use, and it is available in some of the commonly used commercial software packages such as Minitab. It is also less sensitive to mean shifts in the sense that a shift in the mean is less likely to cause the $|S|$ chart to produce false out-of-control signals. On the other hand, the chart is restricted to the case when $n > p$. Tang and Barnett [9] showed such control limits may lead to considerably larger false alarm rate than the nominal value of 0.0027, especially when n is relatively small. In the same paper, they provided accurate calculations of control limits for the cases when $p = 3$ and 4.

Since the $|S|$ chart is a Shewhart chart, it is generally less effective in detecting small changes in the generalized variance. Yeh *et al.* [4] proposed an exponentially weighted moving averages (EWMA) chart based on the EWMA of $\log |S_t|$. Such an EWMA chart is more effective in detecting small changes in the generalized variance. Other existing multivariate control charts for variability derived from the sample generalized variance can be found in works by Guerrero-Cusumano [10], Tang and Barnett [11], Levinson *et al.* [12], Yeh and Lin [13], and Yeh *et al.* [14].

In the case of $n = 1$, none of the aforementioned charts is applicable since the sample covariance matrix cannot be evaluated. The MEWMS chart discussed above is relatively easy to construct. Although most of the commercial software packages currently do not include the MEWMS chart, it was relatively simple for its customized codes. For other charts that are applicable to the case when $n = 1$, readers are referred to the works by Chan and Zhang [15], Khoo and Quah [16], Yeh *et al.* [17], and Reynolds and Cho [18].

Conclusions

We discuss two multivariate control charts for monitoring variability of multivariate normal processes. The $|S|$ chart is applicable when $n > p$, while the MEWMS chart is specifically designed for individual observations ($n = 1$). Examples of other types of charts include the nonparametric procedures designed to detect changes in a covariance matrix as defined by changes of some function of the matrix (Hawkins [19]), the data depth based nonparametric control

charts for nonnormal processes (Liu [20]), and the principle component analysis and dissimilarity index based control charts for multivariate time-series data (Kano *et al.* [21]).

The $|S|$ chart is essentially designed to detect changes in the determinant of the covariance matrix. When a point is outside of the control limits on the $|S|$ chart, caution should be taken when interpreting such a signal (*see Multivariate Control Charts, Interpretation of*). The out-of-control signal is primarily due to a possible increase or decrease in the determinant of the covariance matrix. It does not necessarily imply that there is an increase or decrease in the process variability. Johnson and Wichern [22] gave three covariance matrices for bivariate data that all have the same determinant and yet have very different correlations.

An important area not discussed here is the diagnostic techniques. This is more complicated than the univariate case, mainly due to the complexity of the covariance matrix. Because so many parameters are contained in the covariance matrix and that changes in one or some of the parameters can trigger an out-of-control signal, it is of eminent importance to be able to further pinpoint which of these parameters are out of control. Some potential diagnostics are discussed in Yeh *et al.* [4], and the references therein. In another recent paper, Apley and Shi [23] proposed a technique based on a fault model which transformed the problem of diagnosing into the problem of estimating the number of faults which contribute to changes in process variability and the linear combinations by which each of the p correlated quality characteristics is affected.

References

- [1] Weirda, S.J. (1994). *Multivariate Statistical Process Control*, Wolters-Noordhoff, Groningen.
- [2] Lowry, C.A. & Montgomery, D.C. (1995). A review of multivariate control charts, *IIE Transactions in IIE Research* **27**, 800–810.
- [3] Montgomery, D.C. (2005). *Introduction to Statistical Quality Control*, 5th Edition, John Wiley & Sons, New York.
- [4] Yeh, A.B., Lin, D.K.-J. & McGrath, R.N. (2006). Multivariate control charts for monitoring covariance matrix: a review, *Journal of Quality Technology and Quantitative Management* **3**, 415–436.
- [5] Mitra, A. (1993). *Fundamentals of Quality Control and Improvement*, Macmillan, New York.
- [6] Huwang, L., Yeh, A.B. & Wu, C.-W. (2007). Monitoring multivariate process variability for individual observations, *Journal of Quality Technology* **39**(3), 258–278.
- [7] MacGregor, J.F. & Harris, T.J. (1993). The exponentially weighted moving variance, *Journal of Quality Technology* **25**, 106–118.
- [8] Joshi, M. & Sprague, K. (1997). Obtaining and using statistical process control limits in the semiconductor industry, in *Statistical Case Studies for Industry Process Improvement*, V. Czitrom & P.D. Spagon, eds, SIAM, Philadelphia, pp. 337–356.
- [9] Tang, P.F. & Barnett, N.S. (1996). Dispersion control for multivariate processes-some comparisons, *The Australian Journal of Statistics* **38**, 253–273.
- [10] Guerrero-Cusumano, J.-L. (1995). Testing variability in multivariate quality control: a conditional entropy measure approach, *Information Sciences* **86**, 179–202.
- [11] Tang, P.F. & Barnett, N.S. (1996). Dispersion control for multivariate processes, *The Australian Journal of Statistics* **38**, 235–251.
- [12] Levinson, W., Holmes, D.S. & Mergen, A.E. (2002). Variation charts for multivariate processes, *Quality Engineering* **14**, 539–545.
- [13] Yeh, A.B. & Lin, D.K.-J. (2002). A new variables control chart for simultaneously monitoring multivariate process mean and variability, *International Journal of Reliability, Quality and Safety Engineering* **9**, 41–59.
- [14] Yeh, A.B., Lin, D.K.-J., Zhou, H. & Venkataramani, C. (2003). A multivariate exponentially weighted moving average control chart for monitoring process variability, *Journal of Applied Statistics* **30**, 507–536.
- [15] Chan, L.K. & Zhang, J. (2001). Cumulative sum control charts for the covariance matrix, *Statistica Sinica* **11**, 767–790.
- [16] Khoo, M.B. & Quah, S.H. (2003). Multivariate control chart for process dispersion based on individual observations, *Quality Engineering* **15**, 639–642.
- [17] Yeh, A.B., Huwang, L. & Wu, C.-W. (2005). A multivariate EWMA control chart for monitoring process variability with individual observations, *IIE Transactions on Quality and Reliability Engineering* **37**, 1023–1035.
- [18] Reynolds Jr, M.R. & Cho, G.-Y. (2006). Multivariate control charts for monitoring the mean vector and covariance matrix, *Journal of Quality Technology* **38**, 230–253.
- [19] Hawkins, D.M. (1992). Detecting shifts in functions of multivariate location and covariance parameter, *Journal of Statistical Planning and Inference* **33**, 233–244.
- [20] Liu, R.Y. (1995). Control charts for multivariate processes, *Journal of the American Statistical Association* **90**, 1380–1387.
- [21] Kano, M., Nagao, K., Hasebe, S., Hashimoto, I., Ohno, H., Strauss, R. & Bakshi, R.R. (2002). Comparison of multivariate statistical process monitoring methods with applications to the Eastman challenge problem, *Computers and Chemical Engineering* **26**, 161–174.

- [22] Johnson, R.A. & Wichern, D.W. (2002). *Applied Multivariate Statistical Analysis*, 5th Edition, Prentice Hall.
- [23] Apley, D.W. & Shi, J. (2001). A factor-analysis method for diagnosing variability in multivariate manufacturing processes, *Technometrics* **43**, 84–95.

Multivariate Exponentially Weighted Moving Average (MEWMA) Control Chart.

DENNIS K.J. LIN AND ARTHUR B. YEH

Related Articles

Hotelling's T^2 Chart; Multivariate Control Charts Overview; Multivariate Control Charts, Interpretation of; Multivariate Cumulative Sum (CUSUM) Chart;

Multivariate Control Charts, Interpretation of

Introduction

Statistical process control procedures are designed to work well in a variety of conditions. The general idea is that a stable process containing only common-cause variation is easiest to control, but when the process variation is mixed with special-cause components, statistical input is needed to help monitor the process and detect the changes. One of the goals of a statistical process control (SPC) procedure is to detect the presence of special causes of variation quickly and reliably, and to locate the source of the variation in terms of the process variables so that corrective action may be taken (see **Control Charts, Overview**). This is achieved by using a charting statistic and interpreting the output when it signals that a process change has occurred.

When only one process variable is used to measure process performance, changes in the mean and/or changes in the variation of the variable can be detected using single-variable control charts. When several variables are used to monitor a process, the interpretation of signals in a multivariate control chart is not so simple. In addition to changes in location and variation, there may also be changes in the relationships between the many variables. This occurs when there is a **covariance** or **correlation** change between two variables or among a group of variables. Any of these changes can indicate a signal on the control chart for the multivariate process. For example, a change in the correlation between temperature and pressure might indicate that the observed pressure is either too high or too low for the given observed value of the temperature, relative to the historical process data.

Multivariate process control (see **Multivariate Control Charts Overview**) can be based on observing an individual observation vector, or on observing a statistic, such as a mean vector or covariance matrix estimate, computed from a sample of observations taken at each time period. In general, process industries such as the chemical industry favor monitoring individual observations, while other industries such as the manufacturing industry often monitor

statistics based on the mean of a sample of observations. We will examine signal interpretation for the case where control is based on an individual multivariate observation vector taken in a monitoring stage (labeled phase II). However, the procedures we present can be adapted for control based on the sample mean as well as for phase I operations where a baseline or historical data set (HDS) is being constructed.

Signal Interpretation

One of the more popular multivariate control procedures is based on Hotelling's T^2 statistic (see **Hotelling's T^2 Chart**). Consider a p -dimensional vector of observations, $x' = (x_1, x_2, \dots, x_p)$, that is made on a process at a specified time period. We assume that, when the process is in control, the $\mathbf{x}$ vectors are independent and follow a multivariate **normal distribution** with an unknown mean vector μ and an unknown covariance matrix Σ . The form and distribution of the T^2 statistic for the observation vector $\mathbf{x}$ used in phase II is given by

$$T^2 = (\mathbf{x} - \bar{\mathbf{x}})' \mathbf{S}^{-1} (\mathbf{x} - \bar{\mathbf{x}}) \\ \sim \left[\frac{p(n+1)(n-1)}{n(n-p)} \right] F_{(p, n-p)} \quad (1)$$

The term $F_{(p, n-p)}$ denotes an F distribution with p and $n-p$ **degrees of freedom**, and the estimates, $\bar{\mathbf{x}}$ and $\mathbf{S}$, are obtained from a phase I operation where the created HDS is of size n .

The popularity of the T^2 statistic is due in part to its ease of calculation and also because of the availability of procedures that can be used to interpret its signals. Two well-known procedures are based on using transformations to decompose a signaling T^2 statistic into the sum of p orthogonal components. These orthogonal components can be related to the contributions of individual process variables or groups of variables to the signal. Thus, we are quickly able to determine if the signaling observation vector agrees with the measures of location or variation/covariation of the in-control multivariate distribution.

The oldest of the decomposition procedures is obtained by transforming the signaling T^2 value into a function of the principal components of the estimated covariance matrix, $\mathbf{S}$, or the corresponding estimated correlation matrix. Using the estimated

2 Multivariate Control Charts, Interpretation of

covariance matrix, we define the i th principal component of $\mathbf{S}$ as the linear combination of the p process variables given by

$$z_i = \mathbf{u}_i'(\mathbf{x} - \bar{\mathbf{x}}) \quad (2)$$

for $i = 1, 2, \dots, p$, where $\mathbf{u}_i$ is the i th eigenvector of $\mathbf{S}$ and λ_i is the corresponding eigenvalue [1]. Using this transformation, the T^2 statistic in equation (1) becomes

$$T^2 = \frac{z_1^2}{\lambda_1} + \frac{z_2^2}{\lambda_2} + \dots + \frac{z_p^2}{\lambda_p} \quad (3)$$

Since the principal component in equation (2) is a linear combination of the p process variables, it can be used to identify the groups of variables contributing to the signal. However, it has an inherent weakness in that it cannot provide a straightforward interpretation in terms of specific process variables.

A more efficient procedure, known as the MYT decomposition, is described below. The MYT decomposition [2] is defined by using the following transformation of the p process variables

$$\mathbf{y} = \mathbf{L}^{-1}(\mathbf{x} - \bar{\mathbf{x}}) \quad (4)$$

where $\mathbf{L}$ is a lower triangular invertible matrix with the property that $\mathbf{S} = \mathbf{L}\mathbf{L}'$. Using this transformation, the T^2 statistic in equation (1) becomes

$$\begin{aligned} T^2 &= y_1^2 + y_2^2 + \dots + y_p^2 \\ &= T_1^2 + T_{2,1}^2 + \dots + T_{p,1,2,\dots,p-1}^2 \end{aligned} \quad (5)$$

The T_1^2 term in equation (5) is labeled an unconditional term and the $T_{j,1,2,\dots,j-1}^2$ terms are labeled conditional terms.

Using equation (4), the general form of the unconditional term contained in any phase II decomposition is given as

$$T_j^2 = \frac{(x_j - \bar{x}_j)^2}{s_j^2} \quad (6)$$

where x_j is the j th component of $\mathbf{x}$, and $\bar{x}_j$ and s_j^2 are its corresponding mean and variance estimates as determined using the HDS. Under the assumption of no signal, the T_j^2 statistic is described by an F distribution given as

$$T_j^2 \sim \left(\frac{n+1}{n} \right) F_{(1,n-1)} \quad (7)$$

We use this term to check if the value of the j th variable is outside of its operational range.

Similarly, the general form of the conditional term in any phase II decomposition is given as

$$T_{j,1,2,\dots,j-1}^2 = \frac{(x_j - \bar{x}_{j,1,2,\dots,j-1})^2}{s_{j,1,2,\dots,j-1}^2} \quad (8)$$

This is the square of the value of the j th variable adjusted by the estimates of the mean and variance of the conditional distribution of x_j adjusted for the variables $x_1, x_2, \dots, x_{j-1}$. The estimates, $\bar{x}_{j,1,2,\dots,j-1}$ and $s_{j,1,2,\dots,j-1}^2$, are obtained from the HDS [3]. With no signal present, the distribution of a conditional term is given as

$$T_{j,1,2,\dots,j-1}^2 \sim \left(\frac{(n+1)(n-1)}{n(n-k-1)} \right) F_{(1,n-k-1)} \quad (9)$$

where $k = j - 1$. We use this term to determine if the actual value of the j th variable is close to the value predicted using the values of the remaining variables.

Since the subscripts $\{1, 2, \dots, p\}$ in the T^2 terms on the right-hand side of equation (5) can be replaced by any permutation of $\{1, 2, \dots, p\}$, the corresponding distributions of the T^2 statistics in equations (7) and (9) remain invariant with respect to the permutation.

Computer algorithms based on equation (5) are readily available for determining which of the p process variables are contributing to the signaling T^2 value [4]. Generally, these systems work by examining all possible decompositions of a signaling T^2 value. A total decomposition includes computing all the possible distinct MYT components: the p unconditional terms, T_j^2 ; all two-way conditional terms $T_{i,j}^2$; all three-way terms, $T_{i,j,k}^2$; etc. There are $2^{(p-1)} \times (p)$ such terms, but several efficient computational shortcuts exist for obtaining the relevant values.

For example, a useful stepwise process involves first using the p unconditional terms described in equation (6) to check the tolerance of each individual variable. Any variable with a significant F value as given by equation (7) is declared to be out of tolerance. This variable would then be considered as part of the signal and removed from further consideration. Next the two-way terms are checked using equation (9). Any signaling $T_{i,j}^2$ component implies that the relationship between x_i and x_j differs from that given in $\mathbf{S}$. If there is a signal, both variables

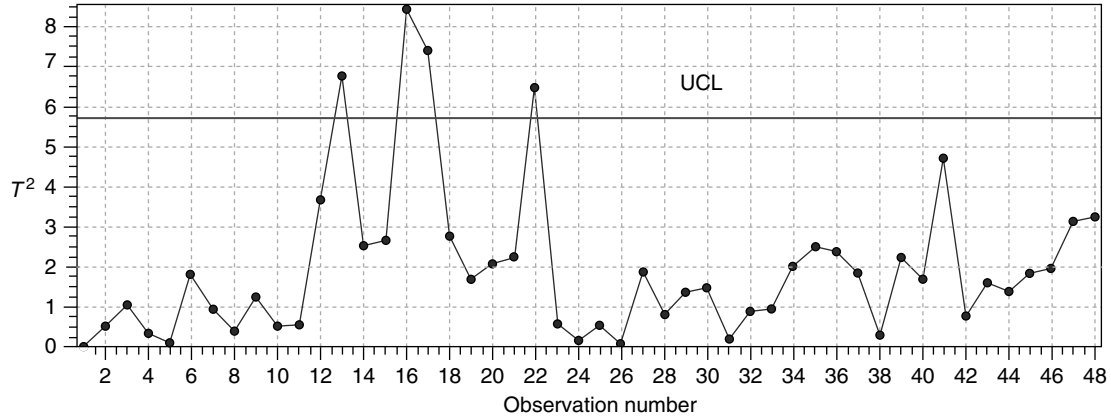


Figure 1 T^2 chart for monitoring the turbine system

are removed from further consideration and the three-way terms are checked and so on.

Data Example

The MYT decomposition procedure is illustrated in the following data example. More details on how to use it and how to interpret its orthogonal components is given in [5]. A steam turbine system is used to generate electricity for industrial use. Our purpose is to show how signal interpretation can be used in locating the source of a signal in this process. To simplify the problem, we consider only two of numerous variables that are used in controlling the turbine system. These variables include the amount of fuel (F) being used and the amount of megawatts (MW) that is being produced by the turbine. Summary statistics for the HDS for the two variables are contained in Table 1.

A typical T^2 control chart used in monitoring the turbine system is given in Figure 1. All T^2

values below the upper control limit (UCL), obtained using an $\alpha = 0.0027$, indicate that performance is in agreement with the baseline conditions. The T^2 values above the UCL indicate signals.

A scatter plot of the turbine data (fuel vs megawatts) along with the corresponding T^2 control ellipse using the formula in equation (1) is presented in Figure 2. Note that the signaling T^2 points identified in Figure 1 are located outside the control ellipse. The two lines in the plot represent the regression of fuel on megawatts (F/M) and the regression of megawatts on fuel (M/F).

In order to illustrate signal interpretation, we apply the MYT decomposition to the four signaling T^2 values corresponding to observations #13, #16, #17, and #22. A total decomposition giving all components and their values for the signaling T^2 values are presented in Table 2. The subscript 1 refers to the megawatt variable while the subscript 2 refers to the fuel variable.

Table 1 Summary statistics for the historical data set

	MW	F
Mean	53.72	692.45
Standard deviation	3.41	30.41
Correlation matrix		
	MW	F
MW	1.000	0.7469
F	0.7469	1.000

Discussion of Results

Observation #13, located at the lower end of the control ellipse, has a significant unconditional component T_1^2 and a significant conditional component $T_{1,2}^2$. The signaling unconditional term, T_1^2 , indicates that the megawatt value is outside the tolerance limits (i.e., range) established using the HDS. The signaling conditional term, $T_{1,2}^2$, indicates that, given the value of fuel for this observation, the value of megawatts is not where it should be relative to the

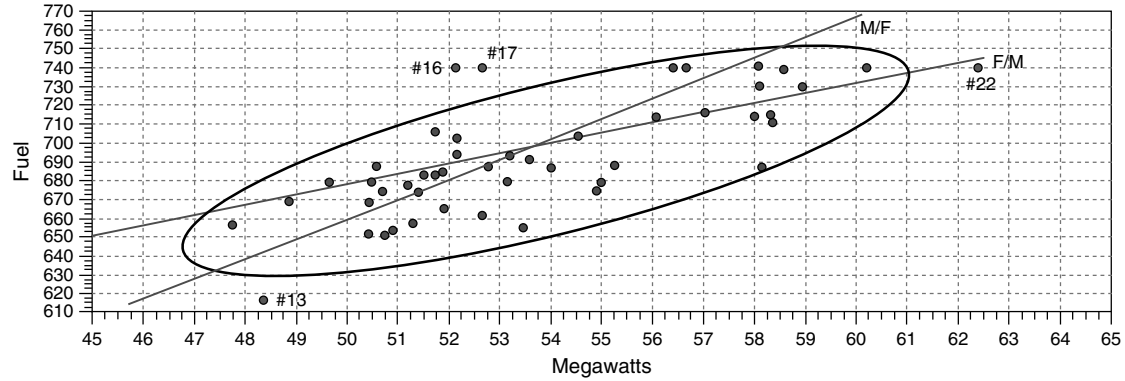


Figure 2 Scatter plot of data

HDS. This is further denoted by the large vertical distance (residual) between the observation and the regression line of megawatts on fuel (M/F) given in Figure 2.

Both of the T^2 conditional terms for observations #16 and #17 are significant. These signals indicate that there is a fouled linear relationship between the megawatts being produced and the fuel being used, and that the two variables are countercorrelated relative to the correlations observed using the HDS. From the plot given in Figure 2, it appears that for these two observations either the megawatts value is too low for the given value of fuel, or the fuel value is too high for the given value of megawatts.

Observation #22, located at the upper end of the control region given in Figure 2, has a significant unconditional term T_2^2 . This indicates that the fuel usage for this observation is outside the tolerance limits (i.e., range) established using the HDS.

When only two variables are being considered, making it possible to plot the control ellipse, we usually find that points located outside the ends of the ellipse, such as observation #22 in our example, have out-of-tolerance variables. Those observations

located around the middle, but outside of the control ellipse, such as observations #16 and #17 in our example, will have fouled relationships between the two variables. Similarly, those observations located between the middle and the end, but outside of the ellipse, such as observation #13 in our example, may have both problems.

Summary

The MYT decomposition has many good features. It provides an orthogonal decomposition of a T^2 value, and the corresponding components indicate which variables in the signaling observation vector disagree with the historical data. The cause of the signal can be directly attributed to a variable or set of variables being out of tolerance, and/or to the correlation between a set of variables being counter to the correlation observed in the baseline data. What it does not answer is whether the cause of the problem is due to a change in the mean vector or to a change in the covariance matrix, since either could produce variables that are out of tolerance and/or

Table 2 Decomposition components of T^2 signals

Observation	#13	#16	#17	#22
T_1^2	6.4940 ^(a)	0.2094	0.0967	2.4639
$T_{2,1}^2$	0.2580	8.2104 ^(a)	7.2936 ^(a)	4.0189
T_2^2	2.4515	2.4441	2.4441	6.2770 ^(a)
$T_{1,2}^2$	4.3005 ^(a)	5.9758 ^(a)	4.9462 ^(a)	0.2059

^(a)Indicates significant using $\alpha = 0.0027$

countercorrelated. This problem also exists when one uses principal components analysis. In these types of decompositions, we are basically identifying the change and not necessarily isolating the mechanism leading to the change.

Other approaches exist for use in interpreting T^2 signals. Several of these have been shown to be related to the components of the MYT decomposition (see [2]), though none are as comprehensive. These include the method in [6, 7] based on using regression-adjusted variables to improve diagnostics as well as the step-down method in [8, 9] that sequentially tests subsets of the variables assuming there is an *a priori* ordering of the variables. Another related procedure [10] is based on using a univariate t statistic to rank the components of an observation vector according to their contribution to the signal. And in [11] a method is given that is a function of a sum of a subset of the conditional T^2 components.

Signal interpretation can also be accomplished using other techniques. Fuchs and Benjamin [12] present a graphical method based on multivariate profile plots, while Sparks *et al.* [13] use the Gabriel biplot. Kouti and MacGregor [14] provide an approach based on usage of normalized principal components and contribution plots. Maravelakis *et al.* [15] propose a modified approach based on principal components. Nedumaran and Pignatiello [16] describe an approach based on Scheffe-type intervals and usage of univariate diagnostic charts for the individual variables as well as univariate charts for the principal components. Although the list of alternative methods for interpreting signals in a multivariate control chart continues to expand, orthogonal decomposition techniques, particularly those based on the MYT decomposition or principal components, remain the most useful and informative.

References

- [1] Jackson, J.E. (1991). *A User's Guide to Principal Components*, John Wiley & Sons, New York.
- [2] Mason, R.L., Tracy, N.D. & Young, J.C. (1995). Decomposition of T^2 for multivariate control chart interpretation, *Journal of Quality Technology* **27**, 99–108.
- [3] Mason, R.L. & Young, J.C. (2002). *Multivariate Statistical Process Control with Industrial Applications*, ASA-SIAM, Philadelphia.
- [4] QualStat V5.0 (1996). *InControl Technologies, Inc.*, Houston.
- [5] Mason, R.L., Tracy, N.D. & Young, J.C. (1997). A practical approach for interpreting multivariate T^2 control chart signals, *Journal of Quality Technology* **29**, 396–406.
- [6] Hawkins, D.M. (1991). Multivariate quality control based on regression-adjusted variables, *Technometrics* **33**, 61–75.
- [7] Hawkins, D.M. (1993). Regression adjustment for variables in multivariate quality control, *Journal of Quality Technology* **25**, 170–182.
- [8] Roy, J. (1958). Step-down procedure in multivariate analysis, *Annals of Mathematical Statistics* **29**, 1177–1187.
- [9] Wierda, S.J. (1994). *Multivariate Statistical Process Control: Groningen Theses in Economics*, Wolters-Noordhoff, Groningen.
- [10] Doganaksoy, N., Faltin, F.W. & Tucker, W.T. (1991). Identification of out-of-control quality characteristics in a multivariate manufacturing environment, *Communications in Statistics: Theory and Methods* **20**, 2775–2790.
- [11] Murphy, B.J. (1987). Selecting out of control variables with the T^2 multivariate quality control procedure, *The Statistician* **36**, 571–583.
- [12] Fuchs, C. & Benjamini, Y. (1994). Multivariate profile charts for statistical process control, *Technometrics* **36**, 182–195.
- [13] Sparks, R.S., Adolphson, A. & Phatak, A. (1997). Multivariate process monitoring using dynamic biplot, *International Statistical Review* **65**, 325–349.
- [14] Kourtis, T. & MacGregor, J.F. (1996). Multivariate SPC methods for process and product monitoring, *Journal of Quality Technology* **28**, 409–428.
- [15] Maravelakis, P.E., Bersimis, S., Panaretos, J. & Psarakis, S. (2002). Identifying the out of control variable in a multivariate control chart, *Communications in Statistics: Theory and Methods* **31**, 2391–2408.
- [16] Nedumaran, G. & Pignatiello, J.J. (1998). Diagnosing signals from T^2 and χ^2 multivariate control charts, *Quality Engineering* **10**, 657–667.

Related Articles

Hotelling's T^2 Chart; Multivariate Control Charts Overview; Multivariate Exponentially Weighted Moving Average (MEWMA) Control Chart; Regression Control Charts.

ROBERT L. MASON AND JOHN C. YOUNG

Autocorrelated Data

Introduction

Traditional statistical process control (SPC) assumes that consecutive points plotted on a control chart are independent (*see Control Charts, Overview*). In the case that individual observations (rather than means computed from subgroups) are being plotted, then the typically assumed process can be described by the simple relationship

$$y_t = \mu + a_t \quad (1)$$

where

y_t = the quality variable of interest at time t

μ = the mean of the process

a_t = the random error at time t .

As early as the 1970s, however, researchers recognized that in practice observations are often serially autocorrelated (see [1–3] as examples). Observations that are autocorrelated are not independent and future observations can be predicted from those in the past. When a process is positively autocorrelated, the observations tend to follow each other or have inertia. Patterns that arise from such serial **autocorrelation** look almost sinusoidal in nature because of the similarity of consecutive observations. Observations that are negatively autocorrelated, on the other hand, show an alternating pattern, where subsequent observations move in the opposite direction to that of previous observations. In manufacturing settings, positive autocorrelation is by far the most commonly encountered deviation from independence. As the use of statistical monitoring tools spreads to new applications such as medical surveillance, however, more examples of negatively correlated process are likely to emerge.

Harris and Ross [4] reported that measured variables from tanks, reactors, and recycle streams in chemical processes show significant serial correlation and Wardell *et al.* [5, 6] motivated their work with an example from an extrusion process. Hence autocorrelation manifests itself in both continuous and discrete processes. The problem has been magnified in more recent times with the advent of automatic sensors. Such sensors allow for continuous measurement in

which the sampling frequency approaches infinity. Such a high sampling frequency almost always results in observations that are not independent, but instead autocorrelated (see [7] for example).

Under such conditions traditional SPC procedures may be ineffective, indeed inappropriate for monitoring, controlling, and improving process quality. A common result of applying standard Shewhart control charts (in particular, individuals and moving range charts) to processes exhibiting positively autocorrelated data is an increase in the number of false alarms, i.e., the number of times that a point falls outside of the control limits when in reality the process characteristics have remained stable. False alarms have the potential to be very costly, especially in environments in which operators are empowered to shut down the operation when a signal is generated from the control chart.

Figure 1(a) shows an individuals control chart applied to simulated data that mimic an AR(1) process with an autoregressive parameter equal to 0.65. Such a process is positively autocorrelated. For the first 50 observations, the mean and the process remained constant (i.e., no parameters changed). At point 51, a shift in the mean of two standard deviations (of the error terms) was introduced (indicated by the vertical line). The inside limits are those that were calculated by using the moving range to estimate the standard deviation. The figure shows multiple points that fall outside the inner limits, which are false alarms. The outer limits were formed by using the sample standard deviation of the first 50 points. No points fall outside those limits, including after the point of the shift. Hence there is a reduction in the number of false alarms, but at the expense of not detecting a true process change. Figure 1(b) will be discussed later.

Because of the difficulties of monitoring autocorrelated processes with traditional control charts, researchers and control chart designers sought ways to account for the autocorrelation. The next section introduces the two most commonly used approaches.

Common Approaches to Monitor Autocorrelated Data

Soon after the problem of autocorrelation was uncovered, researchers began to propose solutions that can broadly be divided into two categories. The first

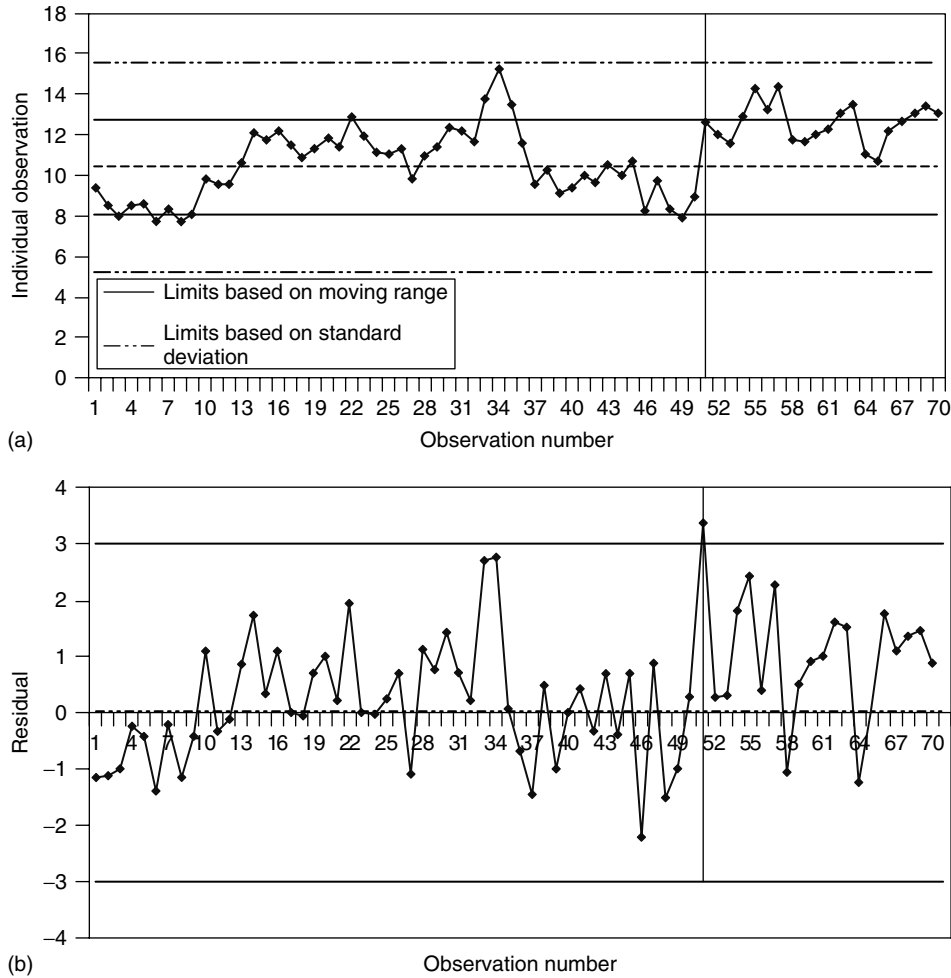


Figure 1 A comparison of a standard Shewhart chart to a residual chart when applied to a positively autocorrelated process. (a) The inner limits of the Shewhart chart are created by using a moving range and produce multiple false alarms. The outer limits are created using an overall standard deviation and do not detect shifts quickly (b) A residual chart typically shows a spike at the point of change and then settles down again after the change

approach was to adjust the control chart limits to compensate for the autocorrelation. In other words, if the problem was an excessive number of false alarms, then the control limits could be widened to reduce the false signals. Vasilopoulos and Stamboulis [3] adjusted Shewhart X-bar chart control chart constants (see **Control Charts for the Mean**) for the case of AR(2) processes. Others followed with similar studies of cumulative sum (CUSUM) (see **Cumulative Sum (CUSUM) Chart**) and exponentially weighted moving average (EWMA) (see **Exponentially Weighted Moving Average (EWMA) Control**

Chart) charts (see [8, 9] for example). The obvious problem with the limit-adjusting approach is that widening the control limits reduces the probability of detecting a true change in the process. Still, the approach is relatively simple to implement and often has better performance than other methods.

The second approach that emerged and gained considerable popularity was to fit the data using an appropriate statistical model and then compute and monitor the **residuals**, or unexplained portion of the model. Although others may have hinted at it earlier,

Alwan and Roberts [10] were the first to formalize the residual-monitoring approach. They introduced the possibility of two “charts”: a common-cause chart that showed the fitted values from the statistical model (but had no control limits); and a special-cause chart that was an individuals control chart of the residuals from the statistical model. If the optimal statistical model was found, the residuals would adhere to the traditional assumptions of SPC and hence a Shewhart control chart would be appropriate.

Alwan and Roberts were quite general in their suggestions of appropriate statistical models that could be used to account for the common-cause variation. They discussed regression models as well as time series, or Box–Jenkins models [11] (*see Time Series Analysis; Large Autoregressive Integrated Moving Average (ARIMA) Modeling*). The time-series approach was the one that was embraced by most subsequent researchers and it was almost forgotten that Alwan and Roberts had proposed a much more general idea.

The time-series approach called for fitting the model with a general autoregressive integrated moving-average (ARIMA) model (*see Large Autoregressive Integrated Moving Average (ARIMA) Modeling*). Most researchers who embraced the Alwan and Roberts approach assumed that the process was stationary and hence ignored the “integrated” portion of the ARIMA model, thereby reducing the analysis to somewhat simpler autoregressive moving-average (ARMA) models. We can write the general ARMA(p, q) model as shown in equation (2) [11]:

$$y_t = \mu(1 - \phi_1 - \cdots - \phi_p) + \phi_1 y_{t-1} + \cdots + \phi_p y_{t-p} + a_t - \theta_1 a_{t-1} - \cdots - \theta_q a_{t-q} \quad (2)$$

where

y_{t-i} = the quality variable of interest at time $t - i$,

$i = 0, 1, \dots, p$

μ = the mean of the process

ϕ_i = the autoregressive coefficients

θ_i = the moving-average coefficients

a_{t-j} = the random error at time $t - j$,

$j = 0, 1, \dots, q$.

The random errors are assumed to be independently and identically distributed (i.i.d.) with mean zero and variance σ_a^2 .

The application of the residual-chart approach requires the **estimation** of the autoregressive and moving-average parameters as well as the standard deviation of the error terms. Once those values are estimated, the model can be used to find fitted values of the process output, which are plotted on the common-cause chart. The difference between the actual observations and the fitted values is referred to as the *residual*, and this is what is plotted on the special-cause or residual chart.

Because of the relatively high sophistication required to fit time-series models to data, others proposed using simpler models to fit the data. For example, Montgomery and Mastrangelo [12] considered using the EWMA (*see Exponentially Weighted Moving Average (EWMA) Control Chart*) to fit the model, no matter what the “true” underlying model was. The downside of doing so was that the residuals were no longer completely independent of each other, but the reduced burden in terms of fitting the model was purported to compensate.

Residual-Chart Performance

While Alwan and Roberts [10] proposed special-cause or residual charts, they did not study the characteristics of such charts. Wardell *et al.* [5, 6] were among the first to fill the gap. By studying the run-length properties (*see Average Run Lengths and Operating Characteristic Curves*) of the special-cause chart, they showed that residual charts have poor performance when processes exhibit strong positive autocorrelation. Subsequent studies of EWMA (*see Exponentially Weighted Moving Average (EWMA) Control Chart*) and CUSUM (*see Cumulative Sum (CUSUM) Chart*) charts applied to the residuals helped [13, 14], but still suffered from the same inherent problem, which is described below. Note that the majority of comments below are true for processes displaying high levels of positive autocorrelation. For low to moderate levels of autocorrelation, residual-chart strategies fare better.

Unlike the performance of a Shewhart chart on independent data, when using residual charts to monitor processes displaying strong positive autocorrelation, if a change in the process is not detected

immediately, the probability of detecting the change subsequently becomes very small. When the shift occurs, the difference between the observed value and the fitted value is relatively large. Within a very short time, however, the statistical model or forecast absorbs the process change and the difference between the actual and fitted values becomes small again. Hence there is a very short window of opportunity over which the residual chart can detect changes in highly correlated data. If the change to the system is large, the residual chart is able to detect the shift quite easily during the short window of opportunity. For smaller changes, however, the shift is not detected during the short window and it therefore takes a very long time before the control chart signals.

Figure 1(b) shows the residual or special-cause chart applied to the residuals of the process described in the introduction. Recall that at observation 51, a shift in the mean was introduced to the simulated process. The figure reveals a spike at point 51 that is large enough to exceed the limits of the residual chart. Note that after the spike the observations settle down again and no other points fall above the limits. Such a pattern is typical of the residuals of a highly autocorrelated process when a mean shift is introduced [5]. If the mean shift is large enough, the residual chart will detect it immediately. If not, however, it typically takes a very long time to detect the shift. For this reason, some researchers [15] have applied hybrid monitoring schemes such as combinations of Shewhart charts and EWMA charts (*see Exponentially Weighted Moving Average (EWMA) Control Chart*) to improve monitoring performance. While these techniques provided improvement for low to moderate levels of positive autocorrelation, high levels of positive autocorrelation remain problematic.

For cases where the shift in the mean is small, more traditional charts (*see Exponentially Weighted Moving Average (EWMA) Control Chart; Cumulative Sum (CUSUM) Chart*) whose limits are adjusted to account for the autocorrelation are likely to have superior performance to that of the residual chart. The tendency for the statistical model to “catch up” to the actual observations makes it difficult for almost any model-based approach to discover process changes quickly. It is an enigma that has caused researchers to look for different ways to attack the problem of autocorrelation.

More Recent Approaches

The poor performance of residual charts at detecting process changes prompted many other studies to try to find charts with better detection capabilities. Some of the approaches were “incremental” changes to the two earlier types of charts, while others were more radical. The more radical approaches involved finding monitoring procedures that neither require parameter estimation nor require modeling [7, 16].

Besides looking for improved methods of monitoring autocorrelated data, studies have also expanded the contexts under which autocorrelated data are of interest. For example, a literature relating statistical monitoring and “engineering” process control has begun to grow [17] (*see Engineering Process Control; Feedback Control*). In these studies, process operators can take advantage of the autocorrelation and specific input variables to control the output variable. Residual charts can still be used in such instances, but the dynamics of the residuals are different depending on the type of control mechanisms used.

Other extensions have included the following:

- The assumption that the model (*see Large Autoregressive Integrated Moving Average (ARIMA) Modeling*) can be fit perfectly has been questioned and the effect of estimation error was studied [14, 18].
- A few researchers have included studies of other measures of interest, such as the variance [19] (*see Moving Range and R Charts; Control Charts for the Standard Deviation*), process capability [20] (*see Multivariate Charts for Variability*), and attribute data [21] (*see Control Charts for Attributes*).
- Studies of multivariate phenomena, including correlation across processes, have begun to emerge [22] (*see Multivariate Control Charts Overview; Hotelling’s T^2 Chart; Multivariate Exponentially Weighted Moving Average (MEWMA) Control Chart; Multivariate Cumulative Sum (CUSUM) Chart; Multivariate Charts for Variability; Profile Monitoring*).

Conclusions

Over the past 25 years or so we have learned a great deal about monitoring autocorrelated data.

Despite the advances, researchers have not been able to identify a “one-size-fits-all” approach that effectively monitors such data. Several important questions remain. A fundamental question that bears further investigation is whether or not we should concentrate on fitting the model to the process or the process to the model [23]. In other words, rather than spending time on trying to find the best way to fit the autocorrelated data, we may be better served by finding ways to eliminate the autocorrelation so that the data fit the traditional assumptions of SPC.

The inherent problem of residual charts makes them difficult to improve. New ways of thinking about the problem are needed. One recent approach that shows promise is the comparison of “context trees” and the use of the Kullback–Leibler statistic [16]. Other approaches that may bear fruit are nonparametric methods, visualization techniques, and integration with the change-point literature (*see Change-Point Methods*). Similarly, almost all of the activity to date has concentrated on monitoring the mean of the process. **Six Sigma** demands an understanding of and reduction of variation in the process. Financial models such as autoregressive conditional heteroscedastic (ARCH) models deal with cases where the variability shows autoregressive behavior. There may be applications of ARCH models in SPC. In any case, much remains to be done to understand the best method to monitor processes that exhibit autocorrelation.

References

- [1] Box, G., Jenkins, G. & MacGregor, J. (1974). Some recent advances in forecasting and control, part II, *Journal of the Royal Statistical Society. Series C* **23**, 158–179.
- [2] Ermer, D., Chow, M. & Wu, S. (1979). A time series control chart for a nuclear reactor, in *Proceedings of the 1979 Annual Reliability and Maintainability Symposium*, Institute of Electrical & Electronic Engineers, New York, pp. 92–98.
- [3] Vasilopoulos, A. & Stamboulis, A. (1978). Modification of control chart limits in the presence of data correlation, *Journal of Quality Technology* **10**, 20–30.
- [4] Harris, T. & Ross, W. (1991). Statistical process control procedures for correlated observations, *Canadian Journal of Chemical Engineering* **69**, 48–57.
- [5] Wardell, D., Moskowitz, H. & Plante, R. (1992). Control charts in the presence of data correlation, *Management Science* **38**, 1084–1105.
- [6] Wardell, D., Moskowitz, H. & Plante, R. (1994). Run-length distributions of special-cause control charts for correlated processes, *Technometrics* **36**, 3–17.
- [7] Runger, G. & Willemain, T. (1995). Model-based and model-free control of autocorrelated processes, *Journal of Quality Technology* **27**, 283–292.
- [8] Yashchin, E. (1993). Performance of CUSUM control schemes for serially correlated observations, *Technometrics* **35**, 37–52.
- [9] Zhang, N. (1998). A statistical control chart for stationary process data, *Technometrics* **40**, 24–38.
- [10] Alwan, L. & Roberts, H. (1988). Time-series modeling for statistical process control, *Journal of Business and Economic Statistics* **6**, 87–95.
- [11] Box, G. & Jenkins, G. (1976). *Time Series Analysis, Forecasting and Control*, 2nd Edition, Holden Day, San Francisco.
- [12] Montgomery, D. & Mastrangelo, C. (1991). Some statistical process control methods for autocorrelated data (with discussion), *Journal of Quality Technology* **23**, 179–204.
- [13] Runger, G., Willemain, T. & Prabhu, S. (1995). Average run lengths for CUSUM control charts applied to residuals, *Communications in Statistics-Theory and Methods* **24**, 273–282.
- [14] Adams, B. & Tseng, I. (1998). Robustness of forecast-based monitoring schemes, *Journal of Quality Technology* **30**, 328–339.
- [15] Lin, S. & Adams, B. (1996). Combined control charts for forecast-based monitoring schemes, *Journal of Quality Technology* **28**, 289–301.
- [16] Ben-Gal, I., Morag, G. & Shimlovici, A. (2003). Context-based statistical process control: a monitoring procedure for state-dependent processes, *Technometrics* **45**, 293–311.
- [17] Vander Wiel, S. (1996). Monitoring processes that wander using integrated moving average models, *Technometrics* **38**, 139–151.
- [18] Rajagopal, R. & Del Castillo, E. (2005). Model-robust process optimization using Bayesian model averaging, *Technometrics* **47**, 152–163.
- [19] Lu, C.-W. & Reynolds, M. (1999). Control charts for monitoring the mean and variance of autocorrelated processes, *Journal of Quality Technology* **31**, 259–274.
- [20] Shore, H. (1997). Process capability analysis when data are autocorrelated, *Quality Engineering* **9**, 615–626.
- [21] Alwan, L. (2003). General data analysis supported SPC, *Quality Engineering* **15**, 371–382.
- [22] Apley, D. & Tsung, F. (2002). The autoregressive T² chart for monitoring univariate autocorrelated processes, *Journal of Quality Technology* **34**, 80–96.
- [23] Hoerl, R. & Palm, A. (1992). Discussion: integrating SPC and APC, *Technometrics* **34**, 268–272.

DON G. WARDELL

Regression Control Charts

Introduction

Standard **control charts** have been widely used in process monitoring. They are often directly applied to the output of a quality characteristic. In practice, there are situations when the process output is affected by an external **covariate**, such as an incoming process variable. When the covariate takes an extreme value, the process output may appear to be in the out-of-control status even though there is actually no operation problem with the process. Intuitively, the results would be error prone in this situation when the standard control charts are directly applied. In contrast, the regression control chart proposed by Mandel [1], which combines regression analysis and control chart theory, can deal with this dilemma. The essential idea of the regression control chart is to chart the process output after it has been adjusted for the effect of the covariate.

The regression control chart has proved useful in a wide variety of applications due to its simplicity and efficiency. A least squares regression program must process the data prior to constructing the control chart. With the availability of modern computation power, the execution time of such a chart is often on the order of a few seconds. Moreover, it offers better sensitivity than a standard control chart that directly monitors the process output without taking the influence of the covariate into consideration. Among various applications, one of the most important application is for monitoring and diagnosing multistage processes because most products produced today are the result of such multistage processes. The stages are often intercorrelated in a multistage process, and the downstream variables depend on the upstream variables (see [2, 3]).

Similar concepts to regression control charts have been used for different purposes. In maximizing the sensitivity to a structured shift when performing multivariate **quality control**, Hawkins [4] proposed a control chart based on regression-adjusted variables. Hawkins [5] further modified his technique for applications in cascade processes by using only upstream observations as the independent variables in the regression adjustment. Hauck *et al.* [6] considered regression adjustment in the multivariate case. The cause-selecting control chart proposed by Zhang

[7] is an updated version of the regression control chart. Wade and Woodall [8] reviewed and extended it for simplifying the interpretation of out-of-control signals in a multistage process. Moreover, the regression control chart has a close relationship to the monitoring method used with linear profiles (*see*, e.g., [9–12]; **Profile Monitoring**). The former detects changes in the Y-intercept based on regression **residuals** while the latter monitors changes in the profile parameters, including the Y-intercept, slope and process variation. Both depend on regression analysis.

The regression control chart is a model-based monitoring procedure. In the above references, the model fitting is done by using ordinary least-squares (OLS) regression. The basic assumptions of OLS regression are normality, constant variance and independence of the observations. In some situations, these assumptions may not be always met. To remedy these assumptions, the model-building algorithm of **generalized linear models** (GLMs) can be applied. For example, Loredó *et al.* [13] incorporated the use of GLM in the regression adjustment framework for correlated data. The use of GLM for Poisson distributed counts in a regression adjustment procedure was explored in Skinner *et al.* [14, 15]. Jearkpaorn *et al.* [16, 17] discussed the use of GLM and robust GLM for a (γ -distributed response variable, respectively. Woodall [18] recently reviewed the potential applications of regression control charts in health-care and public-health surveillance in which the response variable is often non-normal.

The Regression Control Chart

Define X and Y to be the external covariate and the output characteristic of interest, respectively. The model relating X to Y can take many forms. For simplicity, the simple linear regression model

$$Y_i = \beta_0 + \beta_1 X_i + \epsilon_i \quad (1)$$

is used here for discussion. It is assumed that ϵ_i are independent and identically distributed (i.i.d.) normal variables with mean zero and variance σ^2 . In practice, the values of σ , β_0 , and β_1 are unknown and must be estimated using historical data collected from the process. We assume that a sample of n paired observations (X_i, Y_i) , are available in a set of historical data. In classical linear regression, the regressor variable, X , is assumed to be fixed. When

2 Regression Control Charts

the regressor variable is random, then the regression can be regarded as conditional on the value of the regressor (see **Least-Squares Estimation**; [19]).

Denote $\bar{X}$ and $\bar{Y}$ as the sample averages of the X and Y variables, respectively. Also, define $S_{XX} = \sum_{i=1}^n (X_i - \bar{X})^2$ and $S_{XY} = \sum_{i=1}^n (X_i - \bar{X})Y_i$. Then, the estimates of the Y-intercept, β_0 , and the slope, β_1 , can be obtained using the least-squares method (see, e.g., [20]), given by

$$\hat{\beta}_0 = \bar{Y} - \hat{\beta}_1 \bar{X} \text{ and } \hat{\beta}_1 = \frac{S_{XY}}{S_{XX}} \quad (2)$$

Denote the predicted value of Y as $\hat{Y} = \hat{\beta}_0 + \hat{\beta}_1 X$. The residual is the difference between the observed and predicted values; that is,

$$e_i = Y_i - \hat{Y}_i \quad (3)$$

The estimate of σ^2 is

$$\hat{\sigma}^2 = MSE = \frac{\sum_{i=1}^n (e_i - \bar{e})^2}{n - 2} \quad (4)$$

In the least-squares estimate, it follows that

$$\bar{e} = \sum_{i=1}^n e_i / n = 0 \quad (5)$$

The estimates $\hat{\beta}_0$, $\hat{\beta}_1$, and $\hat{\sigma}^2$ are unbiased.

The regression control chart is a Shewhart or other type of control chart applied to the regression residuals. The Shewhart regression control chart with three- σ control limits may be established as follows:

$$LCL = \bar{e} - 3\hat{\sigma}, \quad CL = \bar{e}, \quad UCL = \bar{e} + 3\hat{\sigma} \quad (6)$$

However, note that there is less accuracy in predicting a future Y value when the corresponding X value is further away from the sample mean $\bar{X}$ (see [20]). To account for this, it has been suggested that the parallel limits in equation (6) be revised to the prediction limits [8]

$$LCL = \bar{e} - k\hat{\sigma}_e, \quad CL = \bar{e}, \quad UCL = \bar{e} + k\hat{\sigma}_e \quad (7)$$

The variable

$$\hat{\sigma}_e = \sqrt{MSE \left(1 + \frac{1}{n} + \frac{(X_0 - \bar{X})^2}{S_{XX}} \right)} \quad (8)$$

is the **standard error** of the prediction for a given future observation of $X = X_0$. The constant multiplier, k , can be chosen as $k = t(\alpha/2; n - 2)$, where $t(\alpha/2; v)$ is the upper $\alpha/2$ percentile of the Student's t distribution with v **degrees of freedom**. Although the residuals are dependent due to parameter **estimation**, the simulation study performed by Wade and Woodall [8] shows that this choice can maintain an approximate false-alarm rate of α .

An Example

For illustration, the example on the ABCD Post Office Department in [1] is used here. The variables measured are X = the number of pieces of mail (in millions) handled in a four-week period and Y = the number of man hours (in thousands) used. Table 1 summarizes the data. The first 26 data points are used to estimate the regression parameters and calculate

Table 1 Data on mail processing hours and volume for the ABCD post office

Number	X	Y	Number	X	Y	Number	X	Y	e
1	157	572	14	154	569	27	162	623	31
2	161	570	15	157	564	28	174	664	32
3	168	645	16	164	573	29	171	664	42
4	186	645	17	188	667	30	182	680	21
5	183	645	18	191	700	31	175	650	14
6	184	671	19	180	765	32	169	582	-34
7	268	1053	20	270	1070				
8	180	675	21	180	637				
9	175	670	22	172	650				
10	193	710	23	184	655				
11	184	656	24	179	665				
12	179	640	25	169	599				
13	164	599	26	160	605				

the control limits while the last six are used for future monitoring purposes. As shown by [1], observations 7, 19, and 20 are abnormal. Therefore, they are excluded from the computations used to establish a preliminary control chart.

After excluding the three abnormal points, the fitted relationship between X and Y is

$$\hat{Y} = 50.44 + 3.345X \quad (9)$$

with

$$\hat{\sigma} = 18.92 \quad (10)$$

The next step is to compute the predicted Y values and the residuals for the last six observations, which are also given in Table 1. Given the results with $n = 23$, $S_{XX} = 3097.7$ and $\bar{X} = 174.4$ obtained from the revised data set (i.e., observations 7, 19, and 20 are removed), $\hat{\sigma}_e$ can be computed and then the prediction limits can be established.

For example, consider $\alpha = 0.05$ and the value of k is then set as $k = t(0.025; 21) = 2.08$. Figure 1(a) plots the regression control chart, and an out-of-control signal is triggered at observation 29. This signal is due to hiring a large number of inexperienced employees, as explained by Mandel [1]. For comparison, Figure 1(b) plots a Shewhart chart on the Y variable (i.e., manhours used). The control limits are placed at two estimated standard deviations of Y from the center line (i.e., $\bar{Y} \pm 2\hat{\sigma}_Y = 634 \pm 2 \times 44.78$), which can approximately provide a false-alarm rate of 5%. The Y chart in Figure 1(b) does not indicate an out-of-control situation. This example illustrates that better sensitivity on the regression control chart

can be achieved due to the adjustment for the effect of the external covariate compared to standard control charts.

Concluding Remarks

The regression control chart is a useful statistical process control (SPC) tool in many applications, including monitoring and diagnosing multistage processes. Compared with standard control charts, it may provide better sensitivity to changes in process output. Moreover, the regression control chart has the ability to distinguish between incoming and outgoing problems when it is used in combination with a control chart on the incoming variable. This facilitates the interpretation issue of an out-of-control signal in multistage process monitoring. In contrast, a multivariate control chart cannot make this distinction (see **Multivariate Control Charts Overview**).

An important step of the regression control chart is to estimate the model parameters from a historical data set. The parameter estimation leads to the dependence of the regression residuals although the observations on the output are originally independent. This dependency in turn causes deterioration of the performance of the chart, which has been investigated by Shu *et al.* [21, 22]. The minimum sample size required for the regression control chart to perform like the chart with known parameters is shown to be considerably larger than expected. For this reason, the minimum sample-size requirement may not be satisfied in real situations. Shu *et al.* [21] discussed some remedies for this problem without detailed

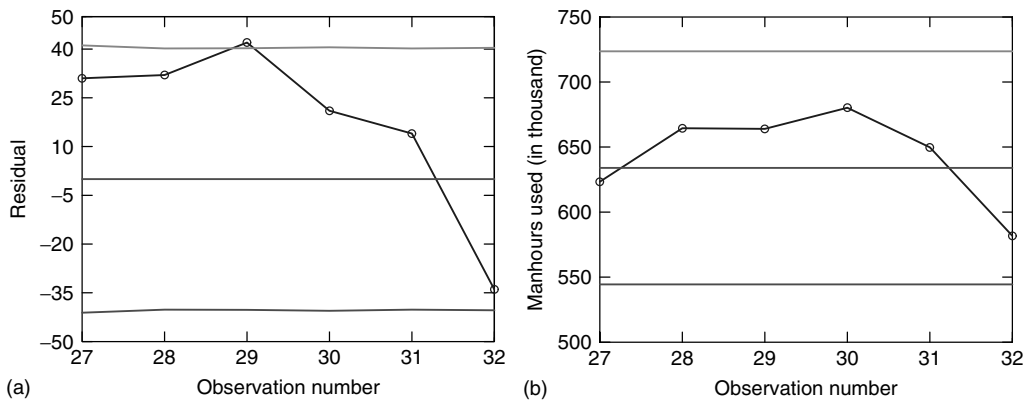


Figure 1 (a) A regression control chart and (b) a control chart for man hours

investigations, although much more work is needed on this issue.

Great efforts have been devoted to phase II methods of regression control charts while relatively less attention has been given to phase I analysis. In practical applications, it is often assumed that the historical data collected from the process is in control or reasons for assignable causes can be identified. However, there could be **outliers** in the data. In the existing phase I methods for a regression control chart, the control limits established to identify outliers are based purely on rules of thumb. The traditional use of three- σ limits would produce a false-alarm rate that would be very different from the desired rate of 0.0027. How to maintain the exact type I error probability in phase I of a regression control chart remains an important issue that requires more discussion. One possible remedy is to use the self-starting control chart [23, 24].

References

- [1] Mandel, B. (1969). The regression control chart, *Journal of Quality Technology* **1**, 1–9.
- [2] Zantek, P., Wright, G. & Plante, R. (2002). Process and product improvement in manufacturing systems with correlated stages, *Management Science* **48**, 591–606.
- [3] Xiang, L. & Tsung, F. (2007). Statistical monitoring of multistage processes based on engineering models, *IIE Transactions* (in press).
- [4] Hawkins, D. (1991). Multivariate quality control based on regression-adjusted variables, *Technometrics* **33**, 61–75.
- [5] Hawkins, D. (1993). Regression adjustment for variables in multivariate quality control, *Journal of Quality Technology* **25**, 170–182.
- [6] Hauck, D., Runger, G. & Montgomery, D. (1999). Multivariate statistical process monitoring and diagnosis with grouped regression-adjusted variables, *Communications in Statistics: Simulation and Computation* **28**, 309–328.
- [7] Zhang, G. (1992). *Cause-Selecting Control Chart and Diagnosis, Theory and Practice*, Aarhus School of Business, Department of Total Quality Management, Aarhus.
- [8] Wade, M. & Woodall, W. (1993). A review of cause-selecting control charts, *Journal of Quality Technology* **25**, 161–169.
- [9] Kang, L. & Albin, S. (2000). On-line monitoring when the process yields a linear profile, *Journal of Quality Technology* **32**, 418–426.
- [10] Kim, K., Mahmoud, M. & Woodall, W. (2003). On the monitoring of linear profiles, *Journal of Quality Technology* **35**, 317–328.
- [11] Mahmoud, M. & Woodall, W. (2004). Phase I analysis of linear profiles with calibration applications, *Technometrics* **46**, 380–391.
- [12] Woodall, W., Spitzner, D., Montgomery, D. & Gupta, S. (2004). Using control charts to monitor process and product profiles, *Journal of Quality Technology* **36**, 309–320.
- [13] Lored, E., Jearkpaorn, D. & Borror, C. (2002). Model-based control chart for autoregressive and correlated data, *Quality and Reliability Engineering International* **18**, 489–496.
- [14] Skinner, K., Montgomery, D. & Runger, G. (2003). Process monitoring for multiple count data using generalized linear model-based control charts, *International Journal of Production Research* **41**, 1167–1180.
- [15] Skinner, K., Montgomery, D. & Runger, G. (2004). Generalized linear model-based control charts for discrete semiconductor process data, *Quality and Reliability Engineering International* **20**, 777–786.
- [16] Jearkpaorn, D., Montgomery, D., Runger, G. & Borror, C. (2003). Process-monitoring for correlated gamma-distributed data using generalized linear model-based control charts, *Quality and Reliability Engineering International* **19**, 477–491.
- [17] Jearkpaorn, D., Montgomery, D., Runger, G. & Borror, C. (2005). Model-based process monitoring using robust generalized linear models, *International Journal of Production Research* **43**, 1337–1354.
- [18] Woodall, W. (2006). The use of control charts in health-care and public-health surveillance, *Journal of Quality Technology* **38**, 89–104.
- [19] Seber, G. (1977). *Linear Regression Analysis*, John Wiley & Sons, New York.
- [20] Montgomery, D. & Peck, E. (1982). *Introduction to Linear Regression Analysis*, John Wiley & Sons, New York.
- [21] Shu, L., Tsung, F. & Tsui, K. (2004). Run-length performance of regression control charts with estimated parameters, *Journal of Quality Technology* **36**, 280–292.
- [22] Shu, L., Tsung, F. & Tsui, K. (2005). Effects of estimation errors on cause-selecting charts, *IIE Transactions* **37**, 559–567.
- [23] Quesenberry, C. (1997). *SPC Methods for Quality Improvement*, John Wiley & Sons, New York.
- [24] Zantek, P., Wright, G. & Plante, R. (2006). A self-starting procedure for monitoring process quality in multistage manufacturing systems, *IIE Transactions* **38**, 293–308.

LIANJIE SHU, KWOK-LEUNG TSUI AND FUGEE
TSUNG

Control Charts, Nonparametric

Introduction

Distribution of Chance Causes

In statistical process control (SPC), the pattern of chance causes is usually assumed to follow some parametric distribution (such as the normal). The charting statistic and the control limits depend on this assumption and as such the properties of these **control charts** are “exact” only if this assumption is satisfied. However, the chance distribution is either unknown or far from being normal in many applications and consequently the performance of standard control charts can be highly affected in such situations. Thus, there is a need for some easy to use, flexible and robust control charts that do not require normality or any other specific parametric model assumption about the underlying chance distribution. Distribution-free or nonparametric control charts can serve this broader purpose.

Nonparametric or Distribution-Free

The term *nonparametric* is not intended to imply that there are no parameters involved, in fact, quite the contrary. While the term *distribution-free* seems to be a better description of what we expect from these charts, that is, they remain valid for a large class of distributions, nonparametric is perhaps the term more often used. In the statistics literature, there is now a rather vast collection of **nonparametric tests** and **confidence intervals**, and these methods have been shown to perform well compared to their normal theory counterparts. Remarkably, even when the underlying distribution is normal, the efficiency of some nonparametric methods relative to the corresponding (optimal) normal theory methods can be as high as 0.955 (see, e.g., Gibbons and Chakraborti [1]). In fact, for some heavy-tailed distributions like the double exponential, nonparametric tests can be more efficient. It may be argued that nonparametric methods will be “less efficient” than their parametric counterparts when one has a complete knowledge of the process distribution for which that parametric method was specifically designed. However, the reality is

that such information is seldom, if ever, available in practice. Thus, it seems natural to develop and use nonparametric methods in SPC and the quality practitioners will be well advised to have these techniques in their tool kits.

We mostly discuss univariate nonparametric control charts designed to track the location of a continuous process; very few charts are available for scale. The field of multivariate control charts (*see **Multivariate Control Charts Overview***) is interesting and the body of literature on nonparametric multivariate control charts is growing. However, in our opinion, it has not yet reached a critical mass and a discussion on this topic is better postponed for the future.

Nonparametric Control Charts

Chakraborti, van der Laan, and Bakir (hereafter CVB) [2] provided a systematic and thorough account of the nonparametric control charts literature. A nonparametric control chart is defined in terms of its in-control run length distribution. If the in-control run length distribution of a control chart is the same for every continuous distribution, the chart is called *nonparametric* or *distribution-free*. CVB summarized the advantages of nonparametric control charts as follows: (a) simplicity, (b) no need to assume a particular parametric distribution for the underlying process, (c) the in-control run length distribution is the same for all continuous distributions (the same is true for the false alarm rate (FAR) and the in-control **average run length** (ARL_0)), (d) more robust and **outlier** resistant, (e) more efficiency in detecting changes when the true distribution is markedly non-normal, particularly with heavier tails, and (f) no need to estimate the variance to set up charts for the location parameter. It is emphasized that from a technical point of view most nonparametric procedures require the population to be continuous in order to be distribution-free and thus in an SPC context we consider the so-called variables control charts.

Terminology and Formulation

Two important problems in usual SPC are monitoring the process mean (*see **Control Charts for the Mean***) and/or the process standard deviation (*see **Control Charts for the Standard Deviation***). In the nonparametric setting, we consider, more generally, monitoring the center or the location parameter

and/or a scale parameter of a process. The location parameter represents a typical value and could be the mean or the median or some other percentile of the distribution; the latter two are especially attractive when the underlying distribution is expected to be skewed. When the underlying distribution is symmetric, the mean and the median are the same and any of these can serve as a measure of the typical. Also in the nonparametric setting, the processes are often implicitly assumed to follow (a) a location model, with a cdf $F(x - \theta)$, where θ is the location parameter or (b) a scale model, with a cdf $F(x/\tau)$, where $\tau (> 0)$ is the scale parameter. Even more generally, one might consider (c) the location-scale model with cdf $F\{(x - \theta)/\tau\}$, where θ and τ are the location and the scale parameter, respectively. Under these model assumptions, the problem is to track θ or τ (or both), based on random samples or subgroups taken (usually) at equally spaced time points. In the usual (parametric) control charting problems F is assumed to be the cdf Φ of the standard **normal distribution** whereas in the nonparametric setting, for variables data, F is some unknown continuous cdf. Although the location-scale model seems to be a natural model to consider paralleling the normal case with mean and variance, both unknown, most of the currently available literature deals mainly with the location model.

As we noted earlier, a comprehensive survey of the literature until about 2000 can be found in CVB. Here, we briefly mention some of the key contributions and ideas and a few of the more recent developments in the area; the literature on nonparametric methods continues to grow at a rapid pace. Most nonparametric charts have been developed for phase II applications. There are generally two phases in SPC. In phase I (also called the *retrospective phase*), typically, preliminary analysis is done which includes planning, administration, data collection, data management, exploratory work including graphical and numerical analysis, goodness-of-fit analysis, and so on, to ensure that the process is in-control. This means that the process is operating at or near some acceptable target value along with some natural variation and no special causes of concern are present. Once this is ascertained, SPC moves to the next phase, phase II, (or the *prospective phase*), where the control limits and/or the estimators obtained in phase I are used for process monitoring based on new samples of data. When the underlying parameters of the process distribution are known or specified, this

is referred to as the *standard(s) known* case and is denoted case K. On the contrary, if the parameters are unknown and need to be estimated it is typically done in phase I, with in-control data. This situation is referred to as the *standard(s) unknown* case and is denoted case U.

There are three main classes of control charts: the Shewhart chart, the cumulative sum (CUSUM; *see Cumulative Sum (CUSUM) Chart*) chart, and the exponentially weighted moving average (EWMA; *see Exponentially Weighted Moving Average (EWMA) Control Chart*) chart and their refinements. Relative advantages and disadvantages of these charts are well documented (see, e.g., Montgomery [3]) in the literature. Analogs or types of these charts have been considered in the nonparametric setting. We describe some of the charts under each of the three classes.

Shewhart-Type Charts

Shewhart-type charts are the most popular in practice because of their simplicity, ease of application, and the fact that these versatile charts are quite efficient in detecting moderate-to-large shifts. Both one-sided and two-sided charts are considered. The one-sided charts are more useful when only a directional shift (higher or lower) in the median is of interest. The two-sided charts on the other hand are typically used to detect a shift or change in the median in any direction.

Case K.

Charts for Location.

Sign control charts. Sign control charts are based on the well-known sign test. The sign test is one of the simplest and most broadly applicable nonparametric tests (see, e.g., Gibbons and Chakraborti [1]) that can be used to test hypotheses (or find confidence intervals) on the median (or any specified percentile) of a continuous population. Suppose that the median of a continuous process needs to be maintained at a specified value θ_0 . Amin *et al.* [4] presented Shewhart-type nonparametric charts for this problem using what are called *within group sign* statistics. This is called the *SN chart*. Let $X_{i1}, X_{i2}, \dots, X_{in}$ be the i th ($i = 1, 2, \dots$) sample (or subgroup) of

independent observations of size $n > 1$. A control chart is a graphic consisting of values of a plotting (or charting) statistic and the associated control limits. The plotting statistic for the SN chart is $SN_i = \sum_{j=1}^n \text{sign}(X_{ij} - \theta_0)$. In order to find the control limits and study chart performance, the distribution of SN_i is necessary. This can be most easily obtained using the relationship $SN_i = 2T_i - n$, where T_i is the usual sign test statistic, which is a count of the number of sample observations greater than or equal to θ_0 . Since the in-control distribution of T_i is binomial with parameters n and $1/2$, it follows that the in-control distribution of SN_i is symmetric about 0 and hence the control limits and the centerline of the two-sided nonparametric Shewhart-type sign chart are: $UCL = c$, $CL = 0$, $LCL = -c$, where c is some positive integer between 0 and n . If the plotting statistic SN_i falls between the control limits, the process is declared to be in-control, whereas if the plotting statistic SN_i falls on or outside one of the control limits, that is, if $SN_i \leq -c$ or $SN_i \geq c$, the process is considered to be out-of-control. The charting constant c is typically obtained for a specified ARL_0 , which, in case K, is equal to the reciprocal of the nominal FAR, α . Thus, using the symmetry of the binomial distribution, c is the smallest integer such that $P_{\theta_0}(SN_i \geq c) \leq \alpha/2$. For example, using Table G of Gibbons and Chakraborti [1] we give some $P_{\theta_0}(T \geq t)$ values in Table 1 that may be considered “small” in a SPC context. The charting constant c is obtained using $c = 2t - n$. For example, for $n = 5$ we get $t = 5$ and thus $c = 5$ for a FAR of $2(0.0312) = 0.0624$ and this is the lowest FAR achievable. However, for $n = 10$ the FAR drops to 0.0020 if $c = 10$.

Looking at the achievable FAR values shown in the second row of Table 1, we see that unless the

sample size n is 10 or more, the sign chart could be somewhat unattractive (from an operational point of view) in SPC applications, where one often stipulates a FAR as small as 0.0027. On the other hand, the sign chart is the simplest of nonparametric charts that works under minimal assumptions. In fact, from the **hypothesis testing** literature, it is known that the sign test (and so the chart) is more robust and efficient when the chance distribution is symmetric like the normal but with heavier tails such as the double exponential.

Example 1 We illustrate the Shewhart-type sign chart using a set of data from Montgomery [3, Tables 5.1 and 5.2] on the inside diameters of piston rings manufactured by a forging process. A part of this data, 15 prospective samples (Table 5.2) each of five observations, is used here. The rest of the data (Table 5.1) will be used later. We assume that the underlying distribution is symmetric with a known median $\theta_0 = 74$ mm. Table 2 shows the SN_i values.

The sign chart is shown in Figure 1 with control limits at ± 5 . Observations 12, 13 and 14 lie on the upper control limit (UCL) which indicates that the process is out-of-control starting at sample 12. It appears most likely that the process median has shifted upwards from the in-control value of 74 mm.

Control charts are often compared on the basis of various characteristics of the run length distribution, such as the ARL. One prefers a longer ARL_0 and a shorter out-of-control ARL under a shift. Amin *et al.* [4] compared the ARL of the classical Shewhart $\bar{X}$ chart and the Shewhart-type sign chart for various shift sizes and underlying distributions. Generally speaking, when the distribution is either asymmetric

Table 1 FAR and ARL_0 of a sign control chart for various values of n and t

n	5	6	7	8	9	10
t	5	6	7	8	9	10
$P_{\theta_0}(T \geq t)$	0.0312	0.0156	0.0078	0.0039	0.0020	0.0010
FAR (α)	0.0624	0.0312	0.0156	0.0078	0.0040	0.0020
ARL_0	16.00	32.00	64.00	128.00	256.00	512.00

Table 2 Sign statistics (SN_i) for the piston-ring data in Montgomery [3]

Sample number	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
SN_i	2	1	-4	3	0	3	3	-1	3	4	1	5	5	5	4

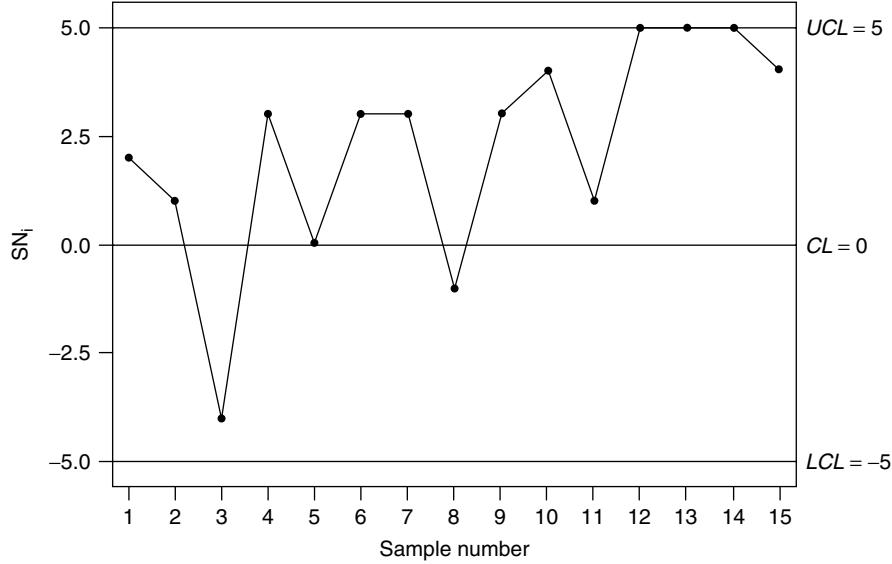


Figure 1 Shewhart-type sign control chart for Montgomery [3] piston-ring data

or symmetric with heavy tails, the sign charts are more efficient while the reverse is true for the normal and the normal-like distributions with light tails. However, one practical advantage of the sign charts (and of all nonparametric charts) is that the FAR (and the ARL_0) remains the same for all continuous distributions. This is not true for any parametric chart designed for a specific distribution, including the Shewhart $\bar{X}$ chart.

Signed-Rank Control Charts. A second popular nonparametric test for the median of a continuous and symmetric population is the signed-rank (SR) (or the Wilcoxon SR) test. The SR test is more powerful than the sign test but while the sign test is applicable for all continuous distributions, the assumption of symmetry must be made, in addition, for the SR test. Moreover, the sign test applies to all **percentiles** but the SR test is proposed only for the median. Details about the SR test can be found, for example, in Gibbons and Chakraborti [1]. Bakir [5] proposed a Shewhart-type nonparametric control chart based on the SR test statistic. This is called the *SR chart*. Let $|X_{i1} - \theta_0|, |X_{i2} - \theta_0|, \dots, |X_{in} - \theta_0|$ denote the absolute values of the deviations (from θ_0) for the i th subgroup and let R_{ij} denote the rank of $|X_{ij} - \theta_0|$ among the n absolute deviations. The plotting statistic of the SR chart is given by

$SR_i = \sum_{j=1}^n \text{sign}(X_{ij} - \theta_0) R_{ij}, i = 1, 2, 3, \dots$. Note that the statistic SR_i is the difference between the sum of the ranks (of the absolute values) corresponding to the positive and the negative deviations, respectively. This is linearly related to the better-known Wilcoxon SR statistic W^+ through the relationship $SR = 2W^+ - (n(n+1)/2)$, where W^+ is the sum of the ranks of the absolute values corresponding to the positive deviations. Since W^+ is known to be distribution-free [1] so is SR_i and the SR chart. The in-control distribution of SR_i is symmetric about 0, so choosing $LCL = -UCL$ seems reasonable and like the Shewhart-type sign chart this results in a (visually appealing) symmetric control chart. Thus as in the case of the sign chart, the control limits and the centerline of the two-sided Shewhart-type SR chart are: $UCL = d, CL = 0, LCL = -d$, where d is some positive integer. If the plotting statistic SR_i for some subgroup falls on or outside of the control limits, the process is declared to be out-of-control at that subgroup. Otherwise, the process is deemed to be in-control. The constant d is found so that a specified FAR (or equivalently an ARL_0) is attained. Bakir [5] calculated the FAR (p_0^+) and the in-control average run length (ARL_0^+) of the upper *one-sided* Shewhart-type signed-rank (SR^+) control chart (process out-of-control if $SR_i \geq UCL = d$) by using the null distribution of the Wilcoxon SR statistic. For

the two-sided SR chart, the FAR and the ARL_0 can be obtained from $p_0 = 2p_0^+$ and $ARL_0 = ARL_0^+/2$, respectively.

The SR chart competes well with the traditional Shewhart $\bar{X}$ chart in case of the normal distribution and is more efficient in case of a heavy-tailed distribution such as the double exponential or the Cauchy. However, the FAR values for the chart are too high (in other words the ARL_0 values are too short) unless the subgroup size n is large, say in the neighborhood of 20. This may not be a problem in applications where $n = 20$ is still considered “small”, but $n = 20$ is not typical in **quality control** applications and thus such charts, like the SN charts, can be somewhat unappealing from a practitioner’s point of view. One way to remedy this problem is to use some signaling rules to enhance the sensitivity of the charts. We consider this next.

Runs Rules, Signaling Rules, and Nonparametric Charts. A “run” is generally used to describe a sequence or succession of objects or observations of the same type in an ordered sequence. For example, if we have a signal for the first time on the fifth sample, there must have been a total of four no signals before the first signal from the first four samples and this constitutes a run of four no signals. In the SPC setting, runs have been typically employed to increase the sensitivity of control charts. Various rules or “tests for special causes” have been considered in the literature for parametric control charts; see for example, the rules associated with the Shewhart control chart in Nelson [6] and in the Western Electric Handbook [7]. See also the discussion in Montgomery [3].

Sign Control Charts with Warning Limits. In the nonparametric SPC literature, Amin *et al.* [4] considered Shewhart-type sign charts with warning limits and runs rules. The two-sided chart with action (control limits) at $\pm a_2$ is defined as follows. The centerline is 0. The warning limits are drawn at $\pm w_2$. A signal is given if r consecutive plotting statistics fall in the region $[w_2, a_2)$ or $(-a_2, -w_2]$, or if any plotting statistic falls outside $\pm a_2$. The average run length can be calculated from the average run lengths for the two one-sided charts. First, using results in Page [8], it can be shown that the ARL_0 of the one-sided (upper or positive direction) chart with warning limit at $w_2 (0 \leq w_2 \leq a_2)$ and control limit (action

limit) at a_2 is given by

$$L_0^+(\theta) = \frac{1 - p_1^r}{1 - p_1 - p_0(1 - p_1^r)} \quad (1)$$

where $p_0 = P(SN_i < w_2|\theta)$, $p_1 = P(w_2 \leq SN_i < a_2|\theta)$. Similarly, the ARL_0 of a one-sided (lower or negative direction) chart with warning limit at $-w_2$ and control limit (action limit) at $-a_2$, $L_0^-(\theta)$, can be found by replacing p_0 and p_1 by q_0 and q_1 where $q_0 = P(SN_i > -w_2|\theta)$ and $q_1 = P(-a_2 < SN_i \leq -w_2|\theta)$. The ARL_0 for the two-sided chart can then be obtained using a result in Roberts [9],

$$\frac{1}{L_0(\theta)} = \frac{1}{L_0^-(\theta)} + \frac{1}{L_0^+(\theta)} \quad (2)$$

In practice some observations can be tied with the specified median. If the number of such cases is small (relative to n) one can drop the tied cases and reduce n accordingly. On the other hand, if the number of ties is large, more sophisticated analysis might be necessary.

These procedures cannot be meaningfully illustrated using the data of Example 1 because the sample size $n = 5$ used there is too small. It may be noted that the highest possible $L_0(\theta)$ values for the basic (without warning limits) two-sided sign chart can be seen to be 2^{n-1} . Thus, achievable values of $L_0(\theta)$ are too small for practical use, unless n is about 10. In Table 3, the charting constants, i.e. the warning and action limits, are shown, along with the achieved ARL values, for the in-control and one out-of-control case. We compare three sign charts, two with warning limits and one without. The ARL values for the two-sided sign chart, without the warning limits, are shown in each case, within parentheses, for reference.

It is seen that adding warning limits to a control chart decreases its average run length. For example, adding a warning limit at 7 to the basic sign chart with an action limit at 10 decreases the ARL_0 approximately 59% (from 512 to 208.97), when $r = 2$, 7% (from 512 to 476.03) when $r = 3$ and 0.002% (from 512 to 511.99) when $r = 6$, respectively. The out-of-control average run length is decreased by approximately 79% (from 162.6 to 35.03) when $r = 2$, 36% (from 162.6 to 103.28) when $r = 3$, and 0.26% (from 162.6 to 162.17) when $r = 6$, respectively. Note that although the out-of-control average run length is reduced significantly (which means a quicker detection of shift) by the addition of warning

6 Control Charts, Nonparametric

Table 3 In-control ARL values for the two-sided sign chart with and without warning limits for $n = 10$

	$r = 2$	$r = 3$	$r = 6$
$w = 7$ and $a = 10$			
$p = 0.5$ (in-control)	208.97 (512)	476.03 (512)	511.99 (512)
$p = 0.6$ (out-of-control)	35.03 (162.6)	103.28 (162.6)	162.17 (162.6)
$w = 7$ and $a = 8$			
$p = 0.5$ (in-control)	42.86 (46.55)	46.37 (46.55)	46.55 (46.55)
$p = 0.6$ (out-of-control)	16.37 (20.82)	20.16 (20.82)	20.82 (20.82)

limits, the ARL_0 is also reduced significantly. This poses a dilemma in practice, since it is desirable to have a high ARL_0 or a low FAR, so one would need to strike a balance. One possibility is to use warning limits closer to the action limits. For example, from the second panel of Table 3, we see that adding a warning limit at 7 to the sign chart with an action limit of 8, decreases the ARL_0 by only 8% (from 46.55 to 42.86) when $r = 2$ and has little effect on ARL_0 when r is reasonably large. Amin *et al.* [4] concluded that for the upper one-sided Shewhart-type sign chart, introduction of warning limits will have little effect on the ARL_0 , but can significantly reduce the out-of-control ARL for small shifts in θ when the warning limits are chosen close to the action limits and r is reasonably large. Similar conclusions are expected to hold for two-sided charts.

In addition to defining warning limits or zones on control charts, runs have been used to define signaling rules in conjunction with usual control limits. The rules considered include the following: A process is declared to be out-of-control when (a) a single point (charting statistic) plots outside the control limit(s) (1-of-1 rule), (b) k consecutive points (charting statistics) plot outside the control limit(s) (k -of- k rule), or (c) k of the last w points (charting statistics) plot outside the control limit(s) (k -of- w rule). Rule (a) is the simplest and is the most frequently used in the literature. Thus, the 1-of-1 rule corresponds to the usual control chart, where a signal is given when a plotting statistic falls outside the control limit(s). Rules (a) and (b) are special cases

of rule (c); rules (b) and (c) have been used in the context of supplementing the Shewhart charts with warning limits and zones.

Signed-Rank Control Charts with Runs-Type Signaling Rules. Chakraborti and Eryilmaz [10] considered simple alternatives to Bakir's [5] class of nonparametric charts, using the SR statistic but incorporating runs rules of the type discussed above to define new signaling rules. These are called the 2-of-2 DR and the 2-of-2 KL charts, both based on the signed-rank statistic SR. The 2-of-2 KL chart signals, for example, when two of the most recent SR statistics both fall either above or below the control limits. The 2-of-2 DR chart is similar, here a signal is indicated when the SR statistics fall either both above or both below or one above (below) and the next one below (above) the control limits. It is shown that the new charts are nonparametric, have much smaller FAR (and thus longer ARL_0) than the 1-of-1 SR chart of Bakir. Moreover, the new charts have better out-of-control performance for some distributions like the normal and the Cauchy. We illustrate these procedures using the same data used in our earlier example.

Example 2 Table 4 shows the signed-rank (SR) statistics for the 15 samples.

The symmetric two-sided control limits for the SR charts are obtained from Chakraborti and Eryilmaz [10]. For $n = 5$, the control limits for the 1-of-1 (Bakir's chart), the 2-of-2 DR and the 2-of-2 KL charts, based on the SR statistic, are set at ± 15 . These

Table 4 Signed-rank statistics (SR_i) for the piston-ring data in Montgomery [3]

Sample number: i	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
SR_i	8	4	-14	7	-3	9	10	-6	12	14	4	15	15	15	14

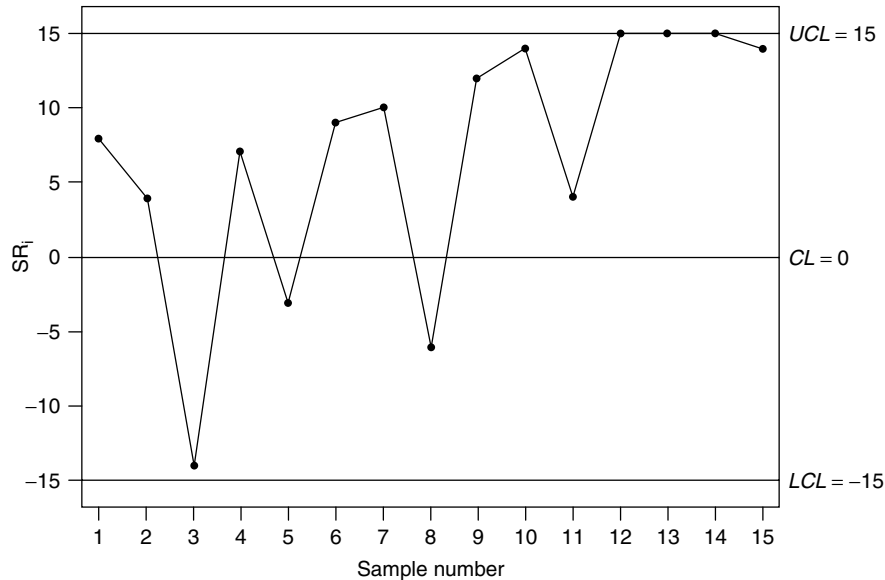


Figure 2 Shewhart-type signed-rank control chart for Montgomery [3] piston-ring data

yield FAR values 0.0626, 0.0039, and 0.0019, respectively. If the control limits were taken to be ± 13 , the FAR would have been much higher: 0.125, 0.0156, and 0.0078, respectively. Although the control limits are the same, namely ± 15 , the signaling rules are quite different operationally and the performance of the resulting charts turn out to be quite different. The 1-of-1 chart signals when the first SR statistic falls outside either of the two control limits; the 2-of-2 KL chart signals when, for the first time, two consecutive SR statistics fall either above or below the two control limits, while the 2-of-2 DR chart signals when, for the first time, two consecutive SR statistics fall outside the control limits, either both above or both below, or one above and the next below, or one below and the next above. On the performance side, note that the 1-of-1 SR chart has a FAR of 0.0626 and an ARL_0 of only 16. Thus many more false alarms will be signaled by this chart leading to a possible loss of time and resources. Compared to that, the 2-of-2 KL chart has a FAR of 0.0019 and an ARL_0 of 526.34, whereas the 2-of-2 DR chart has a FAR of 0.0039 and an ARL_0 of 271.15. Thus both of these run-rule-enhanced charts provide reasonable and practical FARs and can be used in practice, depending on the type of shift one expects.

Figure 2 shows the two-sided control chart with

the symmetric control limits and a plot of the SR statistics.

Thus the DR and KL 2-of-2 SR charts both signal at sample 13, indicating a most likely upward shift in the process median. The 1-of-1 SR chart, on the other hand, signals earlier at sample 12, but note the much higher FAR of 0.0626 (and correspondingly a much lower and less desirable ARL_0 , 15.97) associated with this chart. It is interesting to note that, as shown in Montgomery [3], for these data the Shewhart $\bar{X}$ chart indicates a shift in the mean at sample 11 for these data. However, the key difference is that an application of the Shewhart chart can raise several questions such as the form of the underlying distribution (small $n = 5$), and more importantly about the in-control (stable) performance of the chart in terms of the FAR (or the ARL_0), since it is known that the in-control performance of the Shewhart $\bar{X}$ chart is not robust in typical quality control applications. Compared to this, the proposed nonparametric charts provide a more generally applicable alternative monitoring scheme with a known (stable/robust) in-control performance and a better or equal out-of-control performance.

Charts for Scale. Amin *et al.* [4] also considered a Shewhart-type nonparametric sign chart for the variability based on counts that fall between or outside some prechosen **quartiles**, such as the first and the

third, that are assumed known. They compared their chart to the chart based on S^2 and suggested that the chart based on S^2 is more efficient, but of course the chart based on S^2 is not distribution-free. The overall conclusion is that the nonparametric charts provide a useful alternative to the standard charts when normality is in doubt. More work on nonparametric charts for scale or dispersion will be useful.

In general, control charts (and particularly nonparametric control charts) are simpler to use in case K. In practice, however, the process parameters are sometimes unknown or unspecified. We refer to this situation as case U; a phase I analysis is typically used here first to establish the process in-control and then control limits are calculated from a set of data from the in-control process having estimated any unknown parameters. Such data are referred to as *reference data*. The estimated control limits are then used for prospective process monitoring in phase II. We now discuss some nonparametric control charts useful in case U.

Case U.

Charts for Location. Janacek and Meikle [11] proposed a phase II nonparametric control chart useful in case U. The control limits of this chart are given by two selected order statistics of a phase I reference sample and the charting statistic is the median M_i of the phase II samples taken sequentially. More details are given below.

Precedence Control Charts. Chakraborti, van der Laan, and van de Wiel (hereafter CVV) [12] generalized the work of Janacek and Meikle [11]. They considered using some order statistic of a phase II sample as the charting statistic and control limits constructed from a phase I reference sample. Their work involves a class of two-sample nonparametric statistics, called *precedence statistics* and their Shewhart-type charts are called *precedence charts*. Assume that a reference sample of size m is available from an in-control process. The control limits of the precedence chart are given by two reference sample order statistics, say, $LCL = X_{a:m}$ and $UCL = X_{b:m}$, where $1 \leq a < b \leq m$. The plotting statistic $Y_{j:n}^h$ is the j th order statistic from the h th phase II sample of size n . The plotting statistic is taken to be the median for subsequent illustration but, in fact, it can be any percentile of the phase II sample. CVV provided

recommendations and tables for the implementation of precedence charts and examined the chart performance in terms of the ARL. The overall conclusion is that the Shewhart-type precedence charts are much more robust than their parametric counterparts. The precedence chart, being nonparametric, has the in-control robustness property (such as the same ARL_0 or the FAR for all continuous distributions), whereas as we noted earlier, the Shewhart $\bar{X}$ (and other parametric charts) can suffer serious degradations when the underlying distribution is not as assumed (normal).

Precedence Charts with Signaling Rules. Chakraborti *et al.* [13] considered enhancing the precedence charts with the 2-of-2 type signaling rules. We illustrate these procedures using the piston-ring data from Montgomery [3].

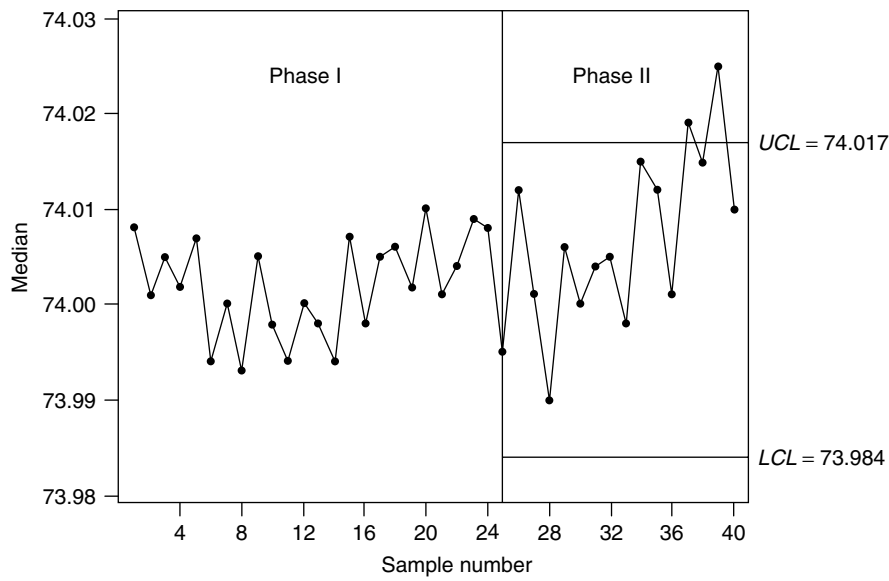
Example 3 Here we use the data given in Table 5.1, of Montgomery [3]. These are the 25 retrospective or phase I samples, each of size five that were collected when the process was thought to be in-control. The traditional Shewhart $\bar{X}$ and R charts provide no indication of an out-of-control condition, so these “trial” limits were adopted for use in on-line process control. These data are considered to be phase I reference data.

In order to implement the control charts, the charting constants are needed. Generally, one finds the chart constants so that a specified ARL_0 , such as 500, or 370, is obtained. For the precedence type charts, symmetric control limits are used so that $b = m - a + 1$ and only one charting constant a (integer, ≥ 1) needs to be found. Possible control limits for the three charts are shown in Table 5 for $m = 125$, $n = 5$ and $j = 3$, along with the corresponding FAR and ARL_0 values. The basic Shewhart-type precedence chart is referred to as the *1-of-1 chart*.

Thus, for an ARL_0 of 500, one can take $a = 7$ and $b = 119$ so that the control limits for the 1-of-1 precedence chart are the 7th and the 119th ordered values of the reference sample. Thus $LCL = X_{7:125} = 73.984$ and $UCL = X_{119:125} = 74.017$, which yield an in-control average run length of 413.80 and a FAR of 0.0044. A plot of the medians for the 1-of-1 chart is shown in Figure 3. It is seen that the 1-of-1 precedence chart signals on the 12th sample in the prospective phase.

Table 5 In-control average run length (ARL_0), false alarm rate (FAR), and chart constant (a) for the 1-of-1, 2-of-2 DR, and 2-of-2 KL precedence charts when $m = 125$, $n = 5$, and $j = 3$

1-of-1			2-of-2 DR			2-of-2 KL		
a	ARL_0	FAR	a	ARL_0	FAR	a	ARL_0	FAR
5	1315.98	0.001865	19	464.38	0.004020	19	819.47	0.002355
6	695.09	0.002948	20	344.73	0.005195	20	608.81	0.003019
7	413.80	0.004358	21	260.69	0.006627	21	460.54	0.003823
8	267.40	0.006164	22	200.46	0.008356	22	354.09	0.004788

**Figure 3** 1-of-1 phase II precedence chart for the Montgomery [3] piston-ring data

For the 2-of-2 DR chart, take $a = 19$ so that $b = 125 - 19 + 1 = 107$ and the resulting limits, $LCL = X_{19:125} = 73.992$ and $UCL = X_{107:125} = 74.012$, yield an ARL_0 and FAR of 464.38 and 0.0040, respectively. Note, however, that if one chooses $a = 20$ so that $b = 106$, the control limits are $LCL = X_{20:125}$ and $UCL = X_{106:125}$ and the ARL_0 decreases to 344.73, whereas the FAR slightly increases to 0.0052. The 2-of-2 DR chart is shown in Figure 4.

Finally, for the 2-of-2 KL chart take $a = 21$ so that $b = 125 - 21 + 1 = 105$ so that $LCL = X_{21:125} = 73.992$ and $UCL = X_{105:125} = 74.011$, and this yields an ARL_0 of 460.54 and a FAR of 0.0038, respectively. This 2-of-2 KL chart is almost identical to the chart in Figure 4 and is thus omitted. Both the 2-of-2 DR and KL charts signal on the 10th sample in the

prospective phase. Note, however, that the achieved FAR values for all three charts are much larger than the nominal FAR of 0.0027.

Mann–Whitney–Wilcoxon Control Charts. While the precedence chart is a step in the right direction, the precedence test is not the most popular or the most powerful of nonparametric two-sample tests. That honor goes to the Mann–Whitney (MW) test (see for example, Gibbons and Chakraborti, 2003 [1]). The MW test (equivalent to the popular Wilcoxon rank-sum test) is a well-known nonparametric competitor of the two-independent-sample t -test. The test is known to be more powerful than the precedence test for light tailed distributions and hence MW charts are expected to be more efficient. Of course,

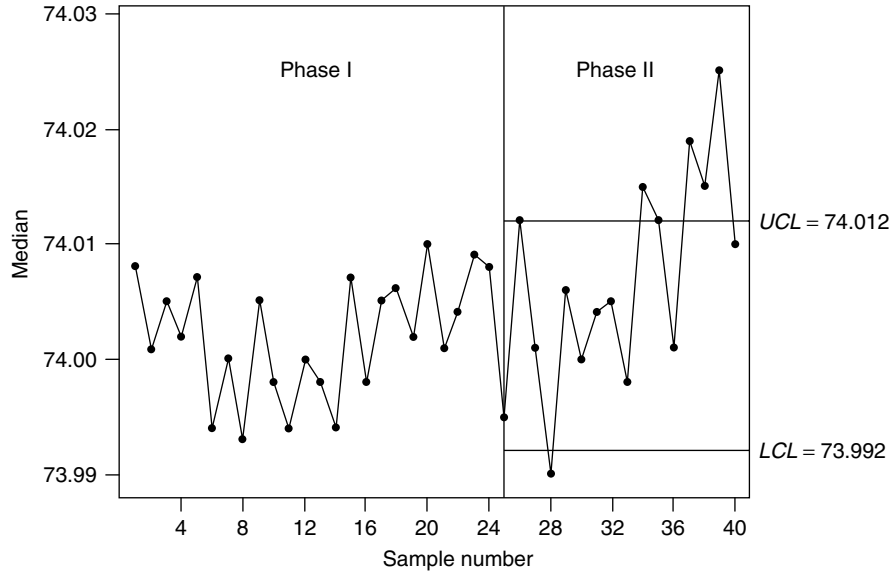


Figure 4 2-of-2 phase II DR chart for the Montgomery [3] piston-ring data

the MW chart is also distribution-free and therefore has the same in-control robustness advantage as the precedence chart, that its in-control distribution is completely known. Park and Reynolds [14] considered control charts for monitoring the location parameter of a continuous process in case U. One of the special cases of their charts is the MW chart based on the well-known Mann–Whitney–Wilcoxon (MWW) statistic. The control limits of these charts are established using phase I reference data. However, they only considered properties of this chart when the reference sample size approaches infinity. Chakraborti and van de Wiel [15] considered the Shewhart-type MW chart, studied its properties, and provided tables for its implementation. These authors show that in some cases the MW chart is more efficient than the precedence chart.

Example 4 For the piston-ring data with $m = 125$ and $n = 5$, Chakraborti and van de Wiel [15] found the upper and the lower control limits of the Shewhart-type MW chart to be 540 and 85, respectively, for a specified $ARL_0 = 400$. The 15

phase II samples and the reference sample lead to 15 MW statistics shown in Table 6 (read from left to right and to left) and the MW control chart is shown in Figure 5.

It is seen that all but three of the test groups, 12, 13 and 14 are in-control. The conclusion from the MW chart is that the medians of test groups 12, 13 and 14 have shifted to the right in comparison with the median of the in-control distribution, assuming that G is a location shift of F . It may be noted that the Shewhart $\bar{X}$ chart shown in Montgomery [3] led to the same conclusion with respect to the means. Of course, the advantage with the MW chart (and with any nonparametric chart) is that it is distribution-free, so that regardless of the underlying distribution, the in-control ARL of the chart is roughly equal to 400 and there is no need to worry about non-normality, as one must for the $\bar{X}$ chart. The distribution-free 1-of-1 precedence chart for this data for an $ARL_0 = 400$ is found to be $LCL = 73.982$ and $UCL = 74.017$ for an attained $ARL_0 \approx 414.0$. Consequently, the precedence chart declares the 12th and the 14th

Table 6 Phase II MW statistics for the piston-ring data in Montgomery [3]

429.0	333.0	142.5	370.5	241.5	410.5	393.0	240.5
471.0	486.0	340.5	561.0	575.5	601.5	484.5	–

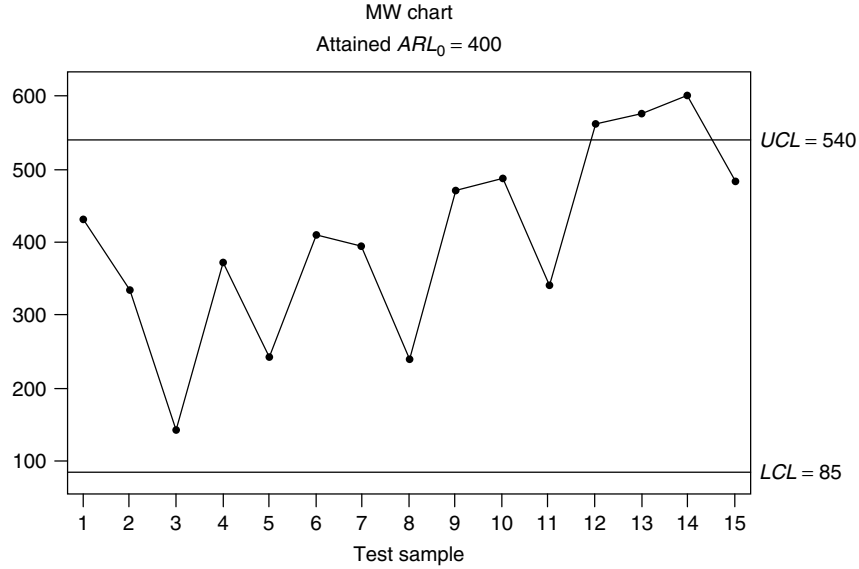


Figure 5 1-of-1 phase II MW chart for the Montgomery [3] piston-ring data

groups to be out-of-control but not the 13th group (see Figure 3), unlike the MW and the Shewhart chart. This is not entirely surprising since the MW test is generally more powerful than the precedence test.

CUSUM-Type Charts

While the Shewhart-type charts are the most widely known and used control charts in practice because of their simplicity and global performance, other classes of charts, such as the CUSUM charts (*see Cumulative Sum (CUSUM) Chart*) are useful and sometimes more naturally appropriate in the process control environment in view of the sequential nature of data collection. These charts, typically based on the cumulative totals of a plotting statistic, obtained as data accumulate, are known to be more efficient for detecting certain types of shifts in the process. The normal theory CUSUM chart for the mean is typically based on the cumulative sum of the deviations of the individual observations (or the subgroup means) from the specified target mean. It seems natural to consider analogs of these charts using the nonparametric plotting statistics discussed earlier. These lead to nonparametric cumulative sum (NPCUSUM) charts.

Case K.

Charts for Location. Bakir and Reynolds [16] proposed a nonparametric CUSUM chart based on the SR statistic introduced earlier to monitor the median of a symmetric distribution with a known in-control value θ_0 . Recall that for the i th random sample, the plotting statistic in the Shewhart-type chart was $SR_i = \sum_{j=1}^n \text{sign}(X_{ij} - \theta_0) R_{ij}$, $i = 1, 2, \dots$. The Bakir–Reynolds grouped signed-rank (GSR) chart instead uses the cumulative sum of the statistic SR_i with a stopping rule. The GSR-CUSUM procedure is distribution free since the in-control distribution of SR_i does not depend on the underlying distribution for all continuous symmetric distributions.

The upper GSR chart uses the plotting statistic $SR_i^+ = \max(0, SR_{i-1}^+ + SR_i - k)$, with a starting value $SR_0^+ = 0$, and a signal is given for the first i when $SR_i^+ \geq h$. The lower GSR chart signals for the first i at which $SR_i^- \leq -h$, where $SR_i^- = \min(0, SR_{i-1}^- + SR_i + k)$ with a starting value $SR_0^- = 0$. The corresponding two-sided symmetric GSR chart signals for the first i at which either of the two inequalities is satisfied that is either $SR_i^+ \geq h$ or $SR_i^- \leq -h$. The statistics SR_i^+ and SR_i^- are reset to zero upon becoming negative and positive, respectively.

Bakir and Reynolds [16] derived the ARL of the chart using a **Markov chain** approach, where

the transition probabilities are calculated using the distribution of the SR statistic. In the in-control case and for small values of n , the distribution of the SR statistic has been tabulated and is available, for example, in Gibbons and Chakraborti [1]. The Markov chain approach is convenient for deriving the run length distribution of CUSUM charts and the associated properties. This approach appears particularly suitable in the nonparametric case since the plotting statistics (such as SR_i^+) are discrete.

One needs the reference value k and the decision value h in order to implement the CUSUM chart. Details regarding how to choose these constants are similar to that in the normal theory case. In fact, Bakir and Reynolds [16] recommended that the optimal k values they calculated for the normal distribution be used in practice unless very heavy tails of the process distribution are indicated. For example, for $n = 6$, one can take $k = 2$ in order to detect a shift of 0.2. Using this value of k , the value of h is chosen to achieve a desired in-control ARL value. They provided tables for the one-sided ARL_0 values, for various combinations of h, k , and n . For example, when n is 6, using $h = 10$ and $k = 11$, the upper GSR-CUSUM chart yields an in-control ARL of 301.01. Comparisons made on the basis of the one-sided ARL, with the Shewhart chart and the normal theory CUSUM chart for various positive shifts under the normal, uniform, double exponential and Cauchy distributions suggest that when observations are naturally collected in groups, the GSR chart is only slightly less efficient than the usual CUSUM chart based on the sample mean when the process is normally distributed, whereas for non-normal distributions the GSR chart can be considerably more efficient. A suitable subgroup size for this nonparametric procedure was suggested to be between $n = 5$ and 10, depending on the shift size and the desired in-control ARL. This recommendation is nearly the same as the subgroup size recommended for the normal theory-based procedures.

Following Bakir and Reynolds [16], Amin *et al.* [4] considered nonparametric CUSUM charts for the median of any continuous population by replacing the SR statistics with the sign statistics SN_i (recall that $SN_i = \sum_{j=1}^n \text{sign}(X_{ij} - \theta_0)$) in the definition of the CUSUM statistic. They also calculated the ARL of the chart using a Markov chain approach, where the transition probabilities are calculated *via* the distribution of the sign statistic, which is of course binomial. Optimal k and h are determined as in the case of the CUSUM of SR statistics by Bakir and Reynolds [16]; using k values for the normal distribution does not lead to large errors. Also, as in the case of the signed-rank CUSUM charts, it is seen that the sign CUSUM chart is generally more efficient than the Shewhart-type sign chart, with or without warning limits. However, more research is clearly necessary, including the full development of two-sided charts and guidelines for practical implementation of the nonparametric sign and the signed-rank CUSUM charts.

We illustrate the sign and signed-rank CUSUM charts using the same data used in Example 1. Recall that $n = 5$; Table 7 shows the upper and lower sign and signed-rank CUSUM statistics. For illustration, we take $k = 2$ for the sign CUSUM chart but for the signed-rank CUSUM chart we take $k = 3$, following Bakir and Reynolds [16], since n is odd. The decision value h for both charts was taken to be 8 for illustration. Figures 6 and 7 show a two-sided sign CUSUM chart and an upper signed-rank CUSUM chart, respectively.

The two-sided sign CUSUM chart signals at sample 13, indicating a most likely upward shift in the process median.

However, the upper signed-rank CUSUM chart signals earlier, at the seventh sample, again indicating a most likely upward shift in the process median. More work and guidance is necessary for the implementation of the sign and signed-rank CUSUM charts.

Table 7 One-sided sign (SN_i^+ and SN_i^-) and signed-rank (SR_i^+ and SR_i^-) statistics for nonparametric CUSUM charts

Sample number	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15
SN_i^+	0	0	0	1	0	1	2	0	1	3	2	5	8	11	13
SN_i^-	0	0	-2	0	0	0	0	0	0	0	0	0	0	0	0
SR_i^+	5	6	0	4	0	6	13	4	13	24	25	37	49	61	72
SR_i^-	0	0	-11	-1	-1	0	0	-3	0	0	0	0	0	0	0

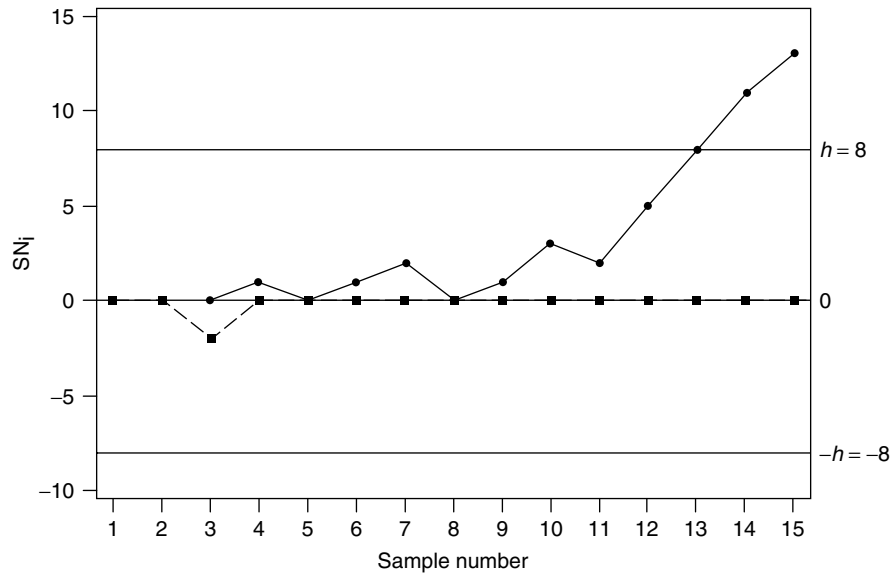


Figure 6 Two-sided tabular sign CUSUM chart for the Montgomery [3]

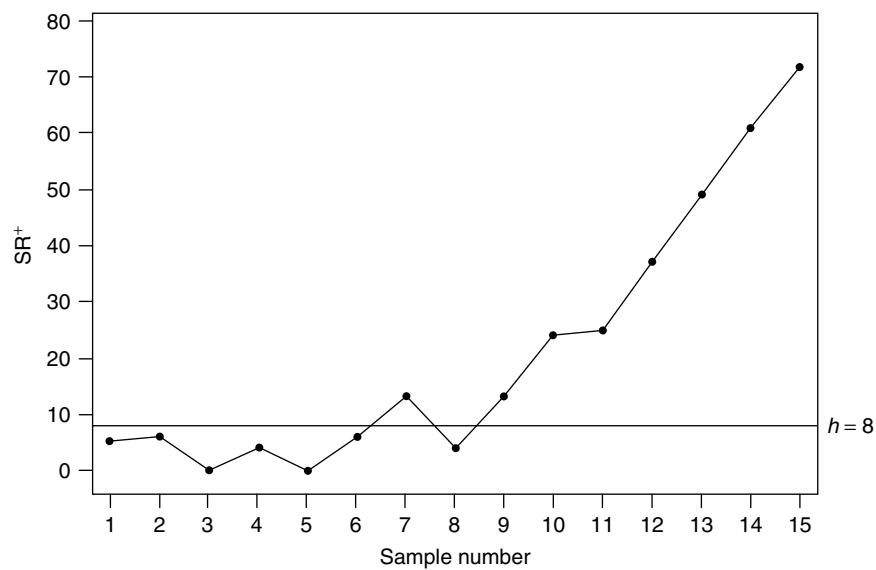


Figure 7 Upper tabular signed-rank CUSUM chart for the Montgomery [3] piston-ring data

It is worth noting that while the CUSUM charts have been shown to be more efficient for small, sustained (persistent) shifts, this efficiency can be achieved only with a proper design of the chart, which means finding the optimum h and k values. This usually requires a specified shift and an ARL_0 .

This “tuning” of CUSUM charts is perhaps a bit more involved for a typical SPC practitioner.

Case U. Bakir [17] considered what are called *SR-like statistics* and used these to consider distribution-free charts. As earlier, in case U, the median of the

in-control distribution (assumed to be symmetric) is unknown and can be estimated by the median of a reference sample, say M . Then for the i th phase II sample, construct the SR-like plotting statistic $\varphi_i^* = \sum_{j=1}^n \text{sign}(Y_{ij} - M) R_{ij}^*$ where R_{ij}^* is the rank of $|Y_{ij} - M|$ among $|Y_{i1} - M|, \dots, |Y_{in} - M|$. This statistic is a direct analog of the plotting statistic SN_i used in case K. It is shown that this statistic is distribution-free and hence charts based on it are also distribution-free. Bakir [17] considered various charts including the Shewhart and the CUSUM based on φ^* . However, more work is necessary on the properties and implementation of these charts.

EWMA-Type Charts

Another popular class of control charts is the EWMA charts. The EWMA charts also take advantage of the sequentially (time ordered) accumulating nature of the data arising in a typical SPC environment and are known to be efficient in detecting smaller shifts but are easier to implement than the CUSUM charts. One can consider nonparametric EWMA-type charts using the plotting statistics or their suitable analogs, considered in cases K and U.

Case K.

Charts for Location. In the spirit of the grouped signed-rank cumulative sum (GSR-CUSUM) charts of Bakir and Reynolds [16], Amin and Searcy [18] considered a nonparametric EWMA chart (grouped signed-rank exponentially weighted **moving average** (GSR-EWMA)) of the SR statistics based on $Z_i = \lambda SR_i + (1 - \lambda)Z_{i-1}$, where λ ($0 < \lambda < 1$) is a smoothing constant. The starting value Z_0 is taken to be the process target value and the process is declared out-of-control at the first subgroup i when Z_i either falls above or below the control limits. Performance of the GSR-EWMA chart is found to be similar to that of the GSR-CUSUM charts. An advantage of the GSR-EWMA is its relative insensitivity to the choice of λ . In general, smaller (larger) λ values result in a faster detection of small (large) shifts. Enhancements such as addition of warning limits improve performance of the chart. **Autocorrelation** does not seem to affect the ARL as much it affects the ARL of an $\bar{X}$ -EWMA chart. Again, more work is necessary on the practical implementation of the charts as well as on adaptations in case U.

Other Charts

There has been other work on nonparametric control charts. Among these, for example, Albers and Kallenberg [19] studied conditions under which the nonparametric charts become viable alternatives to their parametric counterparts. They consider phase II charts for individual observations in case U based on empirical quartiles or order statistics. The basic problem is that for the very small FAR typically used in the industry, a very large reference sample size is usually necessary to set up the chart. They discuss various remedies for this problem.

Another area that has received some attention is control charts for variable sampling intervals (VSI). In a typical control charting environment, the time interval between two successive samples is fixed, and this called a *fixed sampling interval* (FSI) scheme. VSI schemes allow the user to vary the sampling interval between taking samples. This idea has intuitive appeal since when one or more charting statistics fall close to one of the control limits but not quite outside, it seems reasonable to sample more frequently, whereas when charting statistics plot closer to the centerline, no action is necessary and only a few samples might be sufficient. Amin and Widmaier [20] compared the Shewhart $\bar{X}$ charts with sign control charts, under the FSI and VSI schemes, on the basis of ARL for various shift sizes and several underlying distributions like the normal distribution and distributions that are heavy-tailed and/or asymmetric like the double exponential and the gamma. It is seen that the nonparametric VSI sign charts are more efficient than the corresponding FSI sign charts.

Nonparametric charts for individual observations are useful in practice; some work in this area has been done, see for example CVB.

Summary and Conclusions

Nonparametric control charts can be a useful tool for the quality practitioner in a variety of applications where not much is known or can be assumed about the process distribution (say regarding its shape or the variance). Although much has been accomplished in the last few years and a lot of work is currently underway, much more remains to be done. From a practical standpoint, the charting procedures must be made more accessible to the practitioner and to this end, the ease of implementation is vital.

Computer programs, add-ons to popular software packages such as Minitab and SAS, and/or websites would greatly help in this effort. In terms of research, more work needs to be done on charts for scale and on charts for simultaneous monitoring of location and scale parameters. Rank and rank-like statistics can be considered in a CUSUM and/or EWMA framework to design more powerful nonparametric charts. Moreover, the run length distributions of the various charts (both existing and new) need to be fully developed and examined, say in terms of various percentiles. Also, research is necessary on nonparametric control charts for phase I analysis as well as for discrete and categorical data.

References

- [1] Gibbons, J.D. & Chakraborti, S. (2003). *Nonparametric Statistical Inference*, 4th Edition, Revised and Expanded, Marcel Dekker, New York.
- [2] Chakraborti, S., van der Laan, P. & Bakir, S.T. (2001). Nonparametric control charts: an overview and some results, *Journal of Quality Technology* **33**, 304–315.
- [3] Montgomery, D.C. (2001). *Introduction to Statistical Quality Control*, 4th Edition, John Wiley & Sons, New York.
- [4] Amin, R.W., Reynolds Jr, M.R. & Bakir, S.T. (1995). Nonparametric quality control charts based on the sign statistic, *Communications in Statistics: Theory and Methods* **24**, 1597–1623.
- [5] Bakir, S.T. (2004). A distribution-free Shewhart quality control chart based on signed-ranks, *Quality Engineering* **16**, 613–623.
- [6] Nelson, L.S. (1984). The Shewhart control chart—tests for special causes, *Journal of Quality Technology* **16**, 237–239.
- [7] Western Electronic Company. (1956). *Statistical Quality Control Handbook*, Western Electric Corporation, Indianapolis.
- [8] Page, E.S. (1962). A modified control chart with warning lines, *Biometrika* **49**, 171–176.
- [9] Roberts, S.W. (1958). Properties of control chart zone tests, *The Bell System Technical Journal* **37**, 83–114.
- [10] Chakraborti, S. & Eryilmaz, S. (2007). A nonparametric Shewhart-type signed-rank control chart based on runs, *Communications in Statistics: Simulation and Computation* **36**, 335–356.
- [11] Janacek, G.J. & Meikle, S.E. (1997). Control charts based on medians, *The Statistician* **46**, 19–31.
- [12] Chakraborti, S., van der Laan, P. & van de Wiel, M.A. (2004). A class of distribution-free control charts, *Journal of the Royal Statistical Society. Series C: Applied Statistics* **53**, 443–462.
- [13] Chakraborti, S., Eryilmaz, S. & Human, S. (2006). A Phase II Shewhart-type nonparametric control chart based on runs (Submitted).
- [14] Park, C. & Reynolds Jr, M.R. (1987). Nonparametric procedures for monitoring a location parameter based on linear placement statistics, *Sequential Analysis* **6**, 303–323.
- [15] Chakraborti, S. & van de Wiel, M.A. (2003). A nonparametric control chart based on the Mann-Whitney statistic, SPOR Report 2003-24, Technical University of Eindhoven, Eindhoven.
- [16] Bakir, S.T. & Reynolds Jr, M.R. (1979). A nonparametric procedure for process control based on within-group ranking, *Technometrics* **21**, 175–183.
- [17] Bakir, S.T. (2006). Distribution-free quality control charts based on signed-rank-like statistics, *Communications in Statistics: Theory and Methods* **35**, 743–757.
- [18] Amin, R.W. & Searcy, A.J. (1991). A nonparametric exponentially weighted moving average control scheme, *Communications in Statistics: Simulation and Computation* **20**, 1049–1072.
- [19] Albers, W. & Kallenberg, W.C.M. (2004). Empirical non-parametric control charts: estimation effects and corrections, *Journal of Applied Statistics* **31**, 345–360.
- [20] Amin, R.W. & Widmaier, O. (1999). Sign control charts with variable sampling intervals, *Communications in Statistics: Theory and Methods* **28**, 1961–1985.

SUBHABRATA CHAKRABORTI AND MARIEN A.
GRAHAM

Variable Sampling Rate Control Charts

Introduction

Control charts are used to monitor a process to detect changes in the process and its output. The first control charts were developed by Walter A. Shewhart in the 1920s and were used in industrial applications. Control charts are now in widespread use throughout the world and have played a major role in improving quality and productivity in a wide range of applications. The application areas are no longer restricted to industrial and manufacturing settings, but include service industries, government, healthcare, and even law enforcement and athletics.

The traditional practice when using a control chart to monitor a process is to take samples of fixed size from the process using a fixed sampling interval between sampling points. For example, if the objective is to detect changes in the parameters of the distribution of a critical dimension of a component that is being manufactured, then this dimension might be measured for samples of $n = 4$ components taken using a sampling interval of $d = 4$ h. A control chart is operated by plotting a statistic, such as the sample mean or sample variance, after each sample is taken. If any plotted point falls outside of control limits that have been set up for the chart, then this is taken as strong evidence (a signal) that there has been a change in the process. The process is said to be in control if there has been no change in the process and out of control if there has been a change.

With the traditional Shewhart charts the statistic plotted at each sampling point is a function only of the sample at that sampling point. Other types of control charts, such as cumulative sum (CUSUM) charts (*see Cumulative Sum (CUSUM) Chart*) and exponentially weighted moving average (EWMA) charts (*see Exponentially Weighted Moving Average (EWMA) Control Chart*) are based on plotting statistics that are functions of current and past samples. CUSUM and EWMA charts are more effective than the traditional Shewhart charts for detecting small and moderate changes in the process parameters.

Standard control charts can be called *fixed sampling rate* (FSR) control charts because the sampling

rate is constant, but *variable sampling rate* (VSR) control charts vary the sampling rate as a function of the data obtained from the process being monitored. The basic idea behind VSR control charts is that the sampling rate should be increased when there is some indication of a problem with the process, and decreased when there is no indication of a problem with the process. If the indication of a problem with process is strong enough, then a VSR control chart signals in the same way as a standard FSR control chart. A review of VSR control charts, also called *adaptive sampling* control charts, was given by Tagaras [1].

Variable Sampling Interval (VSI) Control Charts

One class of VSR control charts consists of the *variable sampling interval* (VSI) control charts that vary the sampling interval as a function of the data from the process. With a VSI control chart a short sampling interval is used next when there is some indication of a problem with the process, and a longer sampling interval is used next when there is no indication of a problem with the process.

To develop an illustration of a VSI control chart, consider the standard FSR Shewhart $\bar{X}$ chart for monitoring the process mean, where the process observations are normal with mean μ and standard deviation σ . Let μ_0 and σ_0 be the in-control parameter values, which for simplicity, are assumed to be known or estimated with negligible error. Suppose that the standard FSR $\bar{X}$ chart is based on taking samples of size $n = 4$ from the process using a sampling interval of $d = 4$ h, and using 3σ control limits, $\mu_0 \pm 3\sigma_0/\sqrt{n}$. When the process is in control the probability of a signal (a false alarm) is 0.0027, and the *average number of samples to signal* (ANSS) is then $1/0.0027 = 370.4$. Samples are being taken every 4 h, so the in-control *average time to signal* (ATS) is then $(370.4)(4) = 1481.6$ h.

Consider now the application of a VSI $\bar{X}$ chart to this problem. Suppose that it is feasible to sample again as soon as 30 min after the previous sample and it is desirable to maintain the same average sampling rate (the rate of one observation per hour). Consider the VSI $\bar{X}$ chart that is based on two possible sampling intervals, $d_1 = 0.5$ h (30 min) and $d_2 = 6.0$ h, with an additional set of limits at $\mu_0 \pm$

2 Variable Sampling Rate Control Charts

$0.9052\sigma_0/\sqrt{n}$. The sampling rule is to use the long sampling interval d_2 next if the current value of $\bar{X}$ falls within the limits $\mu_0 \pm 0.9052\sigma_0/\sqrt{n}$, and use the short sampling interval d_1 next if the current value of $\bar{X}$ falls outside $\mu_0 \pm 0.9052\sigma_0/\sqrt{n}$ but within the control limits $\mu_0 \pm 3\sigma_0/\sqrt{n}$. A signal is given in the standard way if $\bar{X}$ falls outside $\mu_0 \pm 3\sigma_0/\sqrt{n}$. The value 0.9052 used in defining the inner limits was chosen so that, when the process is in control, the probability of being within these limits is 0.6364, conditional on no signal. This means that, when the process is in control, the probability of being outside $\mu_0 \pm 0.9052\sigma_0/\sqrt{n}$ but within $\mu_0 \pm 3\sigma_0/\sqrt{n}$ is 0.3636, conditional on no signal. Thus, conditional on no signal, the expected sampling interval is $(0.5)(0.3636) + (6.0)(0.6364) = 4.0$ h. Thus, this VSI $\bar{X}$ chart has been designed so that the in-control average sampling interval is the same as the standard FSR $\bar{X}$ chart.

Now consider the situation in which there is some change in the process and this change results in a one standard deviation shift in μ from μ_0 to $\mu_0 + \sigma_0$. The probability of $\bar{X}$ falling outside of $\mu_0 \pm 3\sigma_0/\sqrt{n}$ and resulting in a signal is now 0.1587. The ANSS of both the FSR and VSI $\bar{X}$ charts is thus $1/0.1587 = 6.3030$. Now if $\mu = \mu_0 + \sigma_0$, the probability that the VSI chart uses the short sampling interval has increased to 0.8396, and the probability of using the long sampling interval has decreased to 0.1604. Thus, the average sampling interval of the VSI chart, conditional on no signal, will be $(0.5)(0.8396) + (6.0)(0.1604) = 1.3823$ h. Then, the out-of-control ATS for detecting this shift will be $(6.3030)(1.3823) = 8.71$ h. Note that the standard FSR $\bar{X}$ chart always uses a sampling interval of 4.0 h, so the out-of-control ATS of this chart is $(6.3030)(4.0) = 25.21$ h. Therefore, the VSI $\bar{X}$ chart will detect this shift much faster, on average, than the standard FSR $\bar{X}$ chart, because the shift in the mean has caused the VSI chart to sample at a higher rate.

The first VSI charts to be developed were Shewhart-type charts (see, for example, Reynolds *et al.* [2], Cui and Reynolds [3], and Runger and Pignatiello [4]). Since the development of the VSI Shewhart charts, the VSI idea has been extended to other types of charts. Reynolds *et al.* [5] developed a VSI CUSUM chart, and Saccucci *et al.* [6] developed a VSI EWMA chart.

Most work on VSI charts has concentrated on the problem of monitoring μ , but in practice it is usually

important to monitor both μ and σ . Shamma and Amin [7] investigated the performance of a single VSI EWMA chart that can be used to simultaneously monitor μ and σ . Reynolds and Stoumbos [8] and Stoumbos and Reynolds [9] investigated combinations of VSI Shewhart and EWMA charts for the problem of simultaneously monitoring μ and σ . Reynolds and Cho [39] developed multivariate (see **Multivariate Control Charts Overview**) VSI charts for monitoring the process mean and variability.

In the illustration of the VSI $\bar{X}$ chart discussed above, only two possible sampling intervals, $d_1 = 0.5$ h and $d_2 = 6.0$ h, were allowed. It might seem that it would be better to allow more than two possible sampling intervals, but computational and theoretical results have shown that allowing only two possible intervals is optimal or near optimal. See Reynolds and Arnold [11], Reynolds [12, 13], and Stoumbos *et al.* [14].

VSI Control Charts with Sampling at Fixed Times

The VSI $\bar{X}$ chart discussed as an example above used two possible sampling intervals, 0.5 and 6.0 h, and was compared to a standard FSR $\bar{X}$ chart that sampled every 4.0 h. In some applications, it may be desirable to sample at certain specified fixed times. For example, if process personnel work 8-h shifts, then the standard FSR $\bar{X}$ chart is based on two samples every shift. However, the sampling intervals of the VSI chart are either 0.5 or 6.0 h, depending on the data from the process, so the sampling points for the VSI chart will be unpredictable.

If it is desirable for sampling to be more predictable, the VSI chart sampling decision rules could be modified to specify that samples always be taken at specified fixed times. For example, it could be specified that samples always be taken every 4 h, with additional samples taken within each 4-h interval whenever necessary. This would result in the VSI chart sampling at a higher average rate than the FSR chart. Alternately, it could be specified that the VSI chart samples once during each 8-h work shift, with additional samples taken within the 8-h period only when necessary. In this case, the VSI chart would be sampling less often, on average, than the FSR chart.

See Reynolds [15, 16], Baxley [17], and Stoumbos and Reynolds [18] for discussion of this modification of VSI charts.

Variable Sample Size (VSS) Control Charts

Another class of VSR control charts consists of the *variable sample size* (VSS) control charts that vary the sample size as a function of the data from the process. With a VSS control chart a large sample size is used next when there is some indication of a problem with the process, and a small sample size is used next when there is no indication of a problem with the process.

To develop an illustration of a VSS control chart, consider the situation discussed previously in which the standard FSR Shewhart $\bar{X}$ chart is based on taking samples of size $n = 4$, with a sampling interval of $d = 4$ h, and control limits $\mu_0 \pm 3\sigma_0/\sqrt{n}$.

Consider the applications of a VSS $\bar{X}$ chart to this problem. Suppose that it is feasible to take a sample of eight observations whenever it appears to be necessary and it is desirable to maintain an average sampling rate of one observation per hour. Consider the VSS $\bar{X}$ chart based on two possible sample sizes, $n_1 = 2$ and $n_2 = 8$, and an additional set of limits at $\mu_0 \pm 0.9638\sigma_0/\sqrt{n}$. The sampling rule is to use the small sample size n_1 next if the current value of $\bar{X}$ falls within the limits $\mu_0 \pm 0.9638\sigma_0/\sqrt{n}$, and use the large sample size n_2 next if the current value of $\bar{X}$ falls outside $\mu_0 \pm 0.9638\sigma_0/\sqrt{n}$ but within the control limits $\mu_0 \pm 3\sigma_0/\sqrt{n}$. A signal is given in the standard way if $\bar{X}$ falls outside $\mu_0 \pm 3\sigma_0/\sqrt{n}$. The value 0.9638 used in defining the inner limits was chosen so that, when the process is in control, the probability of being within these limits is 0.6667, conditional on no signal. This means that, when the process is in control, the probability of being outside $\mu_0 \pm 0.9638\sigma_0/\sqrt{n}$ but within $\mu_0 \pm 3\sigma_0/\sqrt{n}$ is 0.3333, conditional on no signal. Thus, conditional on no signal, the expected sample size is $(2)(0.3333) + (8)(0.6667) = 4$. Thus, this VSS $\bar{X}$ chart has been designed so that the in-control average sample size is the same as the sample size for the standard FSR $\bar{X}$ chart. Note that both sets of limits involve $\sqrt{n}$, and the value of n to be used in these limits depends on the sample size being used in the current sample.

When there is a shift in the process mean, the probability of using the large sample size n_2 will increase, while the probability of using the small sample size n_1 will decrease. If the large sample size is used more frequently, then the expected number of samples required for detection will decrease, leading to faster detection of this shift.

The first VSS charts to be developed were Shewhart-type charts (see, for example, Prabhu *et al.* [19] and Costa [20]). Since the development of the VSS Shewhart charts, the VSS idea has been extended to other types of charts. Annadi *et al.* [21] developed a VSS CUSUM control chart.

In the illustration of the VSS $\bar{X}$ discussed above, only two possible sample intervals, $n_1 = 2$ and $n_2 = 8$, were allowed. There are no theoretical results that give the optimal number of possible sample sizes, but using only two possible sample sizes seems to give good performance and provides relatively simple sampling rules for the VSS chart.

VSI–VSS Control Charts

The VSI and VSS approaches to varying the sampling rate can be used simultaneously. The basic idea would be to use a short sampling interval and large sample size next if there is some indication of a problem with the process, and a long sampling interval and small sample size next if there is no indication of a problem. Thus, in addition to the control limits, two other sets of limits could be added to the chart. One set of limits would be used to determine the next sampling interval, and the other to determine the next sample size. The control limits would be used as usual to determine when to signal.

For VSI–VSS Shewhart charts see, for example, Prabhu *et al.* [22], Costa [23], and Park and Reynolds [24]. Arnold and Reynolds [25] developed VSI–VSS CUSUM charts, and Reynolds and Arnold [26] developed VSI–VSS EWMA charts. Aparisi and Haro [27] investigated multivariate VSI–VSS Shewhart-type charts.

Control Charts Based on Sequential Sampling

Another approach to implementing the VSR concept is related to the VSS approach, but the sample size is changed in a different way. With standard VSS

control charts, as usually defined, the sample size for the current sample is determined by past process data, before sampling starts for the current sample. Another approach, called *sequential sampling* allows the sample size at a sampling point to be determined by the data at that sampling point (along with the data at previous sampling points). With sequential sampling, at each sampling point observations are taken in groups of one or more. After each group is taken at a sampling point, a decision is made to either take another group at this sampling point, wait until the next sampling point to sample again, or signal at the current sampling point. Sequential sampling is appropriate for situations in which sampling must be done at equally spaced sampling points, and it is feasible to take samples consisting of multiple groups of observations at a sampling point whenever necessary.

The first work on sequential sampling control charts was for Shewhart charts with a maximum of two groups allowed at each sampling point (see Croasdale [28] and Daudin [29]). Better performance can be achieved if multiple groups are allowed at a sampling point. An extension is the case in which observations are taken individually, with no limit on the number of observations at a sampling point. When the objective is to detect a parameter shift in one direction, a very efficient sequential sampling scheme can be obtained by applying a sequential probability ratio test (SPRT) at each sampling point (see Stoumbos and Reynolds [18, 30, 31] and Reynolds and Stoumbos [32]).

Reynolds and Kim [33] developed multivariate sequential sampling schemes for monitoring the process mean, and Reynolds and Kim [34] developed multivariate sequential sampling schemes for simultaneously monitoring the process mean and variability.

In general, sequential sampling schemes can be comparable in performance to the best VSI or VSS schemes, and all of these schemes will provide very substantial improvements in performance compared to traditional FSR control charts.

Conclusions and Discussion

Using the VSR feature in process monitoring can give very substantial improvements in the ability to detect process changes. A number of approaches to varying the sampling rate as a function of the

process data have been discussed here. The approach that gives the best overall performance depends on which type of control chart is being used and on the constraints placed on sampling. Sampling constraints include the smallest feasible sampling interval, the largest feasible sample size, and requirements that sampling be done at specified fixed sampling points in time.

Earlier in this article the application of the VSI and VSS features was illustrated in the context of the simple Shewhart $\bar{X}$ chart for monitoring μ . In most applications, it will also be important to simultaneously monitor both μ and σ . Combinations of two Shewhart charts, such as the $\bar{X}$ chart and the R chart, have traditionally been used to monitor μ and σ . Although Shewhart charts have the advantage of simplicity, recent research (see, for example, Reynolds and Stoumbos [35–37] and Reynolds and Cho [10]) has shown that combinations of EWMA or CUSUM charts have better overall statistical performance than Shewhart chart combinations for the problem of simultaneously monitoring μ and σ . Thus, we recommend that combinations of EWMA or CUSUM charts be used in applications instead of Shewhart chart combinations.

Recent research (see, for example, Hawkins and Olwell [38], Reynolds and Stoumbos [35–37] and Reynolds and Cho [10]) has also shown that with EWMA or CUSUM charts the best overall statistical performance is achieved by taking small frequent samples instead of taking larger samples less frequently. For example, if there is a choice between taking samples of $n = 4$ every $d = 4$ h or samples of $n = 1$ every $d = 1$ h, then it is better overall to take samples of $n = 1$ every hour. However, samples of size $n = 4$ might be reasonable when there is a significant cost-saving from sampling in groups of $n = 4$ instead of sampling individual observations. Taking samples of $n = 4$ would also be reasonable if Shewhart charts are being used, because the best value of n for Shewhart charts is not the same as the best value of n for EWMA or CUSUM charts.

The fact that $n = 1$ gives the best overall performance for EWMA or CUSUM charts suggests that the best approach to applying the VSR concept is through the VSI feature with $n = 1$, because using the VSS feature assumes that samples of $n > 1$ will sometimes be taken. However, there are a number of situations in which it may be appropriate to use the VSS or sequential sampling features. For example,

suppose that items are produced in batches, a batch takes 4 h to process, and the current practice is to take samples of $n = 4$ from each batch. In this situation, the VSI approach does not seem to be applicable, so a good option would be to use the VSS or sequential sampling features, where the sample size within a batch can vary depending on the process data.

The use of the VSR feature in applications requires that procedures for designing a VSR chart be available. Design procedures can include software or tables that give appropriate VSR chart parameters, such as the limits that determine when to use each sampling interval or sample size. Design procedures are available in many of the references previously cited here, but these design procedures tend to be for relatively simple cases involving the monitoring of only one parameter, such as μ . However, Reynolds and Stoumbos [8] discuss the design of combinations of VSI EWMA charts for the simultaneous monitoring of μ and σ when $n = 1$. More work is needed to develop easy-to-use design procedures for other combinations of VSR control charts.

References

- [1] Tagaras, G. (1998). A survey of recent developments in the design of adaptive control charts, *Journal of Quality Technology* **30**, 212–231.
- [2] Reynolds Jr, M.R., Amin, R.W., Arnold, J.C. & Nachlas, J.A. (1988). $\bar{X}$ charts with variable sampling intervals, *Technometrics* **30**, 181–192.
- [3] Cui, R. & Reynolds Jr, M.R. (1988). $\bar{X}$ charts with runs rules and variable sampling intervals, *Communications in Statistics: Theory and Methods* **17**, 1073–1093.
- [4] Runger, G.C. & Pignatiello, J.J. (1991). Adaptive sampling for process control, *Journal of Quality Technology* **23**, 135–155.
- [5] Reynolds Jr, M.R., Amin, R.W. & Arnold, J.C. (1990). CUSUM charts with variable sampling intervals, *Technometrics* **32**, 371–384.
- [6] Saccucci, M.S., Amin, R.W. & Lucas, J.M. (1992). Exponentially weighted moving control schemes with variable sampling intervals, *Communications in Statistics: Simulation and Computation* **21**, 627–657.
- [7] Shamma, S.E. & Amin, R.W. (1993). An EWMA quality control procedure for jointly monitoring the mean and variance, *International Journal of Quality and Reliability Management* **10**, 58–67.
- [8] Reynolds Jr, M.R. & Stoumbos, Z.G. (2001). Monitoring the process mean and variance using individual observations and variable sampling intervals, *Journal of Quality Technology* **33**, 181–205.
- [9] Stoumbos, Z.G. & Reynolds Jr, M.R. (2005). Economic statistical design of adaptive control schemes for monitoring the mean and variance: an application to analyzers, *Nonlinear Analysis: Real World Applications* **6**, 817–844.
- [10] Reynolds Jr, M.R. & Cho, G.Y. (2006a). Multivariate control charts for monitoring the mean vector and covariance matrix, *Journal of Quality Technology* **38**, 230–253.
- [11] Reynolds Jr, M.R. & Arnold, J.C. (1989). Optimal Shewhart control charts with variable sampling intervals between samples, *Sequential Analysis* **8**, 51–77.
- [12] Reynolds Jr, M.R. (1989). Optimal variable sampling interval control charts, *Sequential Analysis* **8**, 361–379.
- [13] Reynolds Jr, M.R. (1995). Evaluating properties of variable sampling interval control charts, *Sequential Analysis* **14**, 59–97.
- [14] Stoumbos, Z.G., Mittenthal, J. & Runger, G.C. (2001). Steady-state-optimal adaptive control charts based on variable sampling intervals, *Stochastic Analysis and Applications* **19**, 1025–1057.
- [15] Reynolds Jr, M.R. (1996a). Shewhart and EWMA variable sampling interval control charts with sampling at fixed times, *Journal of Quality Technology* **28**, 199–212.
- [16] Reynolds Jr, M.R. (1996b). Variable sampling interval control charts with sampling at fixed times, *IIE Transactions* **28**, 497–510.
- [17] Baxley Jr, R.V. (1995). An application of variable sampling interval control charts, *Journal of Quality Technology* **27**, 275–282.
- [18] Stoumbos, Z.G. & Reynolds Jr, M.R. (2001). The SPRT control chart for the process mean with samples starting at fixed times, *Nonlinear Analysis: Real World Applications* **2**, 1–34.
- [19] Prabhu, S.S., Runger, G.C. & Keats, J.B. (1993). An adaptive sample size $\bar{X}$ control scheme, *International Journal of Production Research* **31**, 2895–2909.
- [20] Costa, A.F.B. (1994). $\bar{X}$ charts with variable sample size, *Journal of Quality Technology* **26**, 155–163.
- [21] Annadi, H.P., Keats, J.B., Runger, G.C. & Montgomery, D.C. (1995). An adaptive sample size CUSUM control chart, *International Journal of Production Research* **33**, 1605–1616.
- [22] Prabhu, S.S., Montgomery, D.C. & Runger, G.C. (1994). A combined adaptive sample size and sampling interval $\bar{X}$ control scheme, *Journal of Quality Technology* **26**, 164–176.
- [23] Costa, A.F.B. (1997). $\bar{X}$ charts with variable sample size and sampling intervals, *Journal of Quality Technology* **29**, 197–204.
- [24] Park, C. & Reynolds Jr, M.R. (1999). Economic design of a variable sampling rate $\bar{X}$ chart, *Journal of Quality Technology* **31**, 427–443.
- [25] Arnold, J.C. & Reynolds Jr, M.R. (2001). CUSUM control charts with variable sample sizes and sampling intervals, *Journal of Quality Technology* **33**, 68–81.

6 Variable Sampling Rate Control Charts

- [26] Reynolds Jr, M.R. & Arnold, J.C. (2001). EWMA control charts with variable sample sizes and variable sampling intervals, *IIE Transactions* **33**, 511–530.
- [27] Aparisi, F. & Haro, C.L. (2003). A comparison of T^2 control charts with variable sampling schemes as opposed to MEWMA chart, *International Journal of Production Research* **41**, 2169–2182.
- [28] Croasdale, P. (1974). Control charts for a double-sampling scheme based on average production run lengths, *International Journal of Production Research* **12**, 585–592.
- [29] Daudin, J.J. (1992). Double sampling $\bar{X}$ charts, *Journal of Quality Technology* **24**, 78–87.
- [30] Stoumbos, Z.G. & Reynolds Jr, M.R. (1996). Control charts applying a general sequential test at each sampling point, *Sequential Analysis* **15**, 159–183.
- [31] Stoumbos, Z.G. & Reynolds Jr, M.R. (1997). Control charts applying a sequential test at fixed sampling intervals, *Journal of Quality Technology* **29**, 21–40.
- [32] Reynolds Jr, M.R. & Stoumbos, Z.G. (1998). The SPRT chart for monitoring a proportion, *IIE Transactions* **30**, 545–561.
- [33] Reynolds Jr, M.R. & Kim, K. (2005). Multivariate monitoring of the process mean vector with sequential sampling, *Journal of Quality Technology* **37**, 149–162.
- [34] Reynolds Jr, M.R. & Kim, K. (2007). Multivariate control charts for monitoring the process mean and variability using sequential sampling, *Sequential Analysis* **26**, 283–315.
- [35] Reynolds Jr, M.R. & Stoumbos, Z.G. (2004a). Control charts and the efficient allocation of sampling resources, *Technometrics* **46**, 200–214.
- [36] Reynolds Jr, M.R. & Stoumbos, Z.G. (2004b). Should observations be grouped for effective process monitoring? *Journal of Quality Technology* **36**, 343–366.
- [37] Reynolds Jr, M.R. & Stoumbos, Z.G. (2005). Should exponentially weighted moving average and cumulative sum charts be used with Shewhart limits? *Technometrics* **47**, 409–424.
- [38] Hawkins, D.M. & Olwell, D.H. (1998). *Cumulative Sum Control Charts and Charting for Quality Improvement*, Springer-Verlag, New York.
- [39] Reynolds Jr, M.R. & Cho, G.Y. (2006b). Multivariate Control Charts for Monitoring the Mean Vector and Covariance Matrix with Variable Sampling Intervals, Technical Report 07-5, Department of Statistics, Virginia Tech.

MARION R. REYNOLDS, JR AND ZACHARY G.
STOUMBOS

Neural Networks: Construction and Evaluation

Introduction

Neural network (NN) stems from the thought of trying to mimic the ability of the human brain in both structure and function. The human brain contains many billions of nerve cells, or neurons. Each neuron is a cell, which connects to other neurons through synapses. Each neuron also uses biochemical reactions to receive, process, and transmit information. Following this idea, NNs have been designed to achieve human brain-like performance and have been applied in many fields of study including statistics.

A NN model (*see Neural Networks in Statistical Process Control*) is usually demonstrated as layers of functional units, called nodes or neurons or perceptrons. There are at least two layers, input and output. However, a NN normally has an additional layer referred to as a “*hidden layer*”. Nodes of the NN are connected by weights. The nodes and weights of the NN are comparable to neurons and synapses, respectively. Hence, the NN model solves a problem by finding optimal weights. Typically, NNs are utilized in three problem areas: (a) classification and diagnosis, (b) function approximation and prediction, and (c) optimization.

The type of NN is defined the NN architecture or topology. Lippmann [1] provides an introduction to six important NNs that can be used for pattern classification: (a) Hopfield net, (b) Hamming net, (c) Carpenter–Grosberg classifier, (d) perceptron, (e) multi-layer perceptron (MLP), and (f) Kohonen self-organizing feature maps. Existing NN architectures currently used are: MLP, radial basis function (RBF), learning vector quantization (LVQ), adaptive resonance theory (ART), auto-associative NNs, and Kohonen self-organizing maps (SOM).

To construct a NN model, there are five important factors for consideration.

1. NN architecture

Different NN architectures have different rules for solving modeling problems. Learning or training is the term used to explain the process of finding

the weights. There are two types of NN training. Supervised learning is known as learning with a teacher. It happens when there is a known target value associated with the input nodes. One of the most popular algorithms used in supervised learning is back-propagation. Unsupervised learning is known as learning without a teacher. It occurs when data lack target output values relating to the pattern of input nodes. The algorithms of unsupervised learning can be found in Lippmann [1].

2. Number of hidden layers

Since the nodes in the input layer only pass data to other layers, computation occurs in both hidden layers and the output layer. Adding additional hidden layers to solve a problem may provide a better result than use of only two-layer, input and output, networks. However, there is no rule of thumb for the suitable number of hidden layers. The more the hidden layers, the larger the NN size, which leads to overfitting (overtraining) problems. On the other hand, too few hidden layers cause underfitting (undertraining) problems.

3. Number of nodes in each layer

Generally, input layer nodes are directly related to the number of observations needed to represent each set of input variables. No guidelines exist for the number of hidden layer nodes as with the number of hidden layers. Overfitting problems occur when there are too many hidden nodes, and underfitting problems when there are too few hidden nodes. Output layer nodes correspond to the anticipated number of distinct patterns or classifications that are to be decided by the network.

4. Type of connection

In a full-connection network every single node in a layer connects to every node in the next layer. Otherwise, it is a partial-connection network. In addition, there is a feed-forward network in which nodes in each layer can only be connected to other nodes in the next layer, while in the recurrent network nodes in each layer can be connected to the nodes in the previous layer, next layer and even nodes in the same layer.

5. Activation function or transfer function

This function occurs at every computational node, both hidden and output. It can be either linear or non-linear, symmetrical or nonsymmetrical. Two common functions widely used are the sigmoid function (logistic function) and hyperbolic tangent function (*see Neural Networks in Statistical Process Control*).

An Example: NNs for Multivariate Quality Control

Traditional **multivariate control charts**, such as the chi-square (*see Hotelling's T^2 Chart*) and the multivariate exponentially weighted moving average (MEWMA) charts (*see Multivariate Exponentially Weighted Moving Average (MEWMA) Control Chart*), are commonly used for monitoring a process and detecting process shifts when there are several correlated quality characteristics. These two traditional techniques make three important assumptions:

1. the data follow a multivariate **normal distribution**,
2. the **covariance** matrix is constant over time, and
3. the observations are independent over time.

The performance of both charts is seriously affected if any of these assumptions is not met. The NN then becomes an alternative for consideration because it does not require such strict assumptions. The NN approach finds the relationships between input variables and output variables from existing data. Therefore, an appropriately designed NN may be easier for practitioners to use and more robust to traditional assumptions, while providing equal or improved performance relative to traditional multivariate statistical process control (SPC) methods.

The bivariate case is illustrated using simulated data written by the interactive matrix language procedure (Proc IML) of Statistical Analysis System (SAS) for generating multivariate observations and assessing the performance of the three different methods based on the average run length (ARLs). For each method, ARLs were calculated for both the in-control and out-of-control state, over the total number of observations, to compare the performance of all three methods. As described briefly in the previous section, there are many NN architectures that may be applied to classification problems, which are focused on a change in structure for mean shifts, assuming constant variance. The MLP is one feasible NN architecture. Advantages of the MLP are its simplicity and flexibility.

Construction of a Neural Network for SPC

1. NN architecture

The feed-forward MLP with back-propagation is a feasible and suitable NN because of its simplicity.

2. Type of connection

A full connection is used following the recommendation of Pugh [2] and Chang and Aw [3] as this approach speeds network convergence.

3. Number of hidden layers

Only one hidden layer is used following the result of Pugh [2, 4] as no significant performance difference is reported for the MLP with one-hidden layer or two-hidden layers.

4. Number of nodes

(a) Input layer nodes

There are two issues that affect the appropriate number of nodes for a NN.

First, if the problem does not address pattern recognition, the input layer contains only the current observation, regardless of preceding and succeeding data. Hence, each input node might be the vector of current quality characteristic mean of each observation sized n . There are then two input nodes. Each of them represents the sample mean of each i th observation for the first and second quality variable, respectively, as the vector of $(\bar{X}_1, \bar{X}_2)$, where $\bar{X}_j = \frac{\sum_{k=1}^n x_{ijk}}{n}$. Summarizing a multivariate sample by using the sample mean vector NN inputs may be considered as pre-processing of the data. Note that this procedure is based on the assumption of a constant variance matrix of two correlated variables. One should also note that the training phase for NN modeling with SPC requires data from both in-control and out-of-control process states. This is quite different from the traditional methods which require phase I data to be from an in-control process.

The other situation is that when pattern recognition is addressed, preceding data are supplied to the NN through the input nodes, in addition to current data. Three consecutive observations then are considered at time periods t , $t-1$, and $t-2$, respectively. In each time period, there is the vector of sample means for the first and second quality variable. Hence, there are six input nodes expressed by $(\bar{X}_{(t,1)}, \bar{X}_{(t,2)}, \bar{X}_{(t-1,1)}, \bar{X}_{(t-1,2)}, \bar{X}_{(t-2,1)}, \bar{X}_{(t-2,2)})$, where $\bar{X}_{(i,j)}$ is the sample mean for the quality characteristic j , $j = 1, 2$, at time period i for $i = t, t-1, t-2$.

(b) Hidden layer nodes

As Guo and Dooley [5] reported, there is no standard rule for deciding the optimum number of hidden layer nodes. The primary concern is

the overtraining or undertraining problem. Thus the performance of NN is illustrated with three and five hidden nodes.

(c) Output layer nodes

Only one output node is required since we are only interested in a change in structure from an in-control state to an out-of-control state. One might consider multiple outputs to signal various changes such as changes in the variance–covariance structures or mean shifts in specific directions.

Four NN architectures are illustrated. They are graphically displayed in Figures 1–4.

5. Cutoff point estimation

At the final step of NN process monitoring, the NN output must be checked against its appropriate cutoff point (control limit) to determine whether the process is in control. Based on the idea of a false alarm rate, α , the cutoff point is determined from the **quantiles** of the predicted values for in-control data. For example, if the false alarm rate, α , equals 0.005 and there are 10 000 and 15 000 observations of in-control data, so the cutoff point is at the 0.005 quantile of the network output distribution. Yi *et al.* [6] used this method to compare **ARL** performances between MLP NNs and $\bar{X}$ control charts for investigating nonnormally distributed quality characteristics.

Since the distribution of the cutoff point is also not known, using a smoothed estimate rather than the sample quantile may be better. The smoothed estimate of the 99.995th percentile, known as the

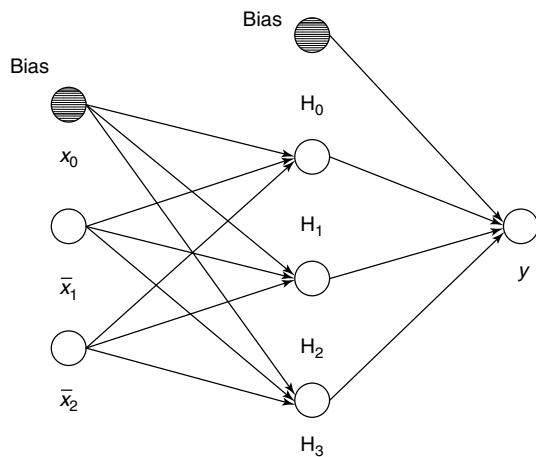


Figure 1 MLP with two input nodes, and three hidden nodes in a hidden layer (NN2(3))

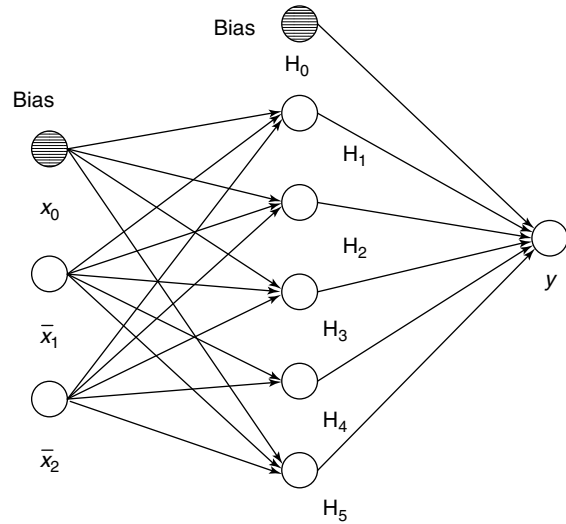


Figure 2 MLP with two input nodes, and five hidden nodes in a hidden layer (NN2(5))

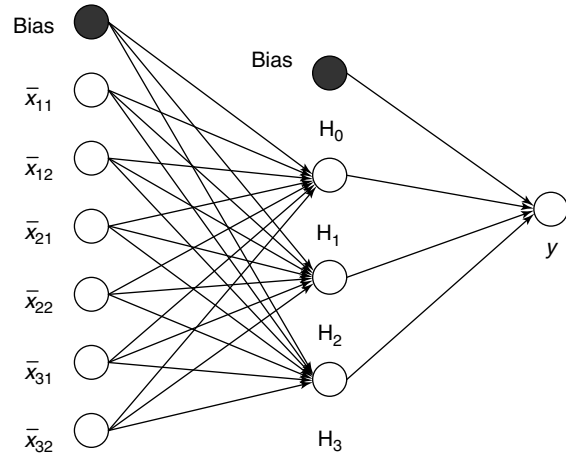


Figure 3 MLP with six input nodes, and three hidden nodes in a hidden layer (NN6(3))

Harrell-Davis [7] estimator, might be provided a starting point for determining more accurate cutoff points through simulation.

Performance of the Illustrated NN Systems

All competing monitoring schemes are designed to provide equal in-control ARLs of 200. The out-of-control ARL of each method is used to compare

4 Neural Networks: Construction and Evaluation

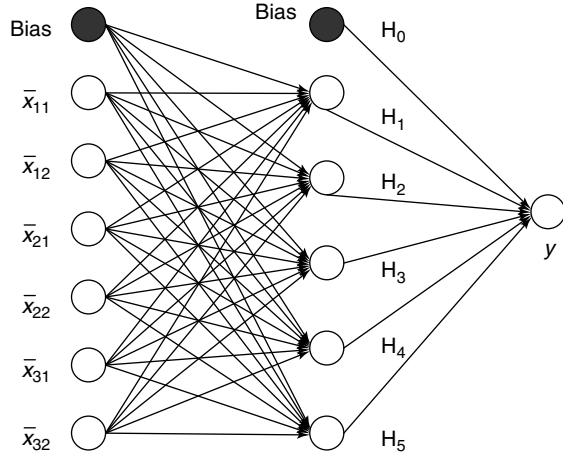


Figure 4 MLP with six input nodes, and five hidden nodes in a hidden layer (NN6(5))

performances under two different shift sizes, small and large shifts, as measured by the noncentrality parameter term, δ , as

$$\delta^2(\mu) = (\mu - \mu_0)' \sum^{-1} (\mu - \mu_0)$$

where μ_0 is the in-control mean vector and μ is the current mean vector. A large value of δ indicates that there are big shifts in the mean. If the value of $\delta = 0$ ($\mu = \mu_0$), the process is in control. In addition to the size of the shift, the performances of the NN methods are evaluated for one-variable shifts along the x-axis and two-variable shifts along the major and minor axes of the bivariate distribution. A major axis shift

is defined in the eigenvector direction associated with the largest eigenvalue of the inverse of covariance (correlation) matrix and a minor axis shift is in the eigenvector direction associated with the smallest eigenvalue of the inverse of covariance (correlation) matrix. NNs are trained using two different sizes of data sets, small and large sizes, consisting of 20 000 and 30 000 observations. Monitoring schemes exhibiting best ARL performances across the various illustrated conditions are provided in Figures 5 and 6.

As for training NN with smaller training set, the NN6(3) provides the best performance at the large shift. For small shifts, the MEWMA method provides the best performance if the variables are not highly correlated. However, if the variables are highly correlated, the NN6(3) provides the best performance. Similar results are obtained for training NN with larger training sets.

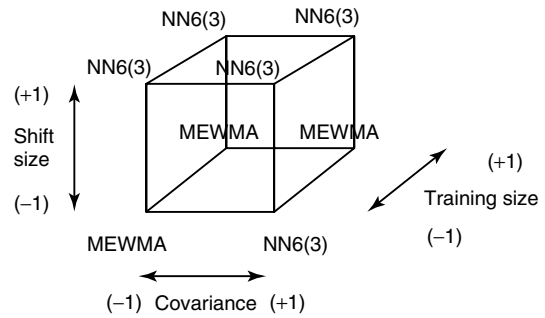


Figure 5 A graphical display of optimal methods for a one-variable shift

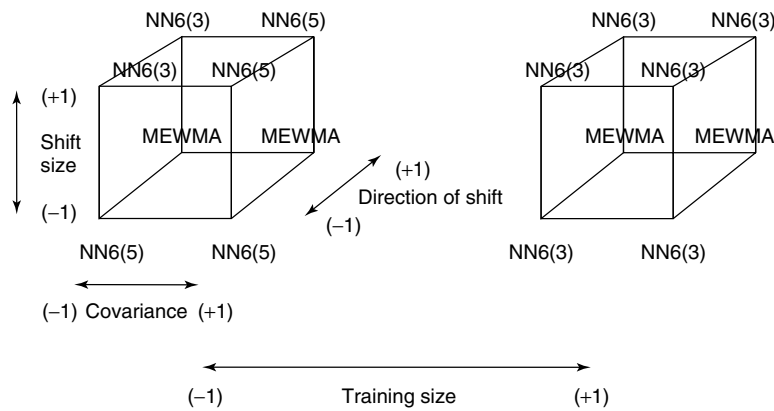


Figure 6 A graphical display of optimal methods for two-variable shifts

For two-variable shifts, four factors, covariance (A), shift size (B), training size (E), and direction of shift (F), each at two levels, are of interest. The graphical display of the optimal method at each condition is shown in Figure 6.

The interested reader can find additional information regarding the use of NN models in multivariate quality in Saithanu [8].

Strengths

The practitioner will find many advantages to using NN-based monitoring schemes. These advantages include good performance and flexibility. The flexibility of this method refers not only to relaxation of distributional assumptions associated with traditional methods, but also includes flexibility in applications. Univariate, multivariate, and regression-based applications (see **Regression Control Charts**), as well as special purpose systems, are legitimate candidates for NN modeling.

Weaknesses

NNs come with some disadvantages or weaknesses. Some of these disadvantages are associated with the lack of experience most practitioners have with NN models. The selection of an appropriate architecture as well as training issues may require steep learning curves for the inexperienced. There is also the matter of model interpretation that is problematic for NN models. Finally, NN models and their development are computational intensive. These computational requirements come in several forms. Two of the most immediate issues are the need for specialized software and the need for large training datasets.

Conclusions

The NN is a very flexible tool that provides competitive performances relative to more traditional process monitoring schemes. The development of these monitoring systems has requirements that are somewhat unique to this method. The need for both in-control and out-of-control data for the training phase, large training datasets and specialized software have been mentioned in this article. The software required to design and develop NN models continues to become more mainstream and available to practitioners. In complex high volume production environments, where large data sets are common, the practitioner may be well served by NN-based monitoring systems.

References

- [1] Lippmann, P.R. (1987). An introduction to computing with neural nets, *IEEE ASSP Magazine* **4**, 4–22.
- [2] Pugh, G.A. (1991). A comparison of neural networks to SPC charts, *Computers and Industrial Engineering* **21**, 253–255.
- [3] Chang, S.I. & Aw, C.A. (1996). A neural fuzzy control chart for detecting and classifying process mean shifts, *International Journal of Production Research* **34**, 2265–2278.
- [4] Pugh, G.A. (1989). Synthetic neural networks for process control, *Computers and Industrial Engineering* **17**, 24–26.
- [5] Guo, Y. & Dooley, K.J. (1992). Identification of change structure in statistical process control, *International Journal of Production Research* **30**, 1655–1669.
- [6] Yi, J., Prybutok, R.V. & Clayton, R.H. (2001). ARL comparisons between neural network models and $\bar{X}$ control charts for quality characteristics that are nonnormally distributed, *Economic Quality Control* **16**, 5–15.
- [7] Harrell, F.E. & Davis, C.E. (1982). A new distribution-free quantile estimator, *Biometrika* **69**, 635–640.
- [8] Saithanu, K. (2006). *Neural networks and multivariate quality control*, Ph.D. Dissertation, University of Alabama, Tuscaloosa.

KIDAKAN SAITHANU

Control Charts for Short Production Runs

Introduction

Shewhart **control charts** are used to monitor both products and processes to determine if and when action should be taken to adjust a process because of changes in centering and/or spread of the quality characteristic being measured. In Shewhart control charting, m subgroups of size n consisting of measurements of a quality characteristic of a part or process are collected. The mean ($\bar{X}$) in combination with the range (R), variance (v), or standard deviation (s) is calculated for each subgroup. When the subgroup size is 1, individual values (denoted by X) are used in combination with moving ranges (denoted by MR) of size 2. The mean of the subgroup means ($\bar{\bar{X}}$) and subgroup ranges ($\bar{R}$), variances ($\bar{v}$), or standard deviations ($\bar{s}$) are calculated and used to determine estimates of the process mean and standard deviation, respectively. When the subgroup size is 1, the mean of the individual values ($\bar{X}$) and moving ranges ($\bar{MR}$) are used for this purpose. These parameter estimates are then used to construct control limits using conventional control chart constants for monitoring the performance of the process. These conventional control chart constants assume that an infinite number of subgroups are available to estimate the process mean and standard deviation.

Hillier [1] presented three situations in which this assumption is invalid. The first is in the initiation of a new process. The second is during the start-up of a process just brought into statistical control again. The third is for a process whose total output is not large enough to use conventional control chart constants. Each of these is an example of a short-run situation. In a short-run situation, little or no historical information is available about a process in order to estimate the process mean and standard deviation to begin control charting. Consequently, the initial data obtained from the early run of the process must be used for this purpose.

In recent years, manufacturing companies have increasingly faced each of these short-run situations. One reason is the widespread application of the just-in-time (JIT) philosophy, which has caused much

shorter continuous runs of products. Other reasons are frequently changing product lines and product characteristics caused by shorter-lived products, fast-paced product innovation, and changing consumer demand. Fortunately, flexible manufacturing technology has provided companies with the ability to alter their processes in order to face these challenges. However, traditional statistical process control (SPC) methodologies do not provide companies with the ability to reliably monitor quality in each of the previously mentioned short-run situations. A common rule of thumb, which has been widely accepted despite evidence that it may be incorrect, states that 20–30 subgroups of size 4 or 5 are necessary before parameter estimates may be obtained to construct control limits using conventional control chart constants. This is a difficult if not an impossible rule to satisfy in a short-run situation. As a result, SPC methodologies have been developed for monitoring quality in each of the previously mentioned short-run situations. Though several of these methodologies are discussed here, the particular focus of this discussion is the two-stage short-run control charting methodology.

Pooling Data

The prevalent methodologies focus on pooling data from different parts onto a single control chart combination (i.e., onto $(\bar{X}, R)$, $(\bar{X}, v)$, $(\bar{X}, \sqrt{v})$, $(\bar{X}, s)$, and (X, MR) control charts) in order to have enough data to satisfy the rule. The difference between $(\bar{X}, \sqrt{v})$ and $(\bar{X}, s)$ control charts is that the former are constructed using the statistic $\sqrt{\bar{v}}$ and the latter are constructed using the statistic $\bar{s}$. Pooling data is the procedure of taking measurements of quality characteristics from different parts, performing a transformation on the measurements, and plotting the transformed measurements from the different parts on the same control chart. Typically, all of the part numbers on the same control chart are produced by one machine or process. Hence, control charting using pooled data is often termed a *process-focused approach* rather than a *product-focused approach* to control charting. Elam [2] provided references for pooled data control charts. For example, see [3, 4].

Pooling data is advantageous because it reduces the number of control charts in use, which greatly simplifies control chart management programs. Also,

2 Control Charts for Short Production Runs

in most cases, control charting can begin almost immediately after the start-up of a process because control limits are known and constant. However, pooling data has several problems. In a true short-run situation, one will often find it difficult to even proceed to pool data. The reason is that, to construct control limits from pooled data, many part types or operations with similar characteristics must be produced or performed, respectively, by the same process.

Another problem is process parameters for each part number are estimated using target or nominal values, tolerances, specification limits, initial subgroups drawn from the process, or historical data. Using target or nominal values is equivalent to using specification limits instead of statistical control limits on control charts, which is a serious mistake. The same can be said for tolerances and specification limits. The reason is that the process target (what you want), the process aim (what you set), and the process average (what you get) are never the same. The magnitude of the differences depends on how well the process is performing. The result is a control chart that, in general, will be useless in delineating special cause variation from common cause variation.

Using initial subgroups drawn from the process to obtain parameter estimates for part numbers begs the original short-run problem that motivates the use of pooled data. If one has enough data (as defined by the common rule of thumb) from a process for a single part to estimate its process parameters, then pooling data is not necessary in the first place. When one has historical data to estimate process parameters for part numbers, then by definition one is not in a short-run situation. Consequently, pooling data is not even necessary, other than to reduce the number of control charts in use.

Finally, an original motivation for pooling data was to satisfy the common rule of thumb. However, satisfying the rule does not guarantee control limits that result in the desired false alarm rate and have a high probability of detecting a special cause signal.

Greater Sensitivity

A second approach to control charting in a short-run situation is using control charts with greater sensitivity (i.e., more statistical **power**) than Shewhart

control charts. In a short-run situation, the total output of the process is not large. Consequently, the quick detection of special cause signals takes on added importance. It is well known that cumulative sum (**CUSUM**) and **exponentially weighted moving average** (**EWMA**) control schemes are more sensitive to detecting small process shifts than Shewhart control charts. Also, economically designed control charts have greater sensitivity. Consequently, these have been adapted for use in short-run situations. Elam [2] provided references for greater sensitivity control charts. For example, Hawkins [5] presents variations of CUSUM and EWMA control schemes to be used in short-run situations, Doganaksoy and Vandeven [6] combine CUSUM and EWMA control schemes with other control chart methodologies for use in short-run situations, and Del Castillo and Montgomery [7] present economically designed control charts for short-run situations.

An advantage of this approach is that it allows for the quick detection of special cause signals. A disadvantage is that initial estimates of the process parameters must be close to their true values in order for the control charts to perform well. Also, the methodologies that comprise this approach are difficult to implement.

Process Inputs

A third approach to control charting in a short-run situation is to monitor and control process inputs rather than process outputs. The assumption upon which this approach is based is that, by correctly selecting and monitoring critical input variables, one can control the output of the process. Elam [2] provided references for process input control charts.

An advantage of this approach is that, since large amounts of process input data may be available even in a short-run situation, Shewhart control charting may be used. A problem with this approach is that critical input parameters for a new part to be produced in a short run may not match all of the critical input parameters for which large amounts of data are available. Also, this approach assumes that the process input nominal values are the same for all products fabricated on that process. If nominal values are different, a transformation of the process input data may be required.

Control Charts with Modified Limits

A fourth approach to control charting in a short-run situation is using control charts with modified limits. In a true short-run situation, the process mean and standard deviation are unknown and must be estimated from a small number of subgroups with only a few samples each drawn from the start-up of a process. When these estimates are used with conventional control chart constants to construct control limits for $(\bar{X}, R)$, $(\bar{X}, v)$, $(\bar{X}, \sqrt{v})$, $(\bar{X}, s)$, and $(\bar{X}, MR)$ control charts, the false alarm probability becomes distorted. Consequently, modified control chart factors need to be used to achieve the desired false alarm probability.

Elam [2] provided references up to 2001 for control charting techniques that fall under this approach. One of these techniques is the Q chart methodology (see [8]). It is advantageous in that, not only does it allow for the pooling of data from different parts, but different statistics may be plotted on the same Q chart. Also, control charting can begin almost immediately after the start-up of a process because control limits are known and constant. Disadvantages of Q charts are their inability to detect a process that starts out of control and their general lack of sensitivity in detecting process changes. Also, process standard deviation estimates used to calculate Q statistics to be plotted on Q charts are unreliable.

Two-Stage Short-Run Control Charts

Hillier [1] developed a two-stage short-run theory of control charting. Though this was one of the first short-run control charting methodologies, focus shifted to the previously discussed approaches in the 1980s and 1990s. The two-stage procedure is used to determine both the initial state of the process and the control limits for testing the future performance of the process. In the first stage, the initial subgroups drawn from the process are used to determine the control limits. The initial subgroups are plotted against the control limits to retrospectively test if the process was in-control while the initial subgroups were being drawn. Any out-of-control initial subgroups are deleted using a delete and revise (D&R) procedure. Once control is established, the procedure moves to the second stage, where the initial subgroups that were not deleted in the first stage are used to determine the control limits for testing if the process

remains in-control while future subgroups are drawn. Each stage uses a different set of control chart factors called *first-stage short-run control chart factors* and *second-stage short-run control chart factors*.

Hillier [1] applied his two-stage short-run control charting methodology to $(\bar{X}, R)$ control charts. The methodology allows for the specification of the desired false alarm probability. Hillier [1] included the methodology for second-stage short-run control charting for $\bar{X}$ charts and R charts presented by Hillier [9] and Hillier [10], respectively. Hillier [9] showed that the probability of a false alarm is exceedingly high when estimates of the process mean and standard deviation based on a small number of subgroups are used together with conventional control chart constants to construct $\bar{X}$ charts for future testing (stage two). To resolve this issue, Hillier [9] derived an equation for A_2^* , the second-stage short-run control chart factor for the $\bar{X}$ chart. Using this factor, which depends on m (the number of subgroups) as well as the subgroup size n , instead of the conventional control chart constant A_2 results in control limits that give the desired false alarm probability. The value A_2^* is related to A_2 in that, as $m \rightarrow \infty$, $A_2^* \rightarrow A_2$.

Hillier [10] showed that the probability of a false alarm is exceedingly high when estimates of the process standard deviation based on a small number of subgroups are used together with conventional control chart constants to construct R charts for future testing (stage two). To resolve this issue, he [10] derived equations for D_4^* and D_3^* , the second-stage short-run upper and lower control chart factors, respectively, for the R chart. Using these factors, which depend on m as well as n , instead of the corresponding α -based (i.e., probability based) conventional upper and lower control chart constants D_4 and D_3 , respectively, result in control limits that give the desired false alarm probability. The value D_4^* is related to D_4 in that, as $m \rightarrow \infty$, $D_4^* \rightarrow D_4$. Similarly, the value D_3^* is related to D_3 in that, as $m \rightarrow \infty$, $D_3^* \rightarrow D_3$.

Hillier [1] incorporated the two-stage procedure with his results in [9, 10] and derived equations to calculate first- and second-stage short-run control chart factors, for $(\bar{X}, R)$ charts. Using these factors, when process parameter estimates come from a small number of subgroups, results in control chart limits that reliably indicate when a process has gone out of control. The first-stage short-run control chart factor

4 Control Charts for Short Production Runs

for the $\bar{X}$ chart is denoted by A_2^{**} . It depends on m as well as n . The value A_2^{**} is related to A_2 in that, as $m \rightarrow \infty$, $A_2^{**} \rightarrow A_2$. The first-stage short-run upper and lower control chart factors for the R chart are denoted by D_4^{**} and D_3^{**} , respectively. Each of these factors depends on m as well as n . As $m \rightarrow \infty$, $D_4^{**} \rightarrow D_4$ and $D_3^{**} \rightarrow D_3$.

Discussion and computer software for designing two-stage short-run control charts may be found in [11–17] for $(\bar{X}, R)$, $(\bar{X}, v)$ and $(\bar{X}, \sqrt{v})$, $(\bar{X}, s)$, and (X, MR) applications, respectively. The following is an example for two-stage short-run $(\bar{X}, R)$ control charts. Table 1 has the data. For $m = 5$ and $n = 4$, the following first-stage short-run control chart factors are obtained from Table A4 in [11]: $A_2^{**} = 0.77660$, $D_4^{**} = 2.11840$, and $D_3^{**} = 0.11338$. First-stage short-run control limits are calculated as follows: $UCL(R) = D_4^{**} \times \bar{R} = 2.11840 \times 0.216 = 0.45757$, $LCL(R) = D_3^{**} \times \bar{R} = 0.11338 \times 0.216 = 0.02449$, $UCL(\bar{X}) = \bar{\bar{X}} + A_2^{**} \times \bar{R} = 1.286 + 0.77660 \times 0.216 = 1.45375$, and $LCL(\bar{X}) = \bar{\bar{X}} - A_2^{**} \times \bar{R} = 1.286 - 0.77660 \times 0.216 = 1.11825$. R for subgroup five ($R = 0.49$) is above $UCL(R)$. After deleting subgroup five, the revised averages in Table 1 are calculated and the following first-stage short-run control chart factors are obtained from Table A4 in [11] (for $m = 4$ and $n = 4$): $A_2^{**} = 0.78832$, $D_4^{**} = 2.07041$, and $D_3^{**} = 0.11848$. Revised first-stage short-run control limits are calculated as follows: $UCL(R) = 2.07041 \times 0.1475 = 0.30539$, $LCL(R) = 0.11848 \times 0.1475 = 0.01748$, $UCL(\bar{X}) = 1.278125 + 0.78832 \times 0.1475 = 1.39440$, and $LCL(\bar{X}) = 1.278125 - 0.78832 \times 0.1475 = 1.16185$. Since none of the remaining values plot out of control (i.e., control has been established), the next step is to construct second-stage control limits using the following second-stage short-run control chart factors from Table A4 in

[11] (for $m = 4$ and $n = 4$): $A_2^* = 1.01772$, $D_4^* = 2.94060$, and $D_3^* = 0.09281$. Second-stage short-run control limits are calculated as follows: $UCL(R) = D_4^* \times \bar{R} = 2.94060 \times 0.1475 = 0.43374$, $LCL(R) = D_3^* \times \bar{R} = 0.09281 \times 0.1475 = 0.01369$, $UCL(\bar{X}) = \bar{\bar{X}} + A_2^* \times \bar{R} = 1.278125 + 1.01772 \times 0.1475 = 1.42824$, and $LCL(\bar{X}) = \bar{\bar{X}} - A_2^* \times \bar{R} = 1.278125 - 1.01772 \times 0.1475 = 1.12801$. These control limits may be used to monitor the future performance of the process.

The two-stage short-run control charting methodology has advantages over the other short-run control charting methodologies. It overcomes their reliance on the common rule of thumb, their use of parameter estimates that are not representative of the process, their assumption that the process starts in-control, and their complex implementation. A disadvantage of the two-stage short-run control charting methodology is its recently investigated issue (see [18, 19]) regarding the dependence between the events that the i th and j th plotted values ($i \neq j$) plot outside the same estimated control limit.

Conclusion

This research with two-stage short-run control charts has resulted in a more comprehensive and industry-accessible methodology for control charting in short-run situations. Consequently, opportunities exist to more fully integrate this methodology with other control charting procedures, namely, multivariate and fuzzy control charting. These particular integrations are beneficial because situations where data is short-run can occur in conjunction with multivariate and fuzzy data.

References

- [1] Hillier, F.S. (1969). $\bar{X}$ - and R -chart control limits based on a small number of subgroups, *Journal of Quality Technology* **1**(1), 17–26.
- [2] Elam, M.E. (2001). *Investigation, extension, and generalization of a methodology for two stage short run variables control charting*, Ph.D. Dissertation, Oklahoma State University, Stillwater.
- [3] Wheeler, D.J. (1991). *Short Run SPC*, SPC Press, Knoxville.
- [4] Lin, S.Y., Lai, Y.J. & Chang, S.I. (1997). Short-run statistical process control: multicriteria part family formation, *Quality and Reliability Engineering International* **13**, 9–24.

Table 1 A numerical example

Subgroup	X_1	X_2	X_3	X_4	$\bar{X}$	R
1	1.17	1.14	1.20	1.18	1.1725	0.06
2	1.38	1.29	1.36	1.44	1.3675	0.15
3	1.20	1.21	1.30	1.14	1.2125	0.16
4	1.40	1.40	1.21	1.43	1.3600	0.22
5	1.12	1.20	1.61	1.34	1.3175	0.49
Averages					1.286	0.216
Revised averages					1.278125	0.1475

-
- [5] Hawkins, D.M. (1987). Self-starting CUSUM charts for location and scale, *The Statistician* **36**(4), 299–316.
 - [6] Doganaksoy, N. & Vandeven, M.. (1997). Process monitoring with multiple products and production lines, *Quality Engineering* **9**(4), 689–702.
 - [7] Del Castillo, E. & Montgomery, D.C. (1996). A general model for the optimal economic design of $\bar{X}$ charts used to control short or long run processes, *IIE Transactions* **28**(3), 193–201.
 - [8] Quesenberry, C.P. (1991). SPC Q charts for start-up processes and short or long runs, *Journal of Quality Technology* **23**(3), 213–224.
 - [9] Hillier, F.S. (1964). $\bar{X}$ chart control limits based on a small number of subgroups, *Industrial Quality Control* **20**(8), 24–29.
 - [10] Hillier, F.S. (1967). Small sample probability limits for the range chart, *Journal of the American Statistical Association* **62**(320), 1488–1493.
 - [11] Elam, M.E. & Case, K.E. (2001). A computer program to calculate two-stage short-run control chart factors for $(\bar{X}, R)$ charts, *Quality Engineering* **14**(1), 77–102.
 - [12] Elam, M.E. & Case, K.E. (2003). Two-stage short-run $(\bar{X}, v)$ and $(\bar{X}, \sqrt{v})$ control charts, *Quality Engineering* **15**(3), 441–448.
 - [13] Elam, M.E. & Case, K.E. (2003). A computer program to calculate two-stage short-run control chart factors for $(\bar{X}, v)$ and $(\bar{X}, \sqrt{v})$ charts, *Quality Engineering* **15**(4), 609–638.
 - [14] Elam, M.E. & Case, K.E. (2005). Two-stage short-run $(\bar{X}, s)$ control charts, *Quality Engineering* **17**(1), 95–107.
 - [15] Elam, M.E. & Case, K.E. (2005). A computer program to calculate two-stage short-run control chart factors for $(\bar{X}, s)$ charts, *Quality Engineering* **17**(2), 259–277.
 - [16] Elam, M.E. & Case, K.E. (2006). Two-stage short-run (X, MR) control charts, *Journal of Modern Applied Statistical Methods* **5**(2), to appear.
 - [17] Elam, M.E. & Case, K.E. (2006). A computer program to calculate two-stage short-run control chart factors for (X, MR) charts, *Journal of Statistical Software* **15**(11), <http://www.jstatsoft.org/>.
 - [18] Elam, M.E. & Tibbs, L.B. (2004). Correlation results for two-stage short-run control charts, in *Proceedings of the Institute of Industrial Engineers Annual Research Conference*, Houston.
 - [19] Elam, M.E. (2005). Correlation results for two-stage short-run (X, MR) control charts, in *Proceedings of the Institute of Industrial Engineers Annual Research Conference*, Atlanta.

MATTHEW E. ELAM

Control Charts for Batch Processes

Introduction

Batch processing in manufacturing can be described as the simultaneous processing of materials or semifinished units that have been grouped together as a single processing entity. Numerous industries rely on batch processing of one kind or another, including some as diverse as specialty chemicals, pharmaceuticals, food products, fabrics, metals, bakeries, semiconductors, and pulp and paper manufacturers. Distinguishing features of batch processing include the following:

1. A setup of the process occurs during which a recipe of materials or group of units is loaded or “charged” into the processing equipment (e.g., vessel, oven, tank, furnace, reactor).
2. Process input variables (X ’s) such as temperature and pressure are ramped up or down to specified conditions following particular time-oriented “trajectories”.
3. The batch is processed over time frames that may be relatively lengthy, from hours to days, or even longer. This constraint is often the motivation for the utilization of batch processing over single-unit processing.
4. Upon completion, the product material is discharged, and the processing equipment may require cleaning and/or maintenance before the next batch is commenced.

Typically, measurements of product quality characteristics (Y ’s) on samples of the product taken from the batch are then obtained. Usually, these measurements are conducted off-line using gauges or test equipment that may be located nearby or in a remote location (e.g., outside laboratory), and there may be numerous characteristics (Y ’s) that require control within tight specification limits. Because such measurements and tests usually cannot be made until the product is accessible, an additional feature of batch processing is the following.

5. Product quality characteristics are often obtained so late relative to the possible occurrence of

process disturbances or “special cause” variation (see **Assignable Cause; Control Charts, Overview**) as to render them useless in identifying quality problems with the affected batch. Moreover, these causes may lead to unacceptable quality of subsequent batches if the cause system is not understood and remedies applied in a timely manner. Because an unacceptable batch may comprise a large amount of product, the financial impact of a poorly controlled batch process on a manufacturing operation can be substantial.

In *discrete manufacturing* batch processing involves the grouping of processing units (parts) from prior steps in the manufacturing process into a batch or “load”. The load of parts is then processed simultaneously. Examples include the deposition of thin films on substrates (wafers) loaded onto “boats” in a diffusion furnace in semiconductor manufacturing, or the heat treating of metal castings in an oven in an investment casting operation. In the steps prior to the batch-processing step, the wafers (or castings) were the processing units. In *continuous manufacturing* the batch process will operate somewhat differently. A recipe of materials (e.g., solid and liquid feedstocks, chemical additives, recycled stock, etc.) is supplied to the batch vessel, where the batch is processed under controlled conditions as described earlier. Generally, there are no prior subunits that have been assembled into one processing entity, as in discrete manufacturing.

Issues in Applications of Control Charts

Since a batch process by definition repeats itself batch to batch, the batch is the primary processing unit, whether formed from groups of discrete units or from mixed streams of feedstocks, chemicals, and so on, in a continuous batch process. However, a secondary processing unit that must be taken into account is the within-batch unit (e.g., wafer or casting or, perhaps, location in a vessel or tank). This requirement is particularly true if measurement of the Y ’s requires sampling of the within-batch secondary units. Therefore, there are two sources of common cause variation and must be accounted for in the Y ’s, batch-to-batch variation from sampling the primary units and within-batch variation from sampling the secondary units. Control chart users must take this

2 Control Charts for Batch Processes

reality into account when they design control charts and select the rational subgroup for a batch process (see **Control Charts, Overview**; **Control Charts, Selection of**).

In discrete manufacturing, where Shewhart-type charts are commonly applied, a common mistake made by SPC practitioners is to assume that n samples of the within-batch secondary units represent a rational subgroup of size n , followed by the application of $\bar{X}$ -bar and R or S charts to the resulting data. The “subgroup of size n ” assumption is based upon two erroneous notions: (a) that the within-batch variation accounts for all of the common cause variation in the process, which is not valid because the batch-to-batch variation is absent and (b) that n independent primary units (repetitions) have been selected from the process. As expected, the resulting charts have unduly narrow control limits that signal that the process is seriously out of control with respect to mean Y . Consider in Figure 1 an example of an $\bar{X}$ -bar chart of three measurements of a critical dimension (CD) at three sites on 74 wafers from a semiconductor manufacturing operation [1].

The chart exhibits chronic out-of-control behavior; however, the subgroup of size 3 is improper, given that the CD measurements are taken from three sites on the same processing unit (wafer), not from three independently produced processing units.

Two options are available to the control chart user

to address this problem. First, if the within-batch sampling is constant from batch to batch, location and dispersion statistics (e.g., mean and standard deviation) can be calculated from the set of n measurements and considered as overall single measurements of batch “performance”. The location metric represents overall batch quality while the dispersion metric represents within-batch variability or “uniformity”. Individuals charts can then be applied to these batch-level individual measurements ([2, 3], as well as [1] for a case study of this approach). Consider an individuals chart of the CD data in Figure 2.

Note that the control limits are much wider than those of the chart depicted in Figure 1. The fact that the batch-to-batch variability is included in, rather than absent from, the estimate of the common cause variability provided by the average moving range explains the difference. (Moreover, the batch-to-batch variability exceeds the within-batch variability, which only exacerbates the deficit.)

A second approach is to calculate the *variance components* of the overall common cause variation in the measurements through a *nested design* supplied by the batch-to-batch and within-batch samples. Assume that from each of n randomly sampled batches, m random within-batch samples are taken. If the within-batch samples are independent,

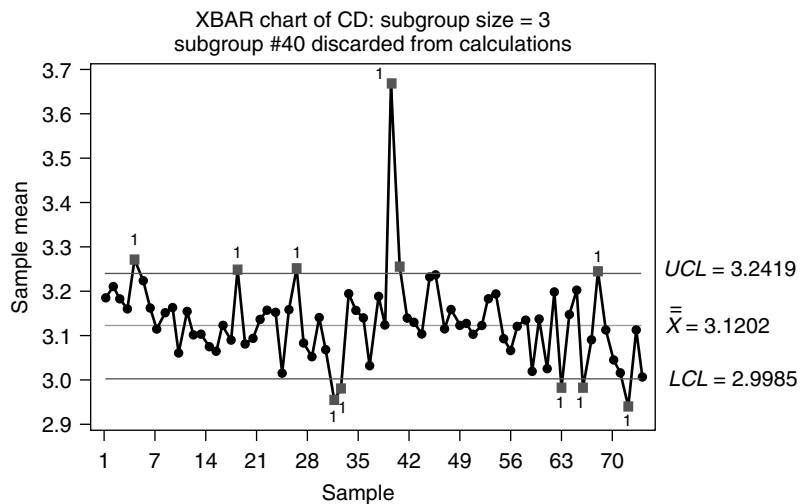


Figure 1 X-bar chart of CD measurements

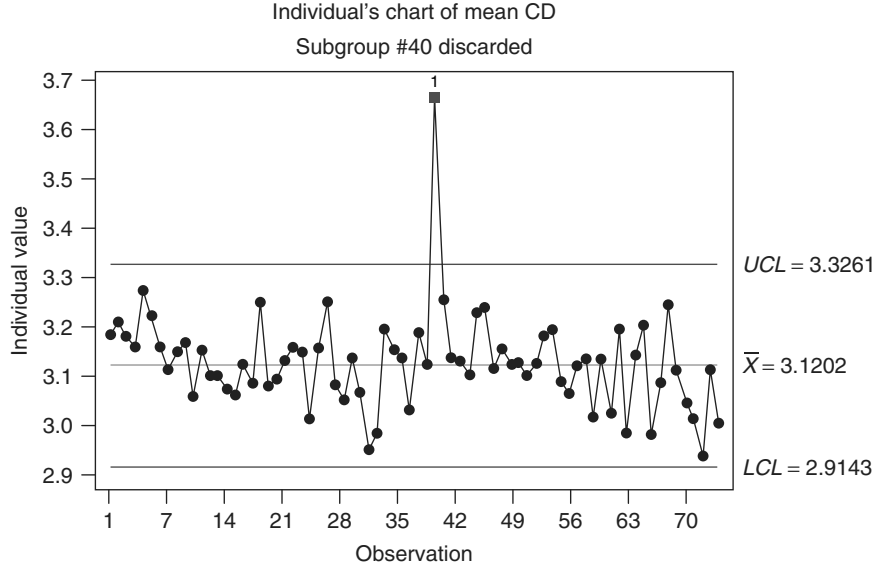


Figure 2 Individuals chart of mean CD

the variance of the mean of the m within-batch measurements is estimated as below (see [2, 4]).

$$\hat{\sigma}_{\bar{x}}^2 = \hat{\sigma}_B^2 + \frac{\hat{\sigma}_W^2}{n} \quad (1)$$

where $\hat{\sigma}_B^2$ = estimated batch variance and $\hat{\sigma}_W^2$ = estimated within-batch variance.

Standard “ 3σ ” control limits for the batch average are given by

$$\bar{\bar{X}} \pm 3 \left(\hat{\sigma}_B^2 + \frac{\hat{\sigma}_W^2}{n} \right)^{1/2}$$

This approach allows for the possibility of varying sample sizes from batch to batch, as well as the introduction of additional components of variability, should that be necessary. Control charts of the within-batch component of variance can also be applied, if desired [2]. If the within-batch variability is random, the standard S chart should suffice; if the variability is nonrandom, however, other control chart schemes may be necessary (e.g., individuals charts on S or a suitable nonlinear transform, S^*).

Use of Multivariate Control Charts

High-quality products are commonly described by more than a few quality characteristics (Y 's), each of

which must be controlled within tight specifications to maintain customer satisfaction. For example, at a polishing process during silicon wafer manufacturing, removal rate, thickness uniformity, and surface quality must be maintained. In paper processing, basis weight, moisture content, and tensile strength, among other properties of the paper must be controlled by the process.

Consider a series of j quality characteristics ($Y_1, Y_2, Y_3, \dots, Y_j$) that are measured on each processing unit (batch). Traditionally, univariate-type control charts (individuals, exponentially weighted moving average (EWMA), cumulative sum (CUSUM), etc.) have been applied to each of the Y_i 's. Two practical issues arise when doing so.

1. A proliferation of control charts to maintain and evaluate can result. This scenario is particularly serious if the manufacturer has different customers who receive products with different target values for each Y_i , and control charts are kept for each Y_i /customer combination. As a consequence, there are manufacturing organizations where thousands of univariate control charts are supposed to be kept and monitored.
2. The Y_i 's are likely to be correlated, in which case, true out-of-control conditions may be overlooked by the assessment of the collection

4 Control Charts for Batch Processes

of individual charts (*see Multivariate Control Charts Overview* and [5] for further discussion).

An answer to this dilemma is the multivariate control chart based upon Hotelling's T^2 statistic. The T^2 statistic collapses the vector of j variables $Y = (Y_1, Y_2, Y_3, \dots, Y_j)$ into the scalar value below

$$T^2 = (Y - \mu)'S(Y - \mu) \quad (2)$$

where $\mu = (\mu_1, \mu_2, \dots, \mu_j)$ is the vector of "expected values" for Y , and S = the observed **covariance** matrix, estimated when Y is "in control". For a complete description of this chart and its application to batch processes, *see Multivariate Control Charts Overview, Hotelling's T^2 Chart* [6, 7].

Consider the CD data analyzed earlier, but imagine that instead of three measurements of the same characteristic, CD, each Y_i is a separate and distinct quality characteristic. The correlation matrix (*see Correlation*) is given below and demonstrates that there are significant pairwise correlations between the three Y 's, an important feature of the data of which the T^2 statistic takes advantage.

1	0.404	0.580
0.404	1	0.388
0.580	0.388	1

A control chart of the T^2 statistic for these data is provided in Figure 3. Clearly, the chart reveals that subgroup #40 is out of control, but in even a more pronounced fashion as compared with that of the individuals chart in Figure 2.

Utilizing the Process Input Variables (X 's)

In recent years, modern manufacturers have installed on-line sensors and gauges, which provide the capability to capture measurements of numerous process input variables (X 's), such as temperature, pressure, line speed, and mixing RPM. MacGregor [5] summarizes the common features of this data. Typically, many more X 's are captured on-line than Y 's are measured in the off-line laboratory. In addition, because these measurements are conducted on-line, they are available instantaneously, while the Y 's can lag by hours. Moreover, the potential sampling rate of these variables is on the order of seconds or less. Finally, the sensors/gauges that capture the X 's do so more precisely and accurately than the gauges used in the laboratory to produce the Y 's. MacGregor [5] also notes that special causes of variation in the Y 's occurring during the batch process generally leave their imprint on the X 's as well. Therefore, massive amounts of relatively cheap and high-quality data may be available in real time on

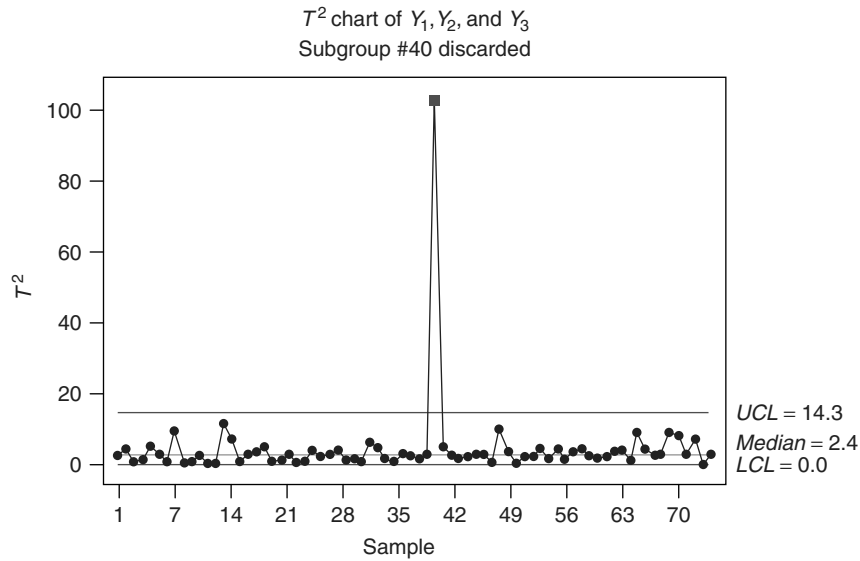


Figure 3 T^2 chart of three CD measurements

these process inputs (X 's), quite apart from the Y 's captured on each batch. This capability provides the potential to address a serious limitation of batch processing SPC mentioned earlier: the inability to respond to special causes while the batch is being processed.

If such data on process inputs are to be utilized for SPC, then some general characteristics of the data need to be recognized.

1. Data overload will occur. Implementing very high sampling rates of large numbers of X 's will generate massive data, quickly overloading the SPC practitioner, especially if traditional univariate control charts are utilized.
2. Multicollinearity in the X data is likely. Although numerous X 's can be captured, each X cannot be independent of the remaining X 's; in fact, many will be correlated, in some cases highly correlated.
3. Missing data is common. For a variety of reasons data on one or more of the X 's may be missing or invalid (e.g., due to a contaminated or completely failed on-line sensor).

Multivariate Projection Methods

Control charting methods for analyzing and monitoring such data must be able to address these particular features to be effective. A class of methods known as *multivariate projection methods*, which reduce high-dimensional data into low-dimensional "latent variables" have been proposed and successfully applied in recent years (see [8–10]). Statistical projection methods, such as principal components analysis (PCA) and partial least squares or projection to latent structures (PLS) are used to generate the latent variables. Frequently, just a handful of latent variables can explain a large proportion of the variability in a large dataset of intercorrelated variables and, because the latent variables are orthogonal, monitoring each of them is valuable. For example, through the use of PCA a set of k coefficients or "loadings", $\mathbf{P} = (P_1, P_2, P_3, \dots, P_k)$ are generated that linearly combine a set of m variables, $\mathbf{X} = (X_1, X_2, X_3, \dots, X_m)$ into k principal components, $\mathbf{T} = (T_1, T_2, T_3, \dots, T_k)$ or "scores", where $\mathbf{T} = \mathbf{P}^T \mathbf{X}$ (see Figure 4).

These scores can be charted separately or analyzed jointly through the T^2 chart.

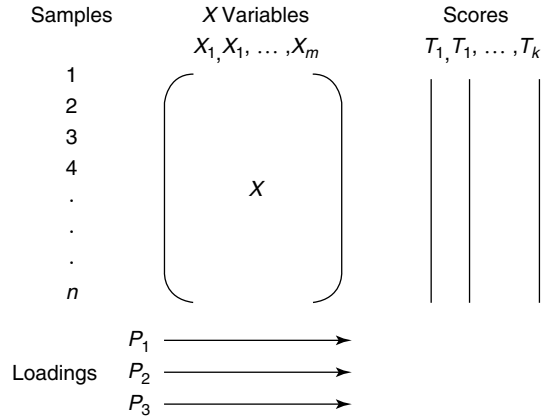


Figure 4 Latent variables generated by principal components analysis

Using PLS both the historical X data and Y data are inputs, and the objective of multivariate projection changes to that of extracting the latent variables in the X space that most explain the variation in the Y space. MacGregor and Kourti [11] provide an overview of the PLS methodology. Kourti and MacGregor [9] describe two case studies in which large-dimensional X and Y spaces were reduced through PLS into two- and three-dimensional "scores". These scores can then be monitored via traditional univariate control charts or the T^2 chart.

An additional chart that is recommended is one that monitors the squared prediction error (SPE), the sum of the SPEs of the new observations, $X_{i,\text{new}}$, from the predicted observations, $\hat{X}_{i,\text{new}}$.

$$SPE = \sum_{i=1}^m (X_{i,\text{new}} - \hat{X}_{i,\text{new}})^2 \quad (3)$$

A control chart for SPE that is easy to construct and interpret is recommended by Nomikos and MacGregor [10]. Any value of the SPE above its upper control limit suggests that the process inputs have shifted away from their predicted values. A "contribution plot", similar to a Pareto chart (see **Pareto Chart**), which identifies those X 's that provide the greatest contribution to variability in the "scores", can also be utilized for diagnosis [11]. Kourti and MacGregor [12] discuss how to best monitor and diagnose process variation using the PLS methodology.

Like any other control chart procedure the setup and application of PLS procedures for SPC requires

the availability of historical data from “good” batches (whose processing was subject only to common cause variability). From such data the appropriate PCA/PLS methodology is applied, and the PCA loadings are calculated and stored. (The user must decide how many dimensions of the data (latent variables or “scores”) to utilize.) From these loadings the new observations are calculated, and appropriate control charts are applied to monitor the batch process.

Nomikos and MacGregor [10] also describe an extension of this approach where the X space is three dimensional; in particular, where the X ’s represent measurements from within-batch sampling of time trajectories of process variables (e.g., pressure or temperature) taken over repeated batches. The methodology can be further extended to include values of initial batch setup variables as well. Applying the PLS methodology to the combined sets of data allows monitoring of a batch process in real time.

Conclusions

Batch processing presents challenges to the unsuspecting control chart practitioner. First, typical sampling of a batch for product quality measurement may require that multiple sources of common cause variation be accounted for in the design of the control chart. The issue can be addressed by ensuring that all sources of common cause variability, particularly the batch-to-batch source, are included in the assessment of the process variability used to construct the control limits. In addition, the user must recognize that the design of the chart must account for both primary and secondary unit subgroup sample sizes.

Second, relying on traditional univariate Shewhart-type charts of product quality measurements taken long after the batch is completed may not be an effective control scheme. Many batch processes require relatively long processing times of volumes in which the manufacturer has a significant financial investment. The manufacturer is motivated to monitor a batch in real time so that process disturbances can be identified quickly and midcourse corrections to the batch can be applied, if possible. Modern batch manufacturers routinely collect on-line massive amounts of data on processing conditions that are available for analysis. Multivariate projection methods can be used to compress a large number of these process variables (X ’s) into a small number of latent variables (T ’s)

that account for a high proportion of the variability in output or product quality (Y ’s) when the process is “in control”. Monitoring the latent variables *via* univariate or multivariate control charts that are easy to construct and interpret provides the potential for the user to respond in real time, instead of after the fact, to those special causes that arise as the batch is being processed.

References

- [1] Joshi, M. & Sprague, K. (1997). *Chapter 24: Obtaining and Using Statistical Process Control Limits in the Semiconductor Industry. Statistical Case Studies for Industrial Process Improvement*, ASA-SIAM, pp. 337–356.
- [2] Woodall, W. & Thomas, E. (1995). Statistical process control with several components of common cause variability, *IIE Transactions* **27**, 757–764.
- [3] Roes, K., Does, R. & Schurink, Y. (1993). Shewhart-type charts for individual observations, *Journal of Quality Technology* **25**, 188–198.
- [4] Box, G., Hunter, J. & Hunter, W. (2005). *Statistics for Experimenters*, 2nd Edition, John Wiley & Sons, New York, pp. 349–352.
- [5] MacGregor, J. (1996). Using on-line process data to improve quality, *ASQ Statistics Division Newsletter* **16**(2), 6–13.
- [6] Mason, R. & Young, J. (2001). *Multivariate Statistical Process Control with Industrial Applications*, ASA-SIAM.
- [7] Mason, R., Chou, Y. & Young, J. (2001). Applying Hotelling’s T^2 statistic to batch processes, *Journal of Quality Technology* **33**, 466–479.
- [8] Kresta, J., MacGregor, J. & Marlin, T. (1991). Multivariate statistical monitoring of process quality performance, *Canadian Journal of Chemical Engineering* **69**, 37–47.
- [9] Kourti, T. & MacGregor, J. (1998). Projection methods for process analysis and statistical process control: examples and experiences from industrial applications, *ASQ Annual Quality Congress Proceedings* **52**, 73–77.
- [10] Nomikos, P. & MacGregor, J. (1995). Multivariate SPC charts for batch processes, *Technometrics* **37**, 41–59.
- [11] MacGregor, J. & Kourti, T. (1995). Statistical process control of multivariate processes, *Control Engineering Practice* **3**, 403–414.
- [12] Kourti, T. & MacGregor, J. (1996). Multivariate SPC methods for process and product monitoring, *Journal of Quality Technology* **28**, 409–428.

Further Reading

- Roes, K. & Does, R. (1995). Shewhart-type charts in non-standard situations, *Technometrics* **37**, 15–24.

DONALD K. LEWIS

Precontrol

Description of the Method

Precontrol monitors the acceptability of a process by comparing individual process values against their specification limits. Its goal is the prevention of nonconforming units. Precontrol, also referred to as *stoplight control*, classifies measurement data into several groups and monitors the process on the basis of the resulting attribute data. In the case of two-sided tolerance, precontrol lines are drawn at the midpoints between the target and the lower (upper) specification limits. The area between the precontrol limits is referred to as the *green zone*, the areas between the precontrol limits and the specification limits are called the *yellow zones*, and the areas outside the specification limits are referred to as the *red zones*.

Initial process acceptability (capability) is established through a simple setup approval procedure. Five consecutive units are taken from the process, and all five measurements must fall within the green zone for the production to commence. If a single yellow result is observed, additional observations are taken and the count of units falling into the green zone is started anew. If two consecutive yellows or a red are encountered, the process is adjusted (see below) and the approval process is restarted.

Following setup approval, two consecutive units are sampled from the process at periodic intervals. If both units fall into the green zone, or if only one unit falls into a yellow zone, production continues. The process is stopped if both units fall into yellow zones, or if one or more unit falls into a red zone. After stopping the process, one must find the cause of the problem and take action. If both units fall into the same yellow zone, it is likely that the process is off-target. In this case, level adjustments are made using feedback control on the deviations from the target. If consecutive yellows (or reds) are on opposite sides of the target, it is likely that the variability has increased. Causes for the extra variability need to be identified and their occurrence must be prevented in the future. The setup approval stage is repeated after each stop. Five consecutive units must be in the green zone before production can resume.

It is suggested that the frequency of the sampling be set such that the average number of sample pairs between stoppages is about six. For example, if samples are taken every hour and stoppages occur on average every 6 h, the sampling interval is set correctly. The sampling interval should be cut in half if stoppages occur on average every 3 h, and it should be doubled if stoppages occur on average every 12 h.

Several modifications of the basic precontrol technique are discussed in the literature. A two-stage sampling approach that takes an additional sample if the initial sample yields ambiguous results is discussed by Salvia [1]. A modification that replaces the specification limits with control limits (which use the standard deviation of process measurements) is proposed by Gurska and Heap [2].

Historical Note

Precontrol was first described by Satterthwaite [3] who attributes the development of the technique to himself and C.W. Carter, W.R. Purcell, and D. Shainin. The method is described in Bhote [4, 5], Juran and Gryna [6], Shainin [7], and Shainin and Shainin [8].

Example

Bhote [4, p. 228] describes an example that deals with the thickness of certain chrome and gold deposits. For chrome the target thickness is 9 units, with specification limits of 6 and 12, resulting in precontrol limits of 7.5 and 10.5. For gold the target thickness is 32, with specification limits of 23 and 41, resulting in precontrol limits of 27.5 and 36.5, respectively. The data in Table 1 indicate no problems with the capability of the thickness of gold deposits. However, problems are indicated for the thickness of chrome deposits. Two consecutive yellows on the same side of the target are observed on April 4. The first sample on April 5 leads to two consecutive units on opposite sides, while in the second sample one of the observations is in the red zone. The process must be investigated and reasons for the unacceptable results must be found. Often this is easier said than done, as considerable detective work may be required. Bhote makes no mention of how the process was adjusted in this particular example.

2 Precontrol

Table 1 Specification limits. Chrome: 6 and 12. Gold: 25 and 41^(a)

	Unit 1		Unit 2		
	Chrome	Gold	Chrome	Gold	
April 3	9.6	32.0	10.2	30.0	
April 4	10.2	33.0	9.6	33.5	
April 4	10.7	33.5	11.2	35.0	Two yellows on same side (chrome)
April 4	7.6	30.5	8.6	25.4	
April 5	10.5	34.0	6.5	34.5	Two yellows on opposite sides (chrome)
April 5	12.2	30.5	10.2	32.0	One red (chrome)
April 6	7.6	33.5	9.6	35.0	

^(a) Values outside the green zone (7.5 and 10.5 for chrome, and 27.5 and 36.5 for gold) are in boldface

Differences between Precontrol and Control Charts

Both precontrol and control charts take periodic observations from a process. Despite their similarities, the two techniques differ with respect to their objectives. The goal of precontrol is the prevention of nonconforming units. It is an algorithm for controlling a process on its tolerances (without explicit consideration of stability), while control charts monitor the stability of the process over time (without specifically looking at its capability). Control charts (*see Control Charts, Overview*), process capability, and links between control charts and capability (*see Control Charts and Process Capability*) are described in more detail elsewhere.

For very capable processes, a control chart may signal lack of control even though the shifted process is still well within its specifications. It has been argued that in this case control charts are of limited value, as operators would question the wisdom of disturbing a process that produces good items.

Strengths of Precontrol

In situations where processes are highly capable, precontrol offers a simple useful technique for checking the capability of the process. Precontrol reduces the temptation of interfering with processes that are still capable. The procedure will not flag unstable processes as long as they remain capable. The approach has the virtue of drawing management attention only to priority problems.

Precontrol was originally developed with machining operations in mind. There, the operator is faced with the problem of first setting up the machine, and

deciding whether the setup has been done properly and whether the machine is ready for full production. In many machining operations the output variable will drift owing to factors such as tool wear. Simple feedback algorithms that bring the process back to its target are usually possible, and operators with considerable machining experience know how to make these adjustments.

Precontrol (just like control charts) forces the operator to take periodic measurements on a process. Quite often it is not the limits or the zones on these charts that are important, but the fact that process information is obtained. Sudden shifts, trends, cycles, and the presence of overcompensation become visible.

Weaknesses of Precontrol

Precontrol is easy to implement, but its *ad hoc* decision rules on grouped data have led to criticism. In particular, it has been questioned whether the small setup inspection sample size of five consecutive items is able to detect processes with low capability. Furthermore, grouping discards information, and there is a real danger that poor processes slip through the certification stage (Ledolter and Swersey [9]).

Also, there are important reasons why one should worry about process instability even though products are still acceptable. (a) Specifications may change (what is good today may not be good tomorrow); (b) unstable processes may not stay capable for long (within specifications today, but not tomorrow); (c) any deviation from a target, even if the unit is within the specification limits, leads to a loss; and

(4) it is difficult to learn about root causes if a process is not in statistical control.

Precontrol is poorly suited for processes with low capability, as its frequent signals will lead to corrective actions that in the absence of profound process understanding amount to little more than process tampering. Processes with low capability must be improved. This is best done by first making sure that assignable (special) causes of variation have been identified and removed. Typically, this requires a thorough improvement cycle that utilizes control charts as well as other improvement techniques. It would be optimistic to expect that the *ad hoc* adjustments of precontrol can fix the problem. Precontrol alone provides little guidance for process improvement.

Further Discussion

Detailed discussion of the strengths and weaknesses of precontrol can be found in Traver [10], Logothetis [11], Mackertich [12], Ermer and Roepke [13], Steiner [14], Ledolter and Swersey [9], and Ledolter and Burrill [15].

References

- [1] Salvia, A. (1988). Stoplight control, *Quality Progress* **21**, 39–42.
- [2] Gurska, G. & Heaphy, M. (1991). Stop light control – revisited, *ASQC Statistics Division Newsletter Fall* **11**(4), 10–11.
- [3] Satterthwaite, F. (1954). A simple effective process control method, Report 41-1, Rath and Strong, Boston.
- [4] Bhote, K. (1988). *Strategic Supply Management*, American Management Association, New York.
- [5] Bhote, K. (1991). *World Class Quality*, American Management Association, New York.
- [6] Juran, J. & Gryna, F. (1988). *Quality Control Handbook*, 4th Edition, McGraw-Hill, New York.
- [7] Shainin, D. (1984). Better than good old Xbar and R charts asked by vendees, in *ASQC Quality Congress Transactions*, American Society for Quality Control, Milwaukee, pp. 302–307.
- [8] Shainin, D. & Shainin, P. (1989). Pre-control versus Xbar and R charting: continuous or immediate quality improvement, *Quality Engineering* **1**, 419–429.
- [9] Ledolter, J. & Swersey, A. (1997). An evaluation of pre-control, *Journal of Quality Technology* **29**, 163–171.
- [10] Traver, R. (1985). Pre-control: a good alternative to Xbar-R charts, *Quality Progress* **18**(9), 11–14.
- [11] Logothetis, N. (1990). The theory of pre-control. A serious method or a colorful naivity? *International Journal of Total Quality Management* **1**, 207–220.
- [12] Mackertich, N. (1990). Precontrol versus control charting. A critical comparison [with discussion by Dorian Shainin], *Quality Engineering* **2**, 253–268.
- [13] Ermer, D. & Roepke, J. (1991). An analytical analysis of pre-control, in *ASQC Quality Congress Transactions*, American Society for Quality Control, Milwaukee, pp. 522–527.
- [14] Steiner, S. (1997). Pre-control and some simple alternatives, *Quality Engineering* **10**, 65–74.
- [15] Ledolter, J. & Burrill, C. (1999). *Statistical Quality Control: Strategies and Tools for Continual Improvement*, John Wiley & Sons, New York.

JOHANNES LEDOLTER

Economic Design of Control Charts

Introduction

Control chart design is a problem of practical and academic interest that involves finding the parameters of a control chart that one uses to monitor a stable process. These parameters include the sample size, control limit width, and sampling frequency. Control chart design is a problem of much importance in industry because with good design a firm may ensure that they produce products or services of good quality, or, on the other hand, minimize the costs of producing products or services of poor quality.

There are three general methods used to design control charts for use in practice today. These are designs using a heuristic, statistical design methods, and economic design methods. Another method, called *economic statistical design*, finds designs that meet statistical constraints as in statistical design but, in addition, finds the statistical design that is economically optimal.

In this chapter I will discuss economic and economic statistical design in some detail but I will also provide a little background on the former two methods. My emphasis will be on what practitioners can draw from many years of academic research. I suspect that if one were to rank these methods based upon their popularity in practice, a winner by a large margin would be heuristic design. If the ranking criteria were, however, popularity in the academic literature in statistics, the winner by an equally large margin would be statistical design. The academic literature in production and management science would probably find the economic criterion to be the most popular.

What value, then, does the economic design and economic statistical design literature have in the sense of allowing practitioners to improve their decision making regarding design? My opinion is that the research allows practitioners to actually design charts under an economic or economic statistical criterion were they to develop some method of measuring costs and other system parameters necessary to find a design. I address this problem and give some solutions later in this paper. Furthermore, there has been some recent research that shows that simple Shewhart

type control charts do not have an economic disadvantage over other more complex methods such as likelihood-ratio cumulative sum (CUSUM) charts unless certain conditions are met. In yet other situations, it may be economically optimal to employ methods other than control charts to monitor quality. This research helps to point the practitioner in the right direction in terms of choosing the optimal method, or policy, for maintaining control even if they did not wish to implement an economic design.

The remainder of the chapter is organized as follows: In the second section I provide a review of the literature on control chart design with an emphasis on economic and economic statistical design. The following section outlines how a practitioner may actually use current methods to find economic and economic statistical designs and shows how the recent literature on policy selection may be used. In the last section I provide a summary.

Economic and Economic Statistical Design

The simplest design one could employ is Shewhart's [1] heuristic, which for variable control charts such as the $\bar{X}$ chart, is to choose samples of size 4 or 5 and to use three sigma control limits (see Duncan [2] for an explanation of heuristic control chart design). One then should sample, say, once every hour for high volume processes. This design has desirable behavioral, statistical, and economic characteristics. More specifically, this design is easy to use and understand, works very well in terms of signaling large shifts very quickly, and, in some cases, is close to a minimum cost design. One can improve a heuristic design very easily using a simple formula for the optimal intersample interval (see McWilliams [3]).

Statistical design involves finding the control limit width and the sample size such that the average run length (ARL) in control is high and the ARL is low when the process is out of control (*see Average Run Lengths and Operating Characteristic Curves*). The ARL is the average number of samples taken before a shift is signaled. Statistical design is appealing because the control chart will give few false signals of poor quality when quality is acceptable and will signal a process shift to a state of poor quality relatively rapidly. Woodall [4] discusses statistical design and a thorough treatment of policy selection

2 Economic Design of Control Charts

using statistical criteria is given by Reynolds and Stoumbos [5].

Economic control chart design involves finding the parameters of a control chart such that all of the costs associated with producing a product or service related to **quality** are minimized (*see Cost of Quality*). These costs include the cost of sampling, the cost of searching for false assignable causes of poor quality, the costs of searching for a true **assignable cause** and repairing the system when it occurs, and the cost of producing poor quality products or services. The particular parameters of interest in control chart design are the sample size, the control limit width, and the sampling frequency.

The problem has received much attention in the academic literature. Perhaps the initial thrust in the development of an economic model of a process controlled by a control chart was provided by Girshick and Rubin [6]. Their model provided the basis for many of the practical models that followed but alone it is difficult to implement because of the complex mathematics involved in finding a solution. One of the models owing allegiance to the Girshick and Rubin [6] model is that of Duncan [2].

Duncan [2] developed a model to find the economic design of an $\bar{X}$ chart. He assumed one would use a control chart in the usual way; that is, one takes a sample of size n at regular intervals of h hours and plots the average, or $\bar{X}$ on a control chart with control limits of standardized width k . Some of the important assumptions of the Duncan model are that only a single assignable cause of poor quality will occur; in the literature this assignable cause is called the *expected shift*. (Later findings showed that this assumption is not restrictive in that a single cause model may be a good approximation of a multiple cause model.) There are a number of other models similar to Duncan's including a model provided by von Collani [7] that makes the **estimation** and solution process simpler.

Many other applications to other types of control chart have been done; some examples are the $\bar{X}/R$ combination by Saniga [8], the CUSUM chart by Taylor [9], and the attribute chart by Chiu [10]. Goel [11] compared the economic performance of the $\bar{X}$ chart to the CUSUM chart (*see Cumulative Sum (CUSUM) Chart*).

Lorenzen and Vance [12] provided an extension of the popular Duncan [2] model. This model is more general in the sense that production can continue or

be stopped during a search for the assignable cause and production can be either continued or stopped during repair. Otherwise, the standard assumptions in economic control chart design are made; these are that the time in control follows the negative exponential distribution, a single shift of known size can occur, and the other cost and system parameters are deterministic. The model is defined as follows where C is the expected cost per hour:

$$C = \{C_0/\lambda + C_1(-t + nE + h(ARL_2) + \delta_1 T_1 + \delta_2 T_2) + sY/ARL_1 + W\} \\ \div \{1/\lambda + (1 - \delta_1)sT_0/ARL_1 - t + nE + h(ARL_2) + T_1 + T_2\} + \{(a + bn)/h\} \\ \times 1/\lambda - t + nE + h(ARL_2) + \delta_1 T_1 + \delta_2 T_2\} \\ \div \{1/\lambda + (1 - \delta_1)sT_0/ARL_1 - t + nE + h(ARL_2) + T_1 + T_2\} \quad (1)$$

The terms are as follows:

n = sample size

h = hours between samples

L = number of standard deviations from control limits to center line for the $\bar{X}$ chart

k = reference value for the CUSUM chart

h = decision interval for the CUSUM chart

g = intersample interval

$tt = [1 - (1 + \lambda g)e^{-\lambda g}]/[\lambda(1 - e^{-\lambda g})]$

$s = e^{-\lambda g}/(1 - e^{-\lambda g})$

ARL_1 = average run length while in control

ARL_2 = average run length while out of control

λ = 1/mean time process in control

δ = number of standard deviations slip when out of control

E = time to sample and chart one item

T_0 = expected search time when false alarm

T_1 = expected time to discover the assignable cause

T_2 = expected time to repair the process

δ_1 = 1 if production continues during searches
= 0 if production ceases during searches

δ_2 = 1 if production continues during repair
= 0 if production ceases during repair

C_0 = quality cost/hour while producing in control

C_1 = quality cost/hour while producing out of control ($> C_0$)

Y = cost per false alarm

W = cost to locate and repair the assignable cause
 a = fixed cost per sample
 b = cost per unit sampled.

Without listing any additional articles here, it is safe to write that there are literally hundreds of applications and extensions of the Duncan [2] and Lorenzen and Vance [12] model and other models for economic design appearing in the academic literature in statistics, production, and operations research.

Lorenzen and Vance [12] and Woodall [13–15] among others, have addressed some of the practical shortcomings of using economic designs. These shortcomings include the fact that the designs are hard to determine mathematically, the cost and other system coefficients are difficult to measure accurately, and the solutions are perhaps difficult to implement because they yield designs that have odd intersample intervals (such as sampling every 1.28 h). A related example is one where an economic design allows for an excessive number of false searches or equivalently, a relatively high probability of a false search. Saniga [16] found that Duncan's [2] 25 example problems have type I error probabilities as high as 0.1926. In this case, the excessive number of false searches that would take place might result in users disregarding the signals from the control chart or even adjusting the process when it should not be adjusted, resulting in an increase in common cause variability.

Another problem with economic design is that an optimal solution may yield designs not in accord with modern principles of **quality management** such as those provided by Deming [17] (see **Total Quality Management (TQM)**). Here, optimal solutions might have not only high type I error probabilities but also high type II error probabilities, which would allow defective material or services to be passed on to the consumer. Deming, in essence, proposes that "tight" quality standards such as "make it right the first time" and "eliminate rework" be implemented. (Deming, considering his emphasis on continuous improvement, might not even be in favor of continuous monitoring.) Even contemporary methods such as Six Sigma (see **Six Sigma**) emphasize extremely "tight" quality standards.

These issues as well as others were addressed in a design called *economic statistical design* by Saniga [16]. Here, statistical constraints restricting ARL in control and out of control as well as

the intersample interval were placed upon an economic model. These statistical constraints could be expressed equivalently in terms of the error probabilities or the average time to signal, or ATS. The resulting solution has little of the shortcomings mentioned by Lorenzen and Vance [12] and Woodall [13–15] in that false searches would be limited, short ARLs and short, or even, regular intersample intervals could be ensured. Economical statistical designs would then be in accord with what is required in practice in the sense of providing designs which would meet the behavioral needs of the user as well as the economic and statistical needs of the firm's strategic decision makers.

An additional and unexpected finding of Saniga [16] is that pure statistical designs could be made more restrictive to improve economic performance; certainly a counterintuitive result.

Designing a Control Chart to Minimize Cost

In practice, perhaps the first choice a practitioner must make when he decides to use a control chart to monitor a stable process is to choose an optimal policy. There are three general policies available: do nothing, use a control chart, or use a regular search policy (see **Control Charts for Attributes; Variable Sampling Rate Control Charts; Control Charts, Selection of**). If the choice is to use a control chart, one must choose between any number of alternative control charts. Saniga and Montgomery [18] investigated the performance of alternative methods in a large experiment in which they found the optimal cost solutions for using a variables or attributes control chart, a regular search policy, or simply doing nothing. The last one is called a *never search* policy. While the last was considered an alternative policy it is true that, as Saniga and Montgomery [18, p. 261] mention, "it is difficult to visualize a real-world production system" with system parameters that would deem this policy optimal and "the practical relevance of this policy is questionable". In this work the authors used an $\bar{X}/R$ chart for variables sampling and a fraction defective chart for attributes sampling.

Saniga and Montgomery [18] found that, generally, for detecting small to medium size shifts the $\bar{X}/R$ combination or the regular search policy is economically optimal with the different choices being

based on search time ratios and sampling cost ratios. For larger shifts a fraction defective chart is optimal in three of the four regions investigated. Only if both the search time ratios and sampling cost ratios were low would an $\bar{X}/R$ combination be preferred.

While Saniga and Montgomery [18] address the issue of choosing between variables and attribute Shewhart charts it is also of interest to examine the possibility of using CUSUM charts instead of Shewhart charts. While CUSUM charts are more difficult to design, implement, and use and there is a higher training cost associated with their use, they do have an advantage of providing the optimal ARL for detecting a process shift for any desired in control ARL. Moustakiedies [19] has provided this proof. Assuming a practitioner is interested only in statistical performance, Reynolds and Stoumbos [5] provide some guidelines for **control chart selection** under a wide variety of circumstances. They suggest that one should use a CUSUM or **exponentially weighted moving average** chart (see Lucas and Saccucci [20]) using samples of size 1. Cost and ease of design or use, for example, are not factors considered by the authors. Consider cost for example. If fixed cost is not zero, which is unlikely using any standard method of cost accounting, a policy with $n = 1$ may not be economically optimal unless fixed cost is small. Of course, there are situations where only statistical performance matters and Reynolds and Stoumbos' [5] recommendation would then be very useful.

Looking back at the economic criterion again, Goel [11] found that CUSUM charts did not have an economic advantage over $\bar{X}$ charts unless sample sizes smaller than optimal were employed. Chiu [10] and Von Collani [7] reached similar conclusions about the cost advantages of CUSUMs over $\bar{X}$ charts. The conclusions of Goel, Chiu, and Von Collani, however, were based upon very small experiments.

A thorough investigation of the regions of cost and other system parameter configurations where a CUSUM chart would be preferred to a particular Shewhart chart (the $\bar{X}$) for monitoring the mean of a process was done by Saniga *et al.* [21]. Here, the authors found economic designs for CUSUM charts and for $\bar{X}$ charts for a wide range of system parameter configurations covered in an experimental design. They assumed all costs for using the different control charts would be the same. Since the costs are the same and the CUSUM chart has an advantage in

ARL the CUSUM chart would always be preferred economically. Saniga *et al.* then used regression trees to determine regions where the CUSUM chart had somewhat of an advantage over the $\bar{X}$ chart. They found, for one part of the three part experiment, that if the fixed cost of sampling is small, the expected shift is small and the cost of either a false search or a search for an assignable cause is large an optimal policy is to use a CUSUM chart. For another part of the experiment they found that small fixed costs of sampling lead to CUSUM charts having a decided advantage and, also, the CUSUM chart may be advantageous even for midrange values of the fixed cost if the shift is small and the time to sample and chart an observation is large. Saniga *et al.* [22] also found that a regular search policy is preferred to a control chart for this large experiment for many cases when the fixed costs of sampling are large and the expected shift size is small.

In an earlier related study, Saniga *et al.* [22] found, using a small experimental study, that CUSUM charts are economically preferred to $\bar{X}$ charts when there are restrictions on sample size and sampling interval, when there are statistical constraints on ARL that would occur as in economic statistical design, and when there are two components of variance (*see Variance Components*).

Once the policy decision is made the practitioner needs to estimate the parameters an economic model will require and subsequently, find the optimal design. Estimation of the cost parameters requires some knowledge of cost accounting; a practical presentation of that subject is available in Jones [23]. Some of the other parameters that need to be estimated can be done simply by sampling over time. For example, optimal design involves providing as input the expected search time when a false alarm occurs. A practitioner can record how long each search takes over a period of time and use the mean as input into the economic model. Determination of the expected shift size can be done in a similar fashion; one can record the mean of the distribution of output when a shift occurs. Finding the mean of these means will provide a good estimate of the expected shift. And in a similar manner the other times can be determined.

Finding the actual economic design parameters can be done in a number of ways. Chiu and Wetherill [24] have developed a simple method to determine the design parameters of an $\bar{X}$ chart assuming that the **power** of the control chart is either 0.90 or 0.95. Von

Collani [25] also provides a simple procedure to find the design parameters of a control chart; his method is based upon using a model developed by him rather than the usual Duncan [2] model.

A number of computer programs are also available for finding economic and economic statistical designs. These include that of Montgomery [26], who provided a program allowing a user to find the optimal economic design of an $\bar{X}$ control chart based upon Duncan's [2] model. A program for finding economic, economic statistical, and statistical designs for attribute charts is provided by Saniga *et al.* [27]. Another program by McWilliams *et al.* [28] can be used to find economic, economic statistical, and statistical designs of $\bar{X}/R$ and $\bar{X}/s$ control charts.

The programs mentioned above are readily available, easy to use, and perhaps equally important, permit the user to do a sensitivity analysis of the solution to see what happens when changes to the optimal solution are made; the new solution may meet the needs of the practitioner better than a direct solution.

Summary

Economic and economic statistical control chart designs are some of the several design methods available to users. In this chapter, I have presented some of the background work leading to our current knowledge about these type of designs. I have also provided users with some guidelines (and references) on actually designing and using a control chart under this criterion. Many additional references on the subject exist and are readily available.

References

- [1] Shewhart, W.A. (1931). *Economic Control of Quality Manufacturers Product*, Van Nostrand Reinhold, Princeton.
- [2] Duncan, A.J. (1956). The economic design of $\bar{X}$ charts used to maintain current control of a process, *Journal of the American Statistical Association* **51**, 228–242.
- [3] McWilliams, T.P. (1994). Economic, statistical, and economic-statistical $\bar{X}$ -chart designs, *Journal of Quality Technology* **26**, 227–238.
- [4] Woodall, W.H. (1985). The statistical design of quality control charts, *The Statistician* **34**, 155–160.
- [5] Reynolds Jr, M. & Stoumbos, Z. (2004). Control charts and the efficient allocation of sampling resources, *Technometrics* **46**, 200–214.
- [6] Girshick, M.A. & Rubin, H. (1952). A Bayes approach to a quality control model, *Annals of Mathematical Statistics* **23**, 114–125.
- [7] Von Collani, E. (1987). Economic process control, *Statistica Neerlandica* **41**, 89–97.
- [8] Saniga, E. (1977). Joint economically optimal design of $\bar{X}$ and R control charts, *Management Science* **24**(4), 420–431.
- [9] Taylor, H.M. (1968). The economic design of cumulative sum control charts, *Technometrics* **10**, 479–488.
- [10] Chiu, W.K. (1974). The economic design of CUSUM charts for controlling normal means, *Applied Statistics* **23**, 420–433.
- [11] Goel, A.L. (1968). *A comparative and economic investigation of $\bar{X}$ and cumulative sum control charts*, University of Wisconsin Press, Unpublished PhD dissertation.
- [12] Lorenzen, T.J. & Vance, L.C. (1986). The economic design of control charts: a unified approach, *Technometrics* **28**, 3–10.
- [13] Woodall, W.H. (1986). The design of CUSUM quality control charts, *Journal of Quality Technology* **18**, 99–102.
- [14] Woodall, W.H. (1986). Weaknesses of the economic design of control charts (letter to the editor), *Technometrics* **28**, 408–410.
- [15] Woodall, W.H. (1986). Conflicts between Deming's philosophy and the economic design of control charts, *Frontiers in Statistical Quality Control* **3**, 242–248.
- [16] Saniga, E.M. (1989). Economic statistical design of control charts with an application to $\bar{X}$ and R charts, *Technometrics* **31**, 313–320.
- [17] Deming, W.E. (1986). *Out of the Crisis*, MIT Press & Center for Advanced Engineering Study, Cambridge.
- [18] Saniga, E. & Montgomery, D. (1981). Economical quality control policies for a single cause system, *AIIE Transactions* **13**(3), 258–264.
- [19] Moustakides, G.V. (1986). Optimal stopping times for detecting changes in distributions, *Annals of Statistics* **14**, 1379–1387.
- [20] Lucas, J. & Saccucci, M. (1990). Exponentially weighted moving average control schemes: properties and enhancements (with discussion), *Technometrics* **32**, 1–29.
- [21] Saniga, E., McWilliams, T., Davis, D. & Lucas, J. (2006). Economic control chart policies for monitoring variables, *International Journal of Productivity and Quality Management* **1**(1), 116–138.
- [22] Saniga, E., McWilliams, T., Davis, D. & Lucas, J. (2006). Economic advantages of CUSUM control charts for variables, in *Frontiers in Statistical Quality Control*, H.-J. Lenz, & P.-T. Wilrich, eds, Physica-Verlag, Vol. 8.
- [23] Jones, S. (1996). In *Cost Accounting: Marks Standard Handbook for Mechanical Engineers*, E.A. Avallone, & T. Baumeister III, eds, McGraw-Hill, 17–11–17–18.

- [24] Chiu, W.K. & Wetherill, G.B. (1974). A simplified scheme for the design of $\bar{X}$ charts, *Journal of Quality Technology* **6**, 63–69.
- [25] Von Collani, E. (1989). *The Economic Design of Control Charts*, B.G. Teubner, Stuttgart.
- [26] Montgomery, D.C. (1982). Economic design of an $\bar{X}$ control chart, *Journal of Quality Technology* **14**, 40–43.
- [27] Saniga, E., Davis, D. & McWilliams, T. (1995). Economic, statistical and economic statistical design of attribute control charts, *Journal of Quality Technology* **27**(1), 56–73.
- [28] McWilliams, T., Saniga, E. & Davis, D. (2001). Economic, statistical and economic statistical design of $\bar{X}$ and R charts and $\bar{X}$ and s charts, *Journal of Quality Technology* **33**, 234–241.

Related Articles

Average Run Lengths and Operating Characteristic Curves; Control Charts for Attributes; Control Charts, Selection of; Cost of Quality; Cumulative Sum (CUSUM) Chart; Six Sigma; Total Quality Management (TQM); Variable Sampling Rate Control Charts; Variance Components.

ERWIN M. SANIGA

Average Run Lengths and Operating Characteristic Curves

Introduction

The average run length (*ARL*) is the expected number of rational subgroups to be collected or the expected number of charting statistics to be plotted on a **control chart** before the first signal is given by the chart. The operating characteristic (*OC*) is the probability that no signal will be given by a chart on the first subsequent sample following a shift and thus failing to detect the shift. Both the *ARL* and the *OC* are measures of how quickly a chart will detect a change in the process and can be used to assess how well the chart performs.

The performance of a control chart, however, is tied to how it is designed. The design of a control chart involves, among other things, the choice of a charting statistic, the choice of control limits (whether it is $k\sigma$ control limits or probability limits), the size of the shift to be detected, and the sample or rational subgroup size as well as the frequency of sampling. These issues need to be addressed before a chart can be implemented.

Operating Characteristic

To illustrate how the *OC* can be of help with design ideas, consider the Shewhart $\bar{X}$ control chart (*see Control Charts for the Mean*) with $k\sigma$ limits for monitoring the process mean of a quality characteristic. Assume the quality characteristic follows a **normal distribution** with both the process mean and the process standard deviation known constants, denoted by μ_0 and σ_0 , respectively. Consider the probability of no signal on the $\bar{X}$ control chart when a shift in the mean of the quality characteristic has occurred. Suppose that the shift in the mean is expressed as $\mu_1 = \mu_0 + \delta\sigma_0$, where μ_1 denotes the new mean after the shift and δ is the size of the shift expressed in standard deviation units. This type of a shift is known as a *sustained shift* (or a step shift) in the process mean. This implies that the new process mean is $\delta\sigma_0$ units different (higher or lower) from the old process mean

for all future subgroups drawn from the process. The probability of no signal on the $\bar{X}$ control chart, given a sustained shift in the process mean, is calculated as

$$\begin{aligned}
 &P(\text{No signal/Shift in process mean}) \\
 &= P\left(\mu_0 - k\frac{\sigma_0}{\sqrt{n}} < \bar{X}_i < \mu_0 + k\frac{\sigma_0}{\sqrt{n}} \mid \mu_1 = \mu_0 + \delta\sigma_0\right) \\
 &= P(-k - \delta\sqrt{n} < Z_i < k - \delta\sqrt{n}) \\
 &= \Phi(k - \delta\sqrt{n}) - \Phi(-k - \delta\sqrt{n}) \\
 &= \beta(k, \delta, n), \text{ say.} \tag{1}
 \end{aligned}$$

This last probability is the *OC* (or β -risk as it is sometimes called) and is a function of three things: the distance of the control limits from the centerline k , the size of the shift δ , and the sample size n . The *OC* is therefore a useful tool to understand a chart's ability to detect a shift δ and simultaneously study the influence of k and n . In the **quality control** environment, however, the sample size n and charting constant k are typically fixed or prespecified so that it is common to use the plot of $\beta(k, \delta, n)$ as a function of δ . This plot is called the *operating characteristic curve* or the *OC curve* and lets us see what effect the process shift δ has on the chart's performance.

Example 1: Assume that the traditional 3σ control limits (or alternatively, 0.00135 probability limits) are used and detecting a shift of size 1.5 is of interest with samples of size $n = 5$. If the process is in control, $\delta = 0$ and the probability of no signal is found by substituting $k = 3$, $\delta = 0$, and $n = 5$ in $\beta(k, \delta, n)$, which yields $\beta(3, 0, 5) = 0.9973$. When a process shift of size 1.5 occurs, the probability of no signal is found by substituting $k = 3$, $\delta = 1.5$, and $n = 5$ in $\beta(k, \delta, n)$. This yields 0.362.

Figure 1 illustrates the *OC* curve where 3σ control limits and samples of size 5 are used. Thus, we are considering the value of $\beta(3, \delta, 5) = \Phi(3 - \delta\sqrt{5}) - \Phi(-3 - \delta\sqrt{5})$ as δ changes. This allows us to observe the influence of the size of the shift on the probability of no signal over a range of values. Note that there is a second horizontal axis showing the process mean for easier interpretation.

On the *OC* curve the β -risk is plotted on the vertical axis for a given δ on the lower horizontal

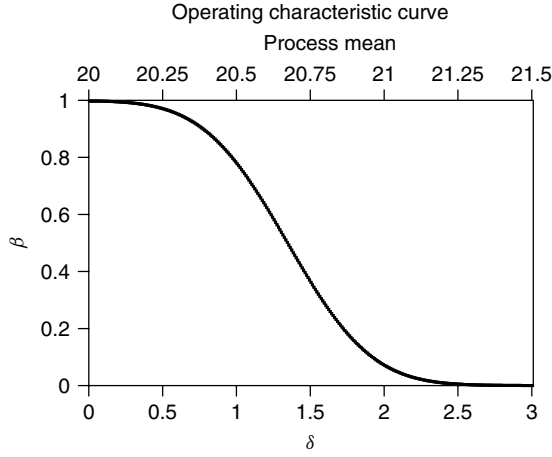


Figure 1 OC curve of an $\bar{X}$ control chart when $k = 3$ and $n = 5$

axis. Because the β -risk is, in fact, a probability and since 3σ control limits are used, the value for the β -risk varies between 0 (when $\delta \rightarrow \infty$) and 0.9973 (when $\delta = 0$). In addition, the upper horizontal axis shows the process mean μ_1 . The relationship between the lower and the upper horizontal axis is linear and is given by $\mu_1 = \mu_0 + \delta\sigma_0 = 20 + \delta 0.5$. However, we mainly focus on the relationship between the β -risk and δ (or the size of the shift in the process mean) and not on the process mean itself.

We can see that as the process mean moves further away from the known in-control value of 20 (when $\delta = 0$) the β -risk decreases. Thus, the larger the sustained shift in the process mean, the smaller the probability that the sample mean (or charting statistic) will plot between the upper control limit and the lower control limit, or alternatively, give no signal. This is a good thing.

As an alternative to the OC curve, we can graph the probability of a signal $1 - \beta(k, \delta, n)$ given that a sustained shift in the process mean has occurred. When substituting $\delta = 0$ in $1 - \beta(k, \delta, n)$ we obtain the false alarm rate (FAR), i.e. the probability of a signal when no shift occurred. In this case, $1 - \beta(3, 0, 5) = 0.0027$. If, however, a process shift of size 1.5 occurs the probability of a signal is $1 - \beta(3, 1.5, 5) = 1 - 0.362 = 0.638$. Therefore, if the process mean has either increased or decreased by 1.5 standard deviations the probability of a signal is no longer 0.0027 but 0.638.

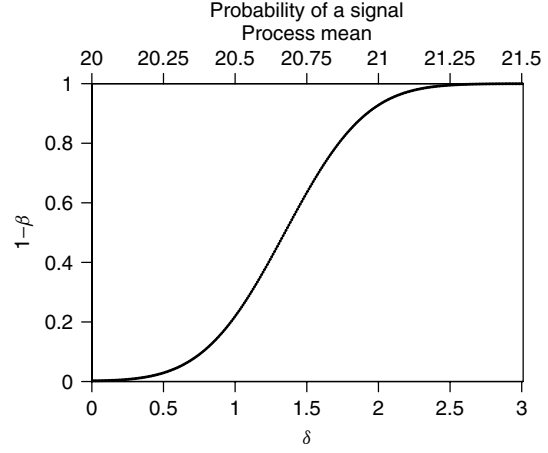


Figure 2 Probability of a signal of an $\bar{X}$ control chart when $k = 3$ and $n = 5$

Figure 2 indicates how the probability of a signal increases as the size of the sustained shift increases. Intuitively an increase in the probability of a signal (or on the other hand, a decrease in the probability for no signal) was expected.

It is interesting to note the similarity between the probability of a signal on an $\bar{X}$ control chart, i.e. $1 - \beta(k, \delta, n)$ as a function of δ and the **power** $\pi(\delta)$ of a test for the null hypothesis $H_0: \mu = \mu_0$ against the alternative hypothesis $H_1: \mu \neq \mu_0$. The test is based on the (test) statistic $Z = \frac{\sqrt{n}(\bar{X} - \mu_0)}{\sigma_0}$ and we reject H_0 in favor of H_1 if $|Z| \geq k$. Now if we assume that $\mu = \mu_0 + \delta\sigma_0$ and calculate the power of this test as a function of δ , we find

$$\begin{aligned} \pi(\delta) &= 1 - \beta = 1 - P(|Z| < k | \mu = \mu_0 + \delta\sigma_0) \\ &= 1 - \Phi(k - \delta\sqrt{n}) + \Phi(-k - \delta\sqrt{n}) \end{aligned} \quad (2)$$

which is of course the probability of a signal $1 - \beta(k, \delta, n)$ obtained earlier.

The OC curve (or, equivalently, the probability of a signal) can also be used as a tool to compare control charting plans. Suppose that an operator would like to compare two control charting plans to monitor the process mean. Both plans use the $\bar{X}$ control chart with 3σ control limits but the first plan uses $n = 5$ taken every 30 min, whereas the second plan uses $n = 6$ taken every 30 min. The question is then what the effect of sampling an extra unit is. The OC curves for both these plans are displayed in Figure 3 with

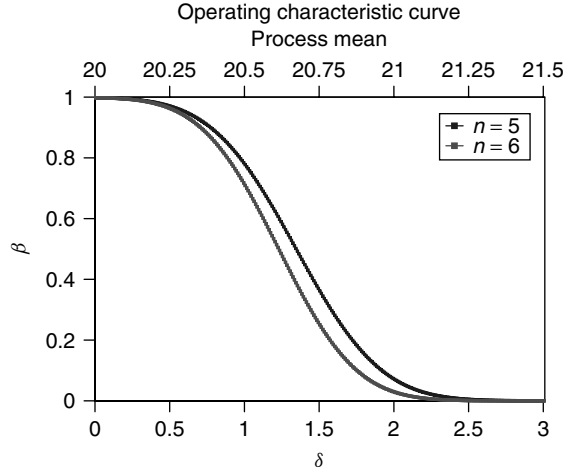


Figure 3 OC curves for (a) plan 1 with $n = 5$ and (b) plan 2 with $n = 6$ for varying δ

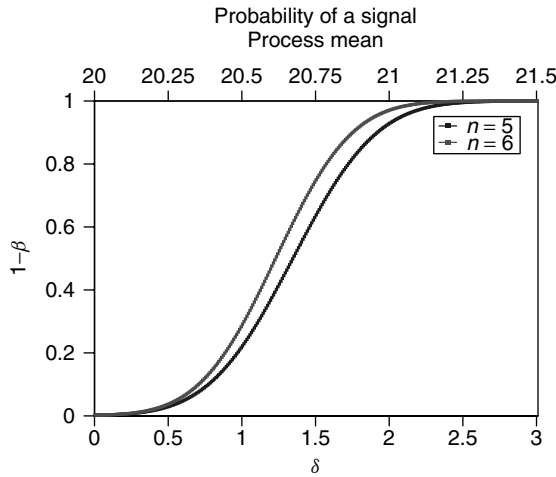


Figure 4 The probability of a signal for (a) plan 1 with $n = 5$ and (b) plan 2 with $n = 6$ for varying δ

the accompanying probabilities of a signal (or $1 - \beta(k, \delta, n)$) shown in Figure 4. In addition, Table 1 lists some values of $\beta(k, \delta, n)$ and $1 - \beta(k, \delta, n)$ for various values of δ .

We see that the second plan has a consistently lower β -risk and a higher probability of a signal for all values of δ . Therefore, if the objective is to detect a shift in the process mean as soon as possible and one can afford the extra cost of sampling, plan 2 will be preferred. In the language of **hypothesis testing**,

this shows that with all other things being equal, the power of a test to detect a shift is higher for a larger sample size.

Because the OC curve and the probability of a signal, given a shift in the process, do not tell the entire story, perhaps a more popular method to evaluate and examine the performance of control charts or make comparisons between control chart plans, is the run-length distribution.

Run-Length Distribution and the Average Run Length

The number of rational subgroups to be collected or the number of charting statistics to be plotted on a control chart before the first out-of-control signal is observed, is the run length of a chart. The discrete random variable defining the run length is called the *run-length random variable* and typically denoted N . The distribution of N is called the *run-length distribution*.

Characteristics of this distribution give the user more insight into the performance of a chart. These characteristics are, for example, its **moments** such as the expected value and standard deviation as well as **percentiles** or **quartiles**.

If no shift occurred ($\delta = 0$) the distribution of N is referred to as the *in-control run-length distribution*. In contrast, if the process did encounter a shift ($\delta \neq 0$) the distribution of N is labeled the out-of-control run-length distribution. To distinguish between these two situations the notations N_0 and N_δ are used. This notation is also used for the characteristics of the run-length distribution.

Assuming that the random samples are independent and that the probability of a signal is the same for all samples (or time points), the run-length distribution is given by

$$P(N = j) = \beta(k, \delta, n)^{j-1} (1 - \beta(k, \delta, n)),$$

$$j = 1, 2, 3, \dots \quad (3)$$

where $\beta(k, \delta, n)$ is the probability of no signal defined in equation (1). This distribution is easily recognized as the geometric distribution with probability of success $1 - \beta(k, \delta, n)$ so that we write, symbolically, $N \sim \text{Geo}(1 - \beta(k, \delta, n))$. The success probability is the probability of a signal or the signal probability and this probability completely characterizes the geometric (run length) distribution.

Table 1 The probability of no signal and the probability of a signal for two competing control chart plans

δ	Probability of no signal		Probability of a signal	
	$\beta(3, \delta, 5)$	$\beta(3, \delta, 6)$	$1 - \beta(3, \delta, 5)$	$1 - \beta(3, \delta, 6)$
0	0.9973	0.9973	0.0027	0.0027
0.25	0.9925	0.9914	0.0075	0.0086
0.50	0.9701	0.9621	0.0299	0.0379
0.75	0.9071	0.8776	0.0929	0.1224
1.00	0.7775	0.7090	0.2225	0.2910
1.25	0.5812	0.4753	0.4188	0.5247
1.50	0.3616	0.2501	0.6384	0.7499
2.00	0.0705	0.0288	0.9295	0.9712
2.25	0.0211	0.0060	0.9789	0.9940
2.50	0.0048	0.0009	0.9952	0.9991
2.75	0.0008	0.0001	0.9992	0.9999
3.00	0.0001	0.0000	0.9999	1.0000

As mentioned, various statistical characteristics or summary measures of the run-length distribution provide insight into how a control chart functions and performs. Remember that we want the chart to signal a shift or a change quickly when a change takes place and not signal too often when the process is actually in control, which is when no shift or change has taken place. We are interested in the typical value as well as the spread or the variation in the run-length distribution. A popular measure of the typical or the central tendency of a distribution is the expected value or the mean or the average. Accordingly, this has been the most popular index or measure of a typical chart performance, called the *average run length (ARL)*. Therefore, as stated before, the *ARL* of a control chart is the expected number of subgroups that must be collected before the control chart signals for the first time. In other words the *ARL* is the expected waiting time for the first signal. When the process is actually in control, the expected number of subgroups that must be collected before the control chart signals (erroneously) for the first time is called the *in-control average run length* and is denoted by ARL_0 . On the other hand, the out-of-control average run length, denoted by ARL_δ , is the expected number of samples to be collected before a control chart signals for the first time, when the process has gone out of control. Obviously, for an “efficient” control chart the in-control average run length ARL_0 should be “large” and the out-of-control average run length ARL_δ should be “small”. Hence, when comparing the performance of two

or more charts the in-control average run length ARL_0 of the charts is usually fixed at an acceptably “high” level so that the number of false alarms or the *FAR* is reasonably “small”. The chart with the smallest or lowest out-of-control average run length ARL_δ for a change or shift of a specified size δ in the process parameter is then selected to be the winner.

From the properties of the geometric distribution the *ARL* for the chart is the expected value of N , therefore $ARL = E(N) = [1 - \beta(k, \delta, n)]^{-1}$.

Therefore, when the signaling events are independent and have the same probability, the *ARL* of the chart is simply the reciprocal of the probability of a signal. It is this simple relationship between the *ARL* and the probability of a signal (or, probability of no signal) that accounts for their popularity as measures of a control charts’ performance.

It is seen that the *ARL* is also a function of k , δ , and n and for the Shewhart $\bar{X}$ chart,

$$ARL(k, \delta, n) = [1 - \Phi(k - \delta\sqrt{n}) + \Phi(-k - \delta\sqrt{n})]^{-1} \quad (4)$$

Consequently, the *ARL* can and often is also used to compare two or more control chart schemes or plans.

Example 2: Suppose we have to choose between the two competing control chart plans from Example 1 using the *ARL* as the performance measure. The *ARL* (for various δ) of both plans are shown in Figure 5 with the *ARL* on the vertical axis and

the size of the shift δ on the lower horizontal axis together with the actual process mean on the upper horizontal axis. The most efficient control chart plan is clearly the one for which the ARL is the smallest, given a particular shift in the process mean.

From Table 2, for instance, we see that if $\delta = 0.5$ then plan 2, with $ARL(3, 0.5, 6) = 26.36$, which uses $n = 6$ observations at each sampling stage, is preferred to plan 1 with $ARL(3, 0.5, 5) = 33.40$. On the other hand, from Table 2 it is seen that if $\delta \geq 2.5$

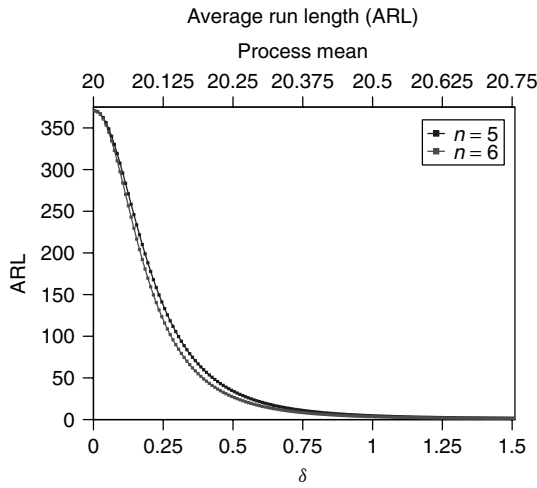


Figure 5 ARL for (a) plan 1 with $n = 5$ and (b) plan 2 with $n = 6$ for varying δ

Table 2 The average run lengths for two competing control chart plans

δ	Plan 1 $ARL(3, \delta, 5)$	Plan 2 $ARL(3, \delta, 6)$
0	370.38	370.38
0.25	133.16	115.87
0.50	33.40	26.36
0.75	10.76	8.17
1.00	4.50	3.44
1.25	2.39	1.91
1.50	1.57	1.33
1.75	1.22	1.11
2.00	1.08	1.03
2.25	1.02	1.01
2.50	1.00	1.00
2.75	1.00	1.00
3.00	1.00	1.00

then either plan 1 or plan 2 may be used since their ARL s are basically equal, i.e. $ARL(3, \delta, 5) \approx ARL(3, \delta, 6) \approx 1$ for $\delta \geq 2.5$.

Note that in order to compare two control chart plans, it is customary to fix the in-control average run length (ARL_0) at some specified or desirable value – in this case 370.38. This was also done in Example 1 when we used $\beta(k, \delta, n)$ and $1 - \beta(k, \delta, n)$ – see Table 1. Although the OC and the ARL enable us to compare different control chart plans, using this approach has a drawback – there is a trade-off between the low occurrence of false alarms and the fast detection of a process change. But, because the OC and the ARL can be used to compare different types of control chart designs (such as the Shewhart, cumulative sum (CUSUM) (see **Cumulative Sum (CUSUM) Chart**), exponentially weighted moving average (EWMA) (see **Exponentially Weighted Moving Average (EWMA) Control Chart**) etc.), which is valuable, this drawback is almost entirely overshadowed.

Other characteristics of the run-length distribution are also of interest. For example, in addition to the mean, one should also look at the standard deviation of the run-length distribution to get an idea about the variation. Using results for the geometric distribution, the standard deviation of the run length, denoted by $SDRL$, is given by

$$SDRL(k, \delta, n) = \sqrt{\beta(k, \delta, n) [1 - \beta(k, \delta, n)]}^{-1} \quad (5)$$

Since the geometric distribution is skewed (to the right) the mean and the standard deviation become questionable measures of central tendency and spread. Other descriptive measures can be more useful. For example, the percentiles such as the median and the quartiles can provide valuable information both about the typical as well as the variation in the run-length distribution. Since the run-length distribution is discrete, the $100^{\text{th}} p$ th percentile ($0 < p < 1$) is defined as the smallest integer j such that the cumulative probability is at least p , that is $P(N \leq j) \geq p$. The median is the 50th percentile so that $p = 0.5$ and the first quartile is the 25th percentile, so that $p = 0.25$. For example, we calculate the following cumulative probabilities for the run-length (geometric) distribution for a signaling probability of 0.2 (Table 3).

Hence, the first quartile of the run-length distribution is 2, the median is $-\log(2)[\log(0.8)]^{-1} = 3.1$,

Table 3 Cumulative probabilities of the run-length distribution for signal probability 0.2

j	1	2	3	4	5	6	7	8	9
$P(N \leq j)$	0.2000	0.3600	0.4880	0.5904	0.6723	0.7379	0.7903	0.8322	0.8658

i.e. 4 and, the third quartile is 7. For the same distribution, the ARL is $0.2^{-1} = 5$ and the standard deviation is $\sqrt{0.8 \times 0.2^{-1}} = 4.47$. The interquartile range equals $7 - 2 = 5$ which is perhaps a better indicator of the spread than the standard deviation.

In the case of other Shewhart-type control charts such as the S^2 , S , R , p , np , c , and u charts the run-length distribution is also geometric with the probability of success equal to the probability of a signal, i.e. $N \sim \text{Geo}(P(\text{Signal}))$. Therefore, as with the $\bar{X}$ we can again use the OC curve and the ARL as a function of the process shift to measure the performance of these control charts. All we need is basically the probability of a signal or, alternatively the probability of a signal. One should note that computing the OC values or ARL for monitoring schemes that combine information over time (e.g. CUSUM and EWMA control charts) or for charts monitoring autocorrelated data can be problematic. In such cases, simulation has proven useful in approximating run-length distributions.

Other Performance Measures

Two other measures of performance of a control chart are the average time to signal (ATS) and the number of individual items that have been produced or inspected I .

If samples are taken over equally spaced time intervals the ATS is given by $ATS = ARL \cdot h$ where h is the time between any two successive samples. Thus, the ATS is the average number of time periods until a signal is generated on the control chart.

The ARL can also be expressed as the number of individual items that have been inspected I , rather than the number of samples taken. If the samples are all of size n , I is given by $I = ARL \cdot n$.

In practice, however, where control charts are often designed for a very specific process or system, evaluating the performance of a control chart can be more difficult and time consuming than what has been presented in the foregoing discussion. A more in-depth analysis of the performance of Shewhart control charts is then typically needed.

Acceptance Sampling and OC Curves

In **acceptance sampling**, for instance, the OC curve is not similar to that of a Shewhart-type of control chart. Instead, for any given fraction nonconforming p in a submitted lot, the OC curve shows the probability P_a that the lot will be accepted on the vertical axis for a given fraction nonconforming p on the horizontal axis.

The question that acceptance sampling tries to answer is also different i.e. given a submitted lot of products, how can we tell if it is acceptable? There are two options when assessing the quality of lots – 100% inspection or acceptance sampling. The purpose of acceptance sampling is to determine whether a submitted lot is likely to contain “acceptable” or “rejected” products based on a random sample only. This involves taking one or more samples from a lot and is generally cheaper than 100% inspection – especially when **destructive testing** is needed. Each item in the sample(s) is then measured or otherwise evaluated and ultimately classified as “accepted” or “rejected”.

The simplest type of sampling plan involves a single sample. However, more advanced procedures, for example, double sampling, multiple sampling, and **sequential sampling**, have also been developed. In judging the diverse choice of sampling plans it is desirable to compare their performance over various quality levels. OC curves serve this purpose.

Ideally, we want a sampling plan that would make the correct decision between good and bad lots 100% of the time. This ideal OC curve will therefore always accept the lot when the actual fraction nonconforming is at or below the acceptable quality level (AQL) and it will always reject the lot when the actual fraction nonconforming is above the AQL . The probability of acceptance will equal one until the AQL is reached. At this point the curve drops vertically from a probability of acceptance of one to a probability of acceptance of zero. Hereafter, the probability of acceptance will equal zero. The ideal OC curve is illustrated in Figure 6.

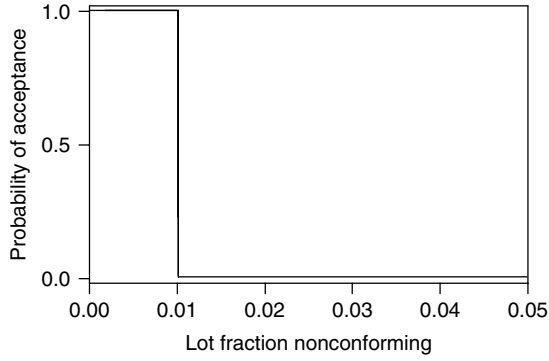


Figure 6 Ideal OC curve with $AQL = 0.01$

Sadly, mistakes will sometimes be made, i.e. bad lots will be accepted or good lots will be rejected and the ideal OC curve can therefore never (or, at most, rarely) be obtained in practice.

OC Curves for Single (Attributes) Sampling Plans

When only one sample is used to declare the acceptability of the lot, the acceptance plan is called a *single-sampling plan* (see **Single Sampling by Attributes and by Variables**). Suppose we needed to inspect a lot of size N which is large (theoretically infinite). A random sample of n items is taken and the number of nonconforming items d is observed. The lot will be accepted if d is less than or equal to an acceptance number c , i.e. the maximum allowable number of nonconforming items in the sample. The probability of lot acceptance is then the probability that d is less than or equal to c , i.e. $P_a = P(d \leq c)$. Because the number of nonconforming items d in a random sample of size n is binomially distributed with parameters n and p , where p denotes the fraction of nonconforming items in the lot, we have that $P_a = P(d \leq c) = \sum_{d=0}^c \binom{n}{d} p^d (1-p)^{n-d}$. The OC curve is obtained if P_a is plotted against p .

Example 3: The probability P_a of no more than $c = 2$ nonconforming items in a sample of $n = 100$, with $p = 0.01$ is $P_a = P(d \leq 2) = \sum_{d=0}^2 \binom{100}{d} 0.01^d \times (1 - 0.01)^{100-d} = 0.92$. This means that if 1% of

the products are nonconforming, the probability of lot acceptance is 0.92. Alternatively, one can interpret this as follows: If 100 lots are submitted from a process that produces 1% nonconforming items it is expected that 92 of the lots will be accepted and 8 will be rejected.

For different choices of n and c the acceptance probability P_a can be calculated similarly, but note that the shapes of the OC curves depend entirely on the choice of n and c . This is shown in Figure 7.

The OC curve basically shows the discriminatory power of the sampling plan. For example, when lots with less than 1% nonconforming items are submitted, they will rarely be rejected, since the probability of acceptance is greater than 0.9 for all plans. However, when lots with more than 6% nonconforming items are submitted, they will almost surely be rejected, since the probability of acceptance is small ($P_a < 0.2$) for all plans. From the slope of the OC curve it is seen that the OC curve approaches the shape of the ideal OC curve as the sample size increases. This implies that the discriminating power of the OC curve increases as sample size increases.

Acceptance sampling plans with $c = 0$, however, are not desirable because the probability of acceptance drops much quicker than for $c > 0$ leading to “easier” rejection of a lot. This is illustrated in Figure 8.

Type-A and Type-B OC Curves

The above-mentioned OC curves are called *type-B*

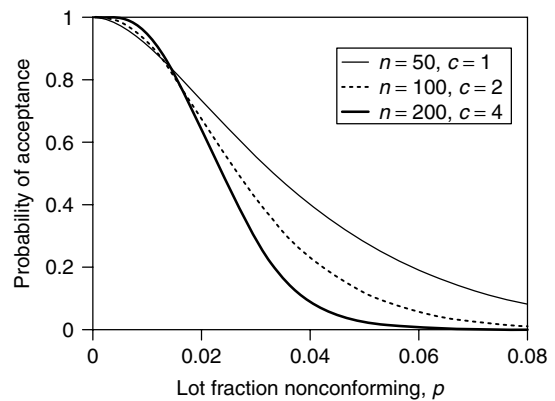


Figure 7 OC curves of a single-sampling plan^a

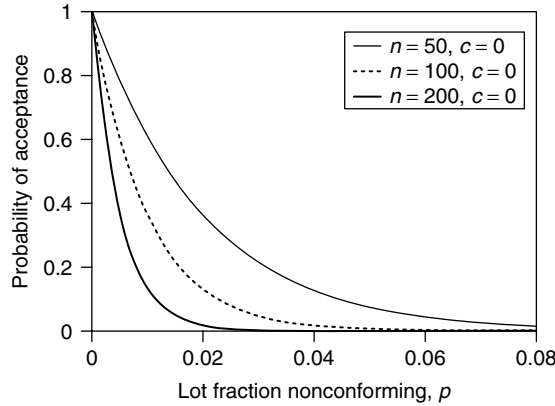


Figure 8 OC curves of a single-sampling plan for different sample sizes and $c = 0$

OC curves. In the structure of these OC curves the binomial distribution is the exact probability distribution for calculating the probability of lot acceptance (the Poisson distribution gives a satisfactory approximation). In contrast, in the structure of type-A OC curves the hypergeometric distribution is used for calculating the probability of lot acceptance (the Poisson or binomial distributions often give satisfactory approximations). While type-A OC curves are concerned with the quality of isolated lots (of finite size), type-B OC curves are concerned with the average quality of a stream of lots (theoretically of infinite size). Type-B OC curves can be used to examine a producer's risk, since a producer typically wants to know what percentage of his lots will be rejected by a sampling plan for any quality level. In contrast, type-A OC curves can be used to examine consumer's risk, since a consumer typically wants to know the risk he is taking when accepting a somewhat bad lot.

It is seen from Figure 9 that as the lot size increases, the type-A OC curve approaches the type-B OC curve. So, there is no practical difference between a type-A OC curve with a large lot size and a type-B OC curve (or, the OC curve for infinite lots, as it is sometimes called). In general, if lot size is at least 10 times the sample size ($n/N \leq 0.1$), "a type-B OC curve". Hence, add OC may be taken as a good approximation to a type-A OC curve. Also shown in Figure 9 are type-A OC curves for $n = 20$, $c = 0$, and various lot sizes N . Note that the OC curves for $N = 500$, $N = 100$, and $N = 50$ are very similar. We can conclude that the lot size has a decreasing effect

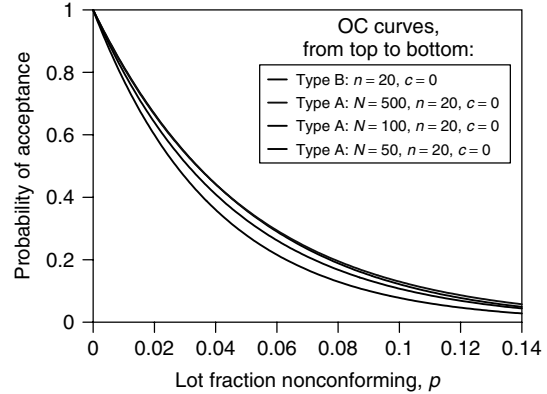


Figure 9 Comparisons of OC curves for different sampling plans

on the OC curve as it increases.

For a given lot fraction nonconforming, the probability of acceptance calculated for a finite lot (by using hypergeometric probabilities) is always less than that calculated for an infinite lot (by using binomial probabilities). Therefore, the type-A curve will always lie below the type-B curve (this is also illustrated in Figure 9).

The acceptance probabilities of the type-A OC curves are calculated using an approximation. As illustration, the probability P_a of no more than $c = 0$ nonconforming items in a sample of $n = 20$ taken from a lot of size $N = 50$ with $p = 0.04$ is

$$\begin{aligned}
 P_a &= P(d \leq 0) \\
 &= \sum_{d=0}^0 \frac{pN!}{d!(pN-d)!} \left(\frac{n}{N}\right)^d \left(1 - \frac{n}{N}\right)^{pN-d} \\
 &= \frac{(0.04)(50)!}{0![(0.04)(50)-0]!} \left(\frac{20}{50}\right)^0 \left(1 - \frac{20}{50}\right)^{(0.04)(50)-0} \\
 &= 0.36
 \end{aligned} \tag{6}$$

This means that if 4% of the products are nonconforming the probability of lot acceptance is 0.36.

OC Curves for Double, Multiple, and Sequential Sampling Plans for Attributes

When two samples are used to declare the acceptability of the lot, the acceptance plan is called a *double-sampling plan* (see **Multiple Sampling**

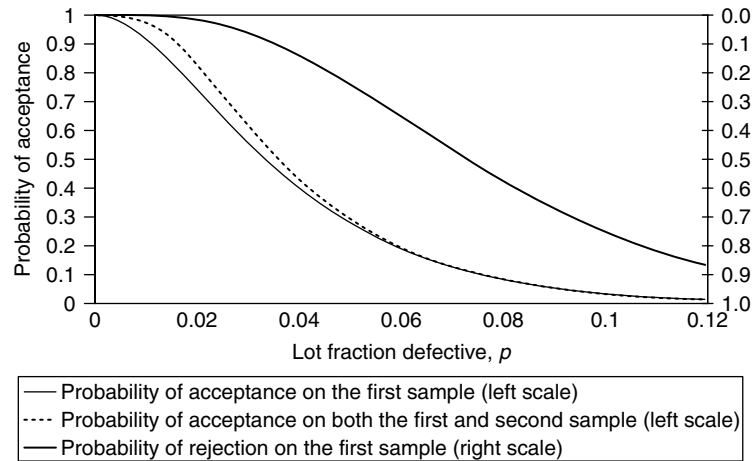


Figure 10 OC curves of a double-sampling plan with $n_1 = 50$, $n_2 = 100$, $c_1 = 1$, and $c_2 = 3$

Plans). Under certain circumstances, a second sample is required before a lot can be accepted or rejected. If the first sample is not good (bad) enough to be accepted (rejected), the decision is based on the first and second samples combined. Let n_1 and n_2 denote the sample sizes of the first and second sample, respectively. Let c_1 and c_2 denote the acceptance number of the first and second sample, respectively. The procedure of a double-sampling plan is as follows. The first sample of size n_1 is examined from a lot of size N . The lot is accepted if the sample contains c_1 or less nonconforming items. The lot is rejected if the sample contains more than c_2 nonconforming items. The second sample of size n_2 is examined if the first sample contained more than c_1 , but less than or equal to c_2 nonconforming items. The lot will be accepted if the combined sample contains c_2 or less nonconforming items; otherwise it will be rejected (Figure 10).

When three or more samples are used to declare the acceptability of the lot, the acceptance plan is called a *multiple-sampling plan*. The procedure of a multiple sampling plan follows. If the number of nonconforming items is less than or equal to some acceptance number, the lot is accepted. However, if the number of nonconforming items is greater than or equal to some rejection number, the lot is rejected. Otherwise, another sample is taken until the final sample is taken where a decision about the acceptability of the lot must be made. When the

acceptance plan allows for a decision each time a new item is inspected, it is called a *sequential sampling plan*.

Acknowledgments

The authors would like to thank STATOMET and the Department of Statistics at the University of Pretoria for making this article possible.

End Note

^aThe acceptance number is kept proportional to the sample size.

Further Reading

- Cowden, D.J. (1957). *Statistical Methods in Quality Control*, Prentice Hall, New York.
- Derman, C. & Ross, S.M. (1997). *Statistical Aspects of Quality Control*, Academic Press, San Diego.
- Duncan, A.J. (1986). *Quality Control and Industrial Statistics*, 5th Edition. Irwin.
- Grant, E.L. & Leavenworth, R.S. (1980). *Statistical Quality Control*, 5th Edition, McGraw-Hill.
- Montgomery, D.C. (2005). *Introduction to Statistical Quality Control*, 5th Edition, John Wiley & Sons, New York.

SCHALK W. HUMAN AND MARIEN A.
GRAHAM

Profile Monitoring

Introduction

Profile monitoring is the use of **control charts** for cases in which the quality of a process or product can be characterized by a functional relationship between a response variable and one or more explanatory variables. These cases appear to be increasingly common in practical applications. For each profile we assume that n ($n > 1$) values of the response variable (Y) are measured along with the corresponding values of one or more explanatory variables (the X s). The reader is referred to Woodall *et al.* [1] for a more detailed and comprehensive review of profile monitoring methods. This is a relatively new area of application, so no software is known to be available for its implementation.

Jin and Shi [2] and Zhou *et al.* [3] used the term *waveform signals* to refer to what we call *profiles*. These signals are often collected by sensors during manufacturing processes, e.g., tonnage stamping in stamping, torque signals in tapping, and force signals in welding. Gardner *et al.* [4] used the term *signature* in place of our use of *profile*.

Calibration processes are often characterized by linear functions, as discussed by Mestek *et al.* [5], for example. The profile monitoring framework also includes applications in which numerous measurements of the same variable, e.g., a dimension such as thickness, are made at several locations on each manufactured part.

In many cases it is most efficient to summarize the in-control performance with a parametric model and monitor for shifts in the parameters of this model. The control charts are then based on the estimated parameters of the model from successive profiles observed over time. For nonparametric models one can alternatively monitor metrics that reflect the discrepancies between observed profiles and a baseline profile established using historical data.

Some General Issues

Phase I Applications

In phase I, one analyzes a historical set of process data. The goals in phase I are to understand

the variation in a process over time, to evaluate the process stability, and to model the in-control process performance. With profile data collected for phase I, one should examine the fit of the hypothesized model for each profile. One should check, for example, for **outliers**. In some cases it may be more appropriate to delete specific points within a profile dataset rather than to discard the entire profile sample.

The principal component analysis, as described by Jones and Rice [6], can be very useful in summarizing and interpreting a set of phase I profile data with identical, equally spaced values of an independent variable X for each profile. With this approach one treats each profile as a multivariate vector of n response Y values and identifies a few mutually orthogonal linear combinations of the Y variables that explain as much variation in the profiles as possible. If these principal components are interpretable, this approach can be very effective in understanding process performance. Jones and Rice [6] recommended plotting the average profile and the profiles with the largest and smallest principal component scores for aiding in the interpretation of the principal components.

Model Selection

Any model selected for a particular application should be as simple as required to adequately model the profile data. Standard recommendations regarding model selection apply. With any model, however, it must be determined how to design a monitoring procedure to detect profile changes over time most effectively, preferably in a way that allows for the interpretation of out-of-control signals.

It is critically important to assess variation between profiles in phase I as well as variation within profiles. Common cause variation is that variation considered to be characteristic of the process and that cannot be reduced substantially without fundamental process changes. It must be determined how much of the profile-to-profile variation is common cause variation and should be incorporated into the determination of the control chart limits to be used for monitoring in phase II. There is common cause variation between profiles if it is not reasonable to expect each set of observed profile data to be adequately represented using the same set of parameters for the assumed model. Process knowledge is always

required in the decision whether or not to include profile-to-profile variation as common cause variation. From our experience, it is often unrealistic to expect that process improvement can be used to remove all profile-to-profile variation.

The Control Chart Statistic(s)

In phase II it is often recommended to monitor the profiles using a separate control chart for each parameter of a parametric model, provided the estimators of the parameters at each sampling stage are independent. If the estimators of the parameters of the model for the profile are dependent, as is more often the case, one can use a T^2 chart (see **Hotelling's T^2 Chart**). It is important to use an appropriate estimator of the variance–covariance estimator in phase I. In particular, if the vectors of estimators are pooled in order to estimate the variance–covariance matrix, then the resulting T^2 chart will tend to be quite ineffective in detecting outlying profiles or sustained shifts in the profiles over time.

If the approach used is nonparametric, then a reasonable approach is to base control charts on metrics that measure the departures of observed process profiles from a baseline profile model developed in phase I. One such metric is the average deviation of the observed profile from the baseline profile calculated over a grid of X values. Gardner *et al.* [4] recommended this approach.

One should make sure that the choice of statistics to monitor does not result in methods that are unable to detect certain types of profile shifts. For example, if one chooses to monitor only a subset of the principal components, or a subset of wavelet coefficients, that were determined using only in-control data, then certain types of out-of-control profile shifts may be undetectable.

In some cases it may be possible to target control charts to detect specified types of process faults. This can be done by inducing the process faults and observing the effect on several profiles. The control chart statistic can then be determined, perhaps using discriminant analysis, to be most effective in detecting a specific fault. Gardner *et al.* [4], for example, used this type of approach. This strategy reduces the dimensionality of the problem and increases the **power** of the charts, but one must be careful not to overlook important types of process faults.

Analysis of Simple Linear Regression Profiles

In this section we review the work done on the simplest case, the one in which the profile is adequately represented by a straight line.

Phase II Methods

Several researchers have studied phase II control charting methods for monitoring a simple linear regression relationship with assumed known values of the Y intercept, slope, and variance parameters. One approach of Kang and Albin [7] involved a multivariate T^2 control chart based on the successive vectors of the least-squares estimators of the Y intercept and slope. Their second method used statistics based on the successive samples of n deviations from the in-control line with a combination of an exponentially weighted moving average (EWMA) chart (see **Exponentially Weighted Moving Average (EWMA) Control Chart**) to monitor the average deviation and a range (R) chart (see **Moving Range and R Charts**) to monitor the variation of the deviations.

Kim *et al.* [8] proposed alternative control charts for monitoring in phase II. Since they recommended coding the X values so that the estimators of the Y intercept and slope are independent, they monitored each of the two regression coefficients using a separate EWMA chart. Also, they proposed replacing the R chart of Kang and Albin [7] by one of two EWMA charts for monitoring a process standard deviation. The statistical performance of the combined use of these three EWMA charts was shown to be superior to that of the phase II methods of Kang and Albin [7] for sustained shifts in the regression parameters. Their method is also much more interpretable since each parameter in the model is monitored with a separate control chart. One could use other types of control charts, such as individuals and **moving average** control charts (see **Control Charts for the Mean**), with the sequences of estimators of the slope and Y intercept.

The approach of Kim *et al.* [8] is illustrated in Figure 1 with Shewhart charts applied to the data shown in Figure 2. These charts show that the fitted line corresponding to sample four was significantly different from the standard line.

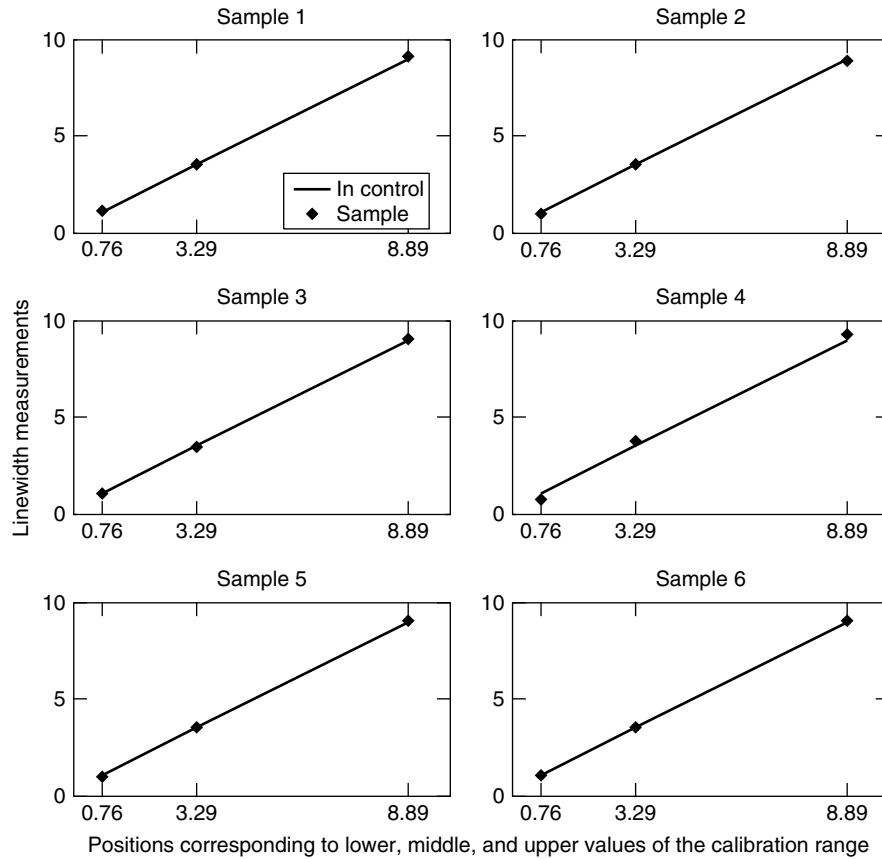


Figure 1 Plot of illustrative phase II calibration data with fitted lines [Reproduced with permission from Gupta *et al.* [9]]

Phase I Methods

Kim *et al.* [8] recommended the use of separate Shewhart charts for the Y intercept, slope, and the error variance in phase I applications provided one codes the X values so that the average X value is zero.

Mahmoud and Woodall [10] proposed a phase I method based on the standard approach of testing collinearity using indicator variables for each profile in a multiple regression model. The performance of several methods was compared for different numbers of shifts of a specified size in the regression parameters. The authors showed that some previously proposed phase I methods are ineffective in detecting shifts in the process parameters due to the ways in which the vectors of estimators were pooled to estimate the variance-covariance matrix. Mahmoud and Woodall [10] illustrated several of the phase I

methods using a calibration dataset given by Mestek *et al.* [5].

Change-point methods (*see Change-Point Methods*) can also be very useful in phase I if one suspects that instability can be modeled adequately by step shifts in the underlying parameter(s). Mahmoud *et al.* [11] developed such a change-point method for the simple linear profile application.

Linear Calibration Applications

As discussed by Kim *et al.* [8], once the parameters are established in a baseline linear calibration study, one would want to detect any change in the Y intercept or slope and any increase in the variation about the regression line since such shifts could correspond to greater inaccuracies in the measurement process. The effect of assignable causes in general, however,

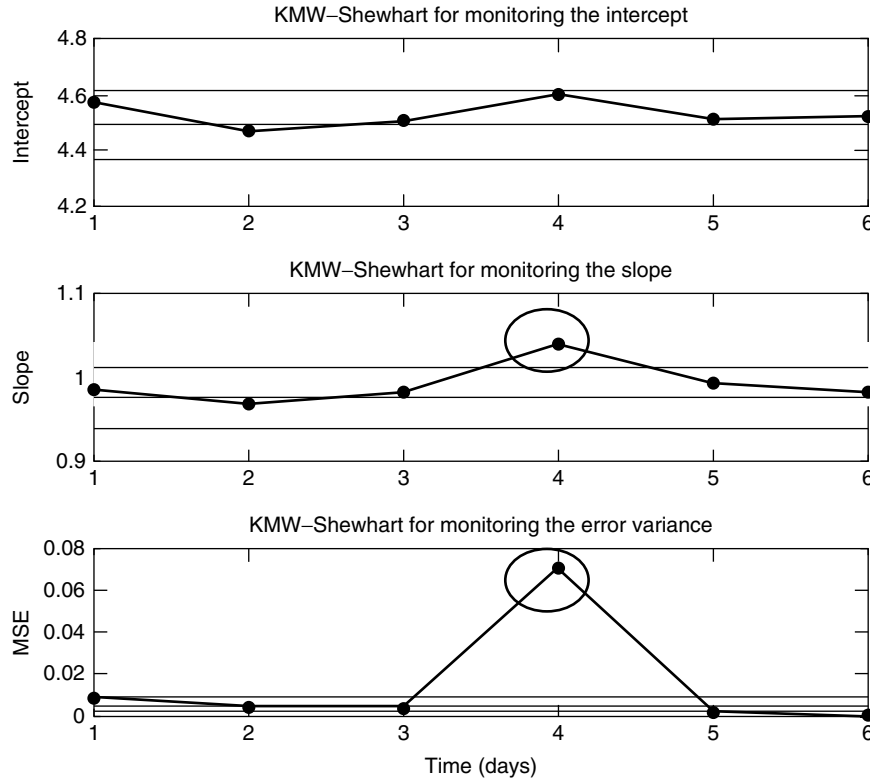


Figure 2 Shewhart control charts using the methods of Kim *et al.* [8] showing a significant deviation from the standard line at sample four [Reproduced with permission from Gupta *et al.* [9]]

will vary from application to application. Sometimes, for example, one may wish to detect isolated outliers.

Croarkin and Varner [12] and ISO 5725-6 [13] recommended a method for monitoring a linear calibration line based on the measurement of three standards at each time period. These are at low, medium, and high values, respectively. The three deviations of the measured values from these standards are plotted simultaneously for each sample on a Shewhart-type control chart. Gupta *et al.* [9] showed, however, that the use of control charts based on the estimated regression parameters is more effective. This latter approach is more interpretable, particularly as the number of standards measured at each time period increases. The NIST/SEMATECH *e-Handbook of Statistical Methods*, available online at <http://www.itl.nist.gov/div898/handbook/>, contains a discussion of the methods proposed by Croarkin and Varner [12].

Other Profile Models and Approaches

As a natural extension of the simple linear regression model, one can consider using control charts to monitor profiles that can be represented by a polynomial regression model or some other multiple linear regression relationship. Jensen *et al.* [14] considered phase II methods for this situation.

In some cases a nonlinear regression model is useful for modeling a profile. A nonlinear model can be used with the vertical board density profile data provided by Walker and Wright [15], where the density of particleboard measured across the diameter of the board is U shaped.

Young *et al.* [16] also considered a vertical density profile application. In their approach, however, the data for each profile were summarized into a top-face, bottom-face, and core average density. Then, a T^2 approach was used for the resulting three-dimensional quality vectors. In general it can be useful to break

the range of the X variable into intervals and to model the profile within each interval. Each “subprofile” could be monitored separately or a T^2 method could be used. This approach seems most promising when assignable causes tend to affect the profile only within certain intervals. It is closely related to the use of wavelet models based on the Haar transform.

Jin and Shi [17] used wavelets to model stamping tonnage signals, profiles which can have a rather complicated shape. They recommended Shewhart charts to monitor changes in wavelet coefficients. In some cases a wavelet coefficient can be associated with a specific process fault. Some additional work was reported by Jin and Shi [2], Lada *et al.* [18], and Zhou *et al.* [3].

Nonparametric, or smoothing, methods do not require a specified functional form for the profile. In nonparametric regression methods, one obtains a smoothed curve that can be expressed as a weighted average of the observed responses. Winistorfer *et al.* [19] used this approach to model vertical density profile data for pressed wood panels.

In a semiconductor application, Gardner *et al.* [4] modeled the spatial signatures (i.e., profiles) of wafers using splines and then used various metrics to compare the profiles of new wafers to an established baseline surface. The metrics used were designed to detect specific equipment faults.

Boeing [20, pp. 140–144] also proposed phase II control charting methods using spline fitting. In their application there were multiple gap measurements made at several locations on a manufactured part. Splines were fit to the responses for each part. The average spline value for the first k splines observed or nominal values are used as the baseline. Metrics such as the maximum difference or the average absolute difference was calculated for each spline relative to the baseline. Individuals X charts and moving range charts are then based on the values of the metrics.

References

- [1] Woodall, W.H., Spitzner, D.J., Montgomery, D.C. & Gupta, S. (2004). Using control charts to monitor process and product quality profiles, *Journal of Quality Technology* **36**, 309–320.
- [2] Jin, J. & Shi, J. (2001). Automatic feature extraction of waveform signals for in-process diagnostic performance improvement, *Journal of Intelligent Manufacturing* **12**, 257–268.
- [3] Zhou, S.Y., Sun, B.C. & Shi, J.J. (2006). An SPC monitoring system for cycle-based waveform signals using Haar transform, *IEEE Transactions on Automation Science and Engineering* **3**, 60–72.
- [4] Gardner, M.M., Lu, J.C., Gyurcsik, R.S., Wortman, J.J., Hornung, B.E., Heinisch, H.H., Rying, E.A., Rao, S., Davis, J.C. & Mozumder, P.K. (1997). Equipment fault detection using spatial signatures, *IEEE Transactions on Components, Packaging, and Manufacturing Technology – Part C* **20**, 295–304.
- [5] Mestek, O., Pavlik, J. & Suchánek, M. (1994). Multivariate control charts: control charts for calibration curves, *Fresenius' Journal of Analytical Chemistry* **350**, 344–351.
- [6] Jones, M.C. & Rice, J.A. (1992). Displaying the important features of large collections of similar curves, *American Statistician* **46**, 140–145.
- [7] Kang, L. & Albin, S.L. (2000). On-line monitoring when the process yields a linear profile, *Journal of Quality Technology* **32**, 418–426.
- [8] Kim, K., Mahmoud, M.A. & Woodall, W.H. (2003). On the monitoring of linear profiles, *Journal of Quality Technology* **35**, 317–328.
- [9] Gupta, S., Montgomery, D.C. & Woodall, W.H. (1996). Performance evaluation of two methods for online monitoring of linear calibration profiles, *International Journal of Production Research* **44**, 1927–1942.
- [10] Mahmoud, M.A. & Woodall, W.H. (2003). Phase I analysis of linear profiles with calibration applications, *Technometrics* **46**, 377–391.
- [11] Mahmoud, M.A., Parker, P.A., Woodall, W.H. & Hawkins, D.M. (2007). A change point method for linear profile data, *Quality & Reliability Engineering International* **23**, 247–268.
- [12] Croarkin, C. & Varner, R. (1982). Measurement assurance for dimensional measurements on integrated-circuit photomasks, NBS Technical Note 1164, U.S. Department of Commerce, Washington, DC.
- [13] ISO 5725-6 (1994). *Accuracy (Trueness and Precision) of Measurement Methods and Results – Part 6*. International Organization for Standardization, Geneva.
- [14] Jensen, D.R., Hui, Y.V. & Ghare, P.M. (1984). Monitoring an input-output model for production. I. The control charts, *Management Science* **30**, 1197–1206.
- [15] Walker, E. & Wright, S.P. (2002). Comparing curves using additive models, *Journal of Quality Technology* **34**, 118–129.
- [16] Young, T.M., Winistorfer, P.M. & Wang, S. (1999). Multivariate control charts of MDF and OSB vertical density profile attributes, *Forest Products Journal* **49**, 79–86.
- [17] Jin, J. & Shi, J. (1999). Feature-preserving data compression of stamping tonnage information using wavelets, *Technometrics* **41**, 327–339.

- [18] Lada, E.K., Lu, J.C. & Wilson, J.R. (2002). A wavelet-based procedure for process fault detection, *IEEE Transactions on Semiconductor Manufacturing* **15**, 79–90.
- [19] Winistorfer, P.M., Young, T.M. & Walker, E. (1996). Modeling and comparing vertical density profiles, *Wood and Fiber Science* **28**, 133–141.
- [20] Boeing Commercial Airplane Group, Material Division, Procurement Quality Assurance Department (1998). *Advanced Quality System Tools*, AQS D1-9000-1, The Boeing Company, Seattle.

Related Articles

Change-Point Methods; Control Charts, Overview; Exponentially Weighted Moving Average (EWMA) Control Chart; Multivariate Control Charts Overview; Regression Control Charts.

WILLIAM H. WOODALL

Control Charts, Selection of

Introduction

Control chart selection is the process of matching an appropriate **control chart** with a given application. Successful process monitoring requires several key decisions. Some of these decisions are driven by strategic plans for quality improvement. Examples include the choice of improvement projects, key processes relevant to the project, and **data collection** strategies, such as when and where samples are obtained. Quality leaders must decide on proper actions in response to control chart signals. In this article we only consider the problem of selecting an appropriate type of control chart for a given application. The guidelines provided in subsequent sections are intended to be guiding principles and

not hard and fast rules to which one must adhere. Even experts in the quality area may disagree on final choices of monitoring schemes. There are common themes upon which most experts will agree and these are the focus in the following sections.

Consider the Data

The logic for selecting control charts is similar to general guidelines used in selecting statistical techniques. An appropriate choice among statistical tools depends on the question one needs to answer and the data available to the decision maker. For **quality control** applications, the typical questions of interest concern changes in the mean and/or variance of key process characteristics. For example, we may wish to ensure that produced items are of the correct length on average, and are consistently near target values.

One must investigate several questions concerning the data prior to selecting the type of control chart. These questions include:

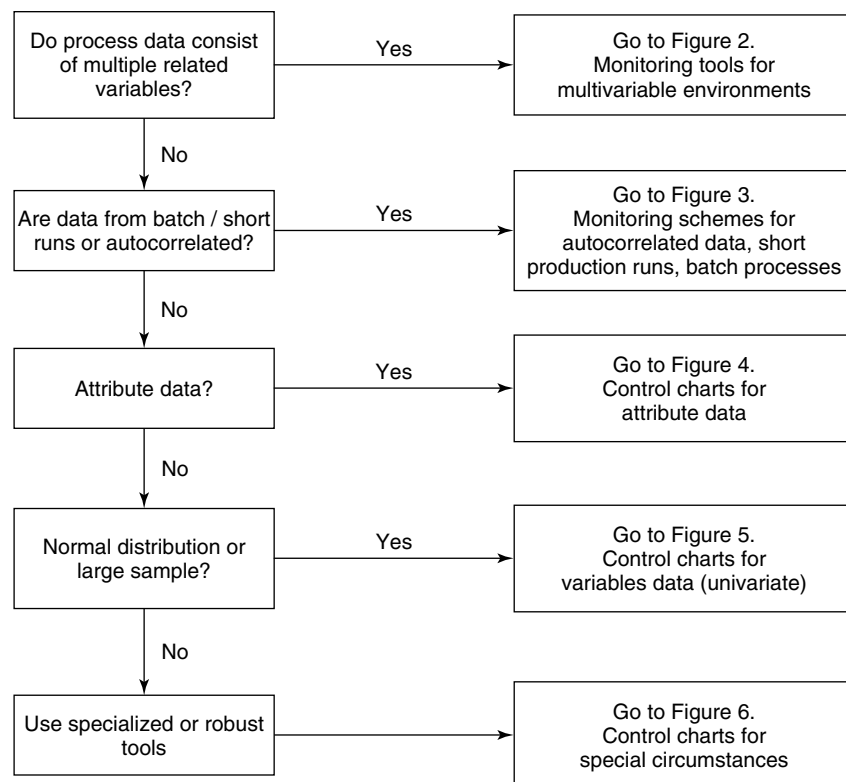


Figure 1 Guidelines for control chart selection

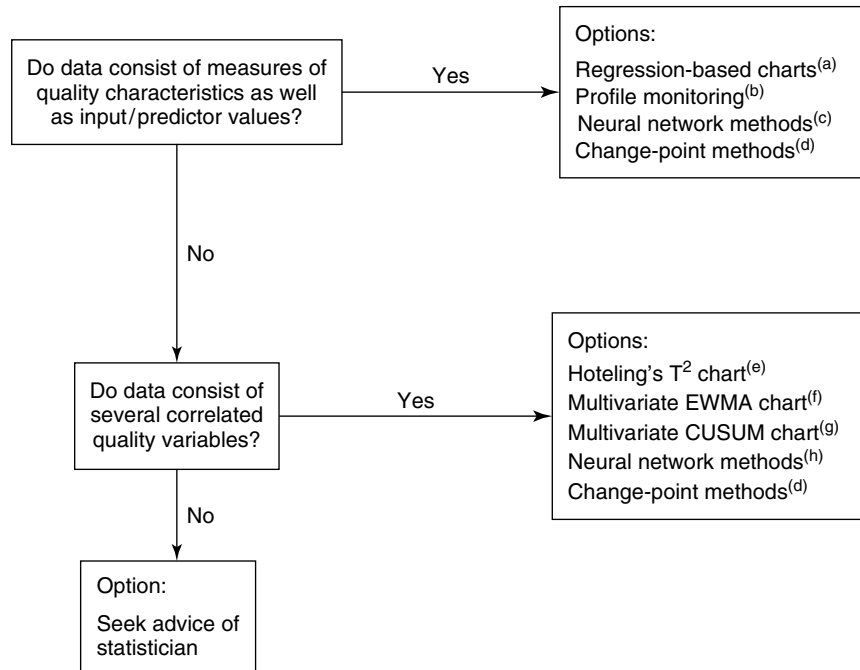


Figure 2 Monitoring tools for multivariable environments ^(a)See **Regression Control Charts** ^(b)See **Profile Monitoring** ^(c)See **Neural Networks: Construction and Evaluation** ^(d)See **Change-Point Methods** ^(e)See **Hotelling's T^2 Chart; Multivariate Charts for Variability; Multivariate Control Charts, Interpretation of** ^(f)See **Multivariate Exponentially Weighted Moving Average (MEWMA) Control Chart** ^(g)See **Multivariate Cumulative Sum (CUSUM) Chart**

1. What type of data is available?
2. How much data is available?
3. Are data associated with long production runs or some other process environment such as **batch processes** or short production runs?
4. Is a single-quality characteristic (univariate process control) or multiple quality characteristics (multivariate process control) being considered?
5. If multiple quality characteristics are important, are data values for the various key quality characteristics correlated with one another?
6. Are data values for a given quality characteristic correlated over time (autocorrelated)?
7. What are the distributional properties of the population as reflected by the data? Are common probability distributions, such as the normal, binomial or Poisson distributions useful approximations?
8. Are key output quality characteristics driven by key process input variables? If so, do we have data on both inputs and outputs?

Important questions concerning the process and goals of the monitoring scheme too exist. Depending on the capability of a process, one must determine the size of a shift deemed important enough to detect quickly. Tools for short production runs or batch processes will differ from tools appropriate for long production runs. Unusual or extreme economic consequences associated with process shifts should be noted.

Types of Data

Variables for which data are collected should reflect the current state of a given process quality characteristic. Variables can typically be classified as either quantitative (continuous) or attribute (qualitative, categorical) variables. Quantitative measurements, sometimes referred to as *variables data*, assume values over some continuum, such as length, width, weight, or time. Hence, quantitative variables are associated with continuous probability distributions such as the **normal distribution** or

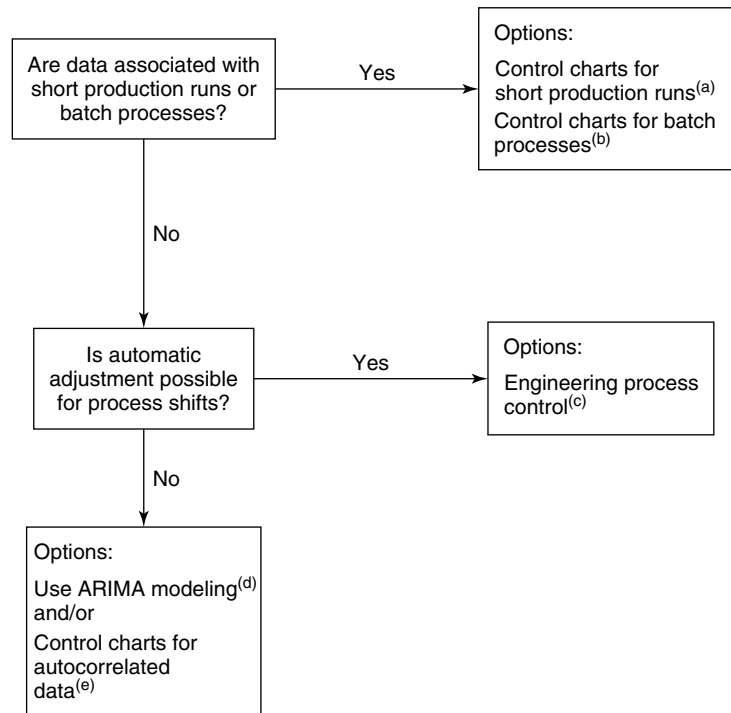


Figure 3 Monitoring schemes for autocorrelated data, short production runs, or batch data ^(a)See **Control Charts for Short Production Runs** ^(b)See **Control Charts for Batch Processes** ^(c)See **Engineering Process Control; Feedback Control** ^(d)See **Large Autoregressive Integrated Moving Average (ARIMA) Modeling** ^(e)See **Autocorrelated Data**

exponential distribution when appropriate. Attribute data reflect more general properties of process output (e.g., number of surface defects or the classification of items as conforming *versus* nonconforming items). Proportions of nonconforming items found in a sample or counts of defects observed over a given time or space are associated with discrete probability distributions, such as the binomial or Poisson distribution, when appropriate. The distinction between variables with continuous or discrete distributions can become blurred. Time is generally associated with a continuous distribution; however, we report our ages as values truncated to integer values giving discrete data values.

Control charts are often designed for a specific type of data and an associated probability distribution. When these distributions are reasonable, control chart performance is predictable. Control charts associated with specific distributional assumptions can provide quite unexpected performances when

these assumptions are violated. In such cases, one might choose monitoring schemes that are robust to the given assumptions such as nonparametric methods.

Amount of Data

The size of samples as well as the sampling frequency influences the amount of data that one encounters. Additionally, the number of variables required to characterize a process impacts the amount of data one must process. Tools associated with relatively large but infrequent samples for monitoring a single process variable may be inappropriate for small samples or infeasible for automated data collection environments with near continuous data collection on several variables.

Univariate versus Multivariate Environments

One might argue that we always collect data over time, so we always have at least two variables of

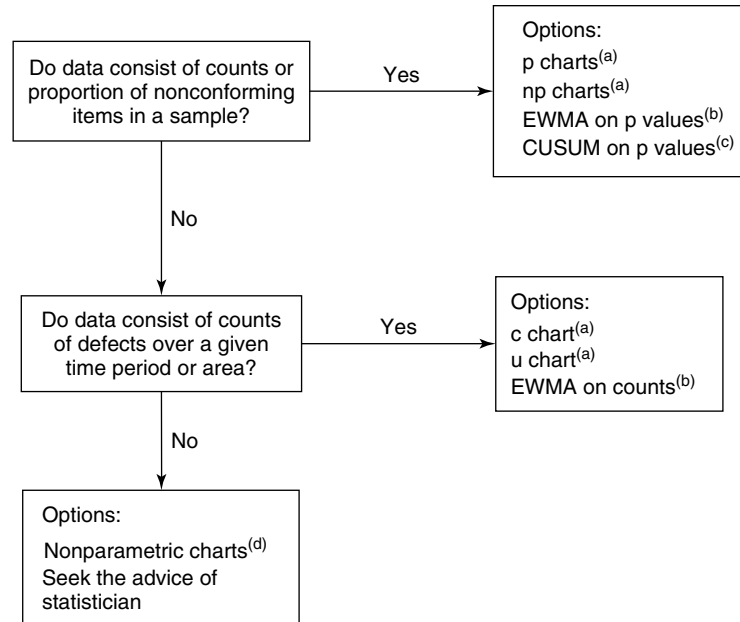


Figure 4 Control charts for attribute data ^(a)See **Control Charts for Attributes** ^(b)See **Exponentially Weighted Moving Average (EWMA) Control Chart** ^(c)See **Cumulative Sum (CUSUM) Chart** ^(d)See **Control Charts, Nonparametric**

interest (the quality characteristic and time). Univariate methods are designed to monitor a single-quality characteristic over time. One must address the issue of **autocorrelation** when using univariate monitoring schemes. If values of a variable are correlated over time, special tools are available for exploiting or at least accounting for this source of variation in the process. If the values are not correlated over time, then the added effort and complexity are not worth the gain in information. Simpler tools will be just as informative.

In other cases, one must process data on multiple related variables in a process. Multivariate methods are designed to monitor multiple variables associated with process output characteristics. If these variables are correlated in a quantifiable manner, specialized tools that exploit the **correlation** are appropriate. If the variables are statistically unrelated, multiple univariate methods may be simpler and efficient to use.

Finally, one can encounter situations involving multiple variables, some of which are process output variables and some are process input (driver, explanatory) variables. Understanding the types of data and relationship among the variables allows the

use of specialized tools such as **profile monitoring** or regression-based monitoring schemes.

Guidelines for Control Chart Selection

This section contains several figures helpful in narrowing the appropriate choices of process monitoring schemes. Figure 1 provides our starting point and guidance to subsequent figures and data considerations. For many situations the practitioner has multiple options that are reasonable monitoring schemes. Other articles within this encyclopedia provide specific strengths and weaknesses of each approach and allow the reader to make final selections of tools.

Figure 2 contains the options available to practitioners who must monitor multivariable data. Multivariable data include the case of multiple correlated quality characteristics (multivariate data) and the case of quality characteristics and their associated explanatory variables.

Figure 3 contains the options appropriate for autocorrelated data, short production runs, and batch processes. For autocorrelated data from long production runs, one must determine if automatic process adjustment is a possibility. If so, the techniques associated

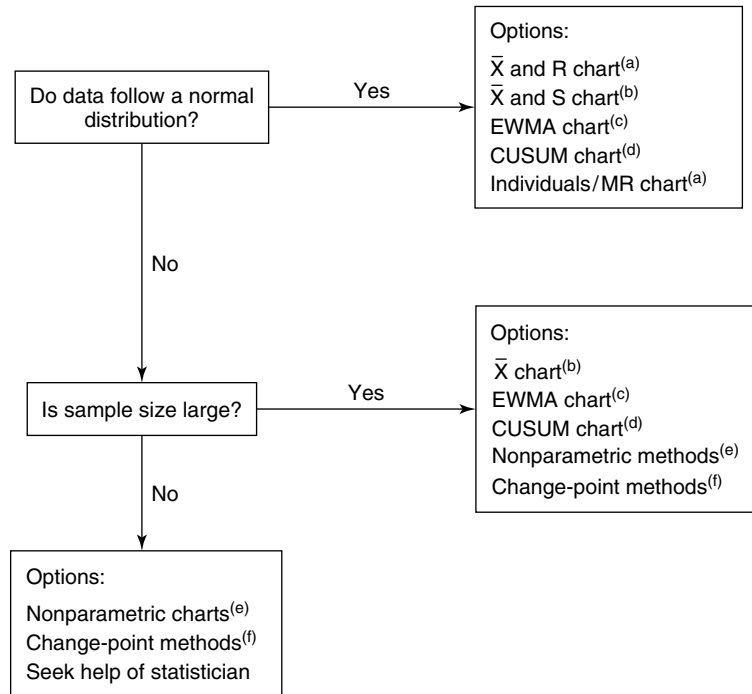


Figure 5 Control charts for variables data (univariate) ^(a)See **Control Charts for the Mean; Moving Range and R Charts** ^(b)See **Control Charts for the Mean; Control Charts for the Standard Deviation** ^(c)See **Exponentially Weighted Moving Average (EWMA) Control Chart** ^(d)See **Cumulative Sum (CUSUM) Chart** ^(e)See **Control Charts, Nonparametric** ^(f)See **Change-Point Methods**

with engineering process control (see **Engineering Process Control**) may be more appropriate than traditional statistical process control (see **Control Charts, Overview**).

Figure 4 contains options for attribute data. The most common types of attribute data are proportions of nonconforming items (associated with the binomial distribution) and counts of defects (associated with the Poisson distribution).

Figure 5 contains options for variables data. The appropriate choice among techniques is often driven by normality assumptions or large sample sizes. For small samples which are not approximately normal, robust monitoring techniques are required.

Figure 6 contains some of the options available to practitioners in special situations. These situations include environments where none of the prior figures is appropriate, economic considerations are an overriding concern in control chart development, or control schemes must be built on specification limits.

Although the figures above capture the great majority of cases encountered in practice, they are not all-inclusive. There are some situations which do not fall precisely into any of the environments described above, and similarly, there are good monitoring schemes that are not included in any of the figures. However, practitioners will find the foregoing guidelines a starting point for further investigation, and for the study of process monitoring tools.

Conclusions

The selection of appropriate monitoring schemes requires an understanding of process behavior and basic concepts underlying data and statistics. Control charts range from simple yet powerful statistical tools to quite sophisticated technical monitoring systems. Choosing an appropriate monitoring scheme among many can seem like a complicated process. In some circumstances, the selection process can indeed

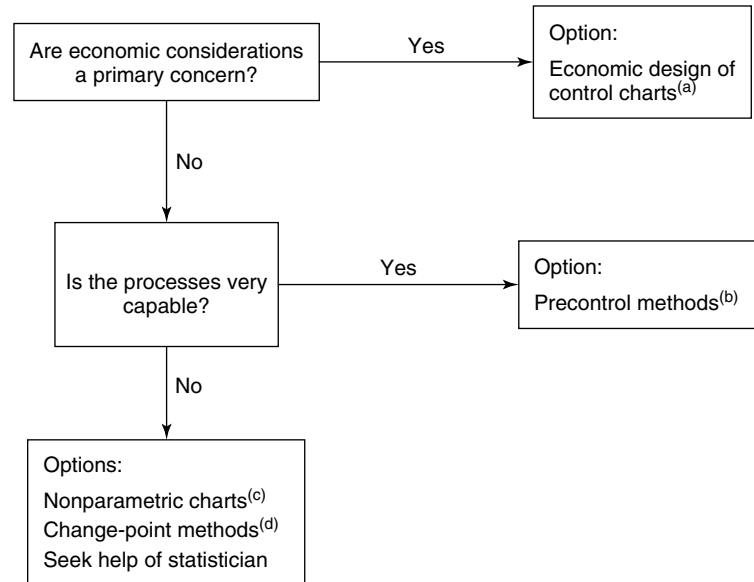


Figure 6 Control charts for special circumstances ^(a)See **Economic Design of Control Charts** ^(b)See **Precontrol** ^(c)See **Control Charts, Nonparametric** ^(d)See **Change-Point Methods**

be challenging. Practitioners who find the process difficult should know that statistical assistance is generally available through a number of sources (including, in North America, local chapters of the American Society for Quality (ASQ) and the American Statistical Association (ASA)). Finally, one should find comfort in the knowledge that suboptimal (but reasonable) monitoring schemes applied to good process data is typically far, far superior to no monitoring scheme at all.

Related Articles

Autocorrelated Data; Average Run Lengths and Operating Characteristic Curves; Change-Point Methods; Control Charts for Attributes; Control Charts for Batch Processes; Control Charts for Short Production Runs; Control Charts for the

Mean; Control Charts for the Standard Deviation; Control Charts, Nonparametric; Control Charts, Overview; Cumulative Sum (CUSUM) Chart; Economic Design of Control Charts; Engineering Process Control; Exponentially Weighted Moving Average (EWMA) Control Chart; Feedback Control; Hotelling's T^2 Chart; Large Autoregressive Integrated Moving Average (ARIMA) Modeling; Moving Range and R Charts; Multivariate Charts for Variability; Multivariate Control Charts Overview; Multivariate Control Charts, Interpretation of; Multivariate Cumulative Sum (CUSUM) Chart; Multivariate Exponentially Weighted Moving Average (MEWMA) Control Chart; Neural Networks: Construction and Evaluation; Precontrol; Profile Monitoring; Regression Control Charts; Variable Sampling Rate Control Charts.

BENJAMIN M. ADAMS

Engineering Process Control

Introduction

Engineering **process control** (EPC) (*see Feedback Control*) and **statistical process control** (SPC) (*see Control Charts, Overview*) are techniques for reducing process variation while maintaining target values during production. While both techniques share similar goals, the strategic and mathematical foundations for the two methods are quite different. The goal of this article is to delineate the two methods and provide guidance for their appropriate use.

EPC reduces variation by quantifying process deviations from target. Target values are maintained by process adjustments to compensate for deviations from target values. Underlying this approach is the ability to estimate process levels (current and future) and the existence of an adjustment variable to “control” the process level. The primary focus is reduction of variability by compensating for process shifts. Minimal attention is given to identification and removal of the source cause of the shifts. Process adjustments are typically automated and driven by process measurement rather than operator intervention.

SPC reduces variation by estimating process levels, also. The primary objective is to determine whether the deviations from target values result from common cause variation inherent in the production process, or from special cause variation such as process shifts or malfunctions. Special cause variation results from assignable causes that may be identified and removed. This identification and removal of special causes of variation lead to long-term reduction in process variation and improved process stability. The primary focus is reduction of variation through greater process understanding and “control”. Process adjustments are not typically automated, but require operator investigation and intervention.

The brief descriptions above suggest that EPC and SPC are complementary rather than alternative techniques. The approaches can be combined with good results. Consider the problem of maintaining the temperature of a room at a target value on a summer day. The typical air conditioner exhibits the character of EPC. The room temperature is monitored

(electronically or mechanically) and adjustments are made accordingly. If the temperature is too high, the air conditioner turns on. When the temperature drops too far below target levels, the air conditioner turns off. The inherent drift in temperatures from the coolest to the hottest part of the day is minimized by automatic adjustments by the air conditioning unit. Suppose a window or door in the room were opened and inadvertently left open. The air conditioner would simply continue to run in an effort to compensate for the temperature increase. The air conditioning unit itself would not typically signal that a special cause source of variation (the open door or window) potentially exists. SPC might be used in conjunction with the EPC controller to identify special causes of variation, unaccounted for by EPC. One might monitor the rate or length of EPC activity to identify and intervene in special cause situation such as the open door or window. This example is a gross simplification of most EPC applications, but illustrates the idea of combining EPC and SPC. The combination of EPC and SPC is sometimes called *algorithmic SPC*. Strengths and weaknesses of EPC and SPC are provided in the subsequent sections.

Engineering Process Control – The PID Controller

Proportional–integral–derivative (PID) control is an EPC tool which reduces process variability by adjusting some measurable process variables based on process feedback and compensation. PID control is named after its three terms: proportional, integral, and derivative. Each of the terms performs a different task and has a different effect on the process output. The proportional term responds to the immediate error, but cannot guarantee achieving the target value. In other words, the proportional term can reduce the difference between the output and the target value, but it cannot eliminate the steady state error. The integral term accounts for eliminating the steady state error and it also erases disturbances. However, the integral term does decrease the process stability since it adds up all past errors. The derivative term predicts the future process behavior, and therefore can stabilize the process. A PI controller is also widely used in situations where a derivative term is not needed. The most used form of PID is the parallel form, where the P, I and D terms are made parallelized with each

other and given the same error input. The output is often expressed in the ideal form:

$$\text{Output} = K_p e(t) + K_i \int_0^t e(\tau) d\tau + K_d \frac{de}{dt} \quad (1)$$

where $e(t)$ is the error signal, which is the difference between the target value and the measured output. The values K_p , K_i , and K_d are constants used to “tune” the PID controller. The application of PID has grown tremendously since the invention of PID control in 1910 and the Ziegler–Nichols (Z–N) tuning rule in 1942 [1]. Today, the research and development of PID control focuses on getting the best out of PID [2] and finding the next key technology or methodology for PID tuning [3]. A complete list of PID tuning patents since 1971 can be found in an article by Li *et al.* [4]. O’Dwyer compiles 245 tuning rules in his book, *Handbook of PI and PID controller tuning Rules*, among which 104 are for PI controllers and 141 are for PID controllers. The choice of tuning rules depends on the process type, the parameters, the uncertainties, and so on [5]. For a system that can be taken offline for tuning, the input is increased in a small step, and the output is measured as a function of time, then control parameters are calculated [6]. For a system that cannot be taken offline for tuning, the Z–N type of tuning can be used, in which the controller settings are expressed explicitly in terms of ultimate gain and ultimate period [5]. Modern industrial process largely relies on tuning optimization software to ensure consistent results.

Strengths of EPC

EPC techniques have a number of strengths. Knowledge of EPC behavior and the available research in the area is quite rich. The application of EPC methods has a long and successful history across a wide range of industries.

EPC techniques are especially well suited for situations where process levels drift over time. Such processes produce auto-correlated data (see **Time Series Analysis**). It is worth noting that monitoring auto-correlated data (see **Autocorrelated Data**) is a particularly challenging problem for SPC practitioners. For auto-correlated processes PID control compensates and regulates the process variability

automatically by adjusting measurable variables. PID assumes (a) the next observation is predictable in the process, (b) there are variables that we can manipulate in order to affect the process, and (c) the effect of this manipulated variable can be estimated so that we can determine how much control action to apply [7]. When these conditions are met, PID control can be very effective. In these situations, EPC is more suitable than traditional SPC methods. Process deviation can be compensated by adjusting other known variables, and consequently, the process output can be kept close to the target values by a series of adjustments to these variables [7].

Weaknesses of EPC

PID control does exhibit some limitations relative to SPC. One disadvantage of PID control is that it works only with measurable variables. As mentioned above, PID control adjusts some measurable variables based on feedback in order to compensate the output deviation and keep the output close to its target value. However, there are situations where the inspected item is either classified as “defective” or “nondefective” and quality characteristics are not represented numerically. For example, a sensor used in a computer system is either functional or nonfunctional. Quality characteristics of this type are called *attributes* (see **Control Charts for Attributes**). PID has no ability to control this type of quality characteristics. SPC, on the other hand, can be used under these circumstances by implementing **control chart** for fraction nonconforming (p chart), control chart for nonconformities (c chart), or control chart for nonconformities per unit (u chart).

The implementation of PID largely relies on PID tuning which decides the constants used in the P, I, and D terms to model the process. The PID tuning eventually determines how well PID controllers adjust the process. In practice, many PID controllers do not tune to their optimal parameters. Seborg *et al.* reported that 30% of PID controllers in the pulp and paper industry gave poor performance due to poor tuning [8]. Another common PID problem is “integral windup”. The windup is due to the controller states becoming inconsistent with the saturated control signal, and as a result, future correction being ignored [9]. This might be fixed by introducing a more aggressive differential term or disabling the

integral term until measured variables have entered the proportional band [6]. There are also large problems associated with the derivative term [9]. One of the common derivative term problems is that small amounts of noise can cause large amounts of change in the output. Setting up the derivative term is so complicated that 80% of PID controllers have this term switched off, or omitted completely [2].

Many PIDs control a valve or similar mechanical device. Wear of the valve or device can be a major reason of poor PID performance and increased maintenance costs. Seborg *et al.* reported that 20% of PID controllers in the pulp and paper industry gave poor performance owing to valve problems [8]; Martin reported that half of the valves in the chemical process industry needed to be fixed [10]. Finding diagnostic tools to identify potential valve problems is still at its early stage [11–14]. It is almost impossible to restore all valves to their expected performance, and efforts may involve a large amount of money.

The failure or misperformance of PID usually has severe consequences. The PID controller actually affects final products as it adjusts measurable variables in order to change the output. If there is anything wrong with the PID controller itself, either mistuning, valve problems as mentioned above, or any nonfunctional parts in PID, it may continuously produce defective products until they are recovered by either product inspectors (if there are any) or final users. In the later case, returns or after sale service might be expected and demand extra company resources.

For a more extreme example, consider the case of a motorist using automatic cruise control. The failure or misperformance of a cruise control system, which is a PID controller, might put the person behind the wheel in danger. Hydroplaning is a well-known phenomenon. In rain, a layer of water builds up beneath tires. As you drive at a high speed, the car begins sliding on this layer, and tires lose contact with the road sometimes. If a cruise control system is activated, it will try to spin the tires faster to compensate the reduction in speed. The driver may lose control when tires contact the road again. An accident happens as a consequence.

The ability of EPC and SPC to handle the false alarms also differs. False alarms occur because the process variability cannot be totally eliminated. For single unusual measurements associated with no process abnormality, PID control will respond to this

false alarm and adjust its current correct settings to incorrect settings based on false information. Consequently, defective products may be produced. PID control does have the ability to correct this error over time. However, defective products have already been produced.

PID control automatically adjusts to keep the process on target. No action is taken to identify assignable causes in the process. One is unable to determine if there is a problem in the production line, such as poor operation, raw material problems, and so on. On the other hand, SPC can alert management when there are problems in the process. The control charts will have out of control data, which warns the management that there are problems in the process and corresponding actions need to be taken to correct the problems.

Conclusions

The choice between EPC and SPC methods is driven by the characteristics of the production process and the goals of the practitioner. EPC methods are especially well suited for stable processes whose mean levels tend to drift in a predictable manner and for which automatic measurement and adjustment is possible. SPC methods are more appropriate for environments where automated control is not feasible, or when assignable causes of process shifts must be identified and eliminated. While the approaches of EPC and SPC have been viewed as competing philosophies and methodologies, one may reasonably view them as complementary tools. The strengths of one approach complement the weaknesses of the other.

As a final note, using EPC and SPC in combination is a particularly powerful option. The simultaneous use of these methods is sometimes called *algorithmic SPC*. The interested reader may find more information in Wander Weil *et al.* [15].

References

- [1] Ziegler, J.G. & Nichols, N.B. (1942). Optimum settings for automatic controllers, *Transactions of ASME* **64**(8), 759–768.
- [2] Digest, I.E.E. (1996). Getting the best out of PID in machine control, in *Digest IEE PG16 Colloquium (96/287)*, London, 24 October 1996.

- [3] Marsh, P. (1998). Turn on, tune in – Where can the PID controllers go next, *New Electronics* **31**(4), 31–32.
- [4] Li, Y., Ang, K.H. & Chong, G.C.Y. (2006). Patents, software, and hardware for PID control, *IEEE Control Systems Magazine* **26**(1), 42–54.
- [5] Yu, C.C. (2006). *Autotuning of PID Controllers—A Relay Feedback Approach*, Springer-Verlag London.
- [6] http://en.wikipedia.org/wiki/PID_controller 2007.
- [7] Montgomery, D.C. (2005). *Introduction to Statistical Quality Control*, 5th Edition, John Wiley & Sons.
- [8] Seborg, D.E., Edgar, T.F. & Mellichamp, D.A. (2004). *Process Dynamics and Control*, 2nd Edition, John Wiley & Sons, New York.
- [9] Li, Y., Ang, K.H. & Chong, G.C.Y. (2006). PID control system analysis and design, *IEEE Control Systems Magazine* **26**(1), 32–41.
- [10] Martin, J.J. (1979). Cubic equation of state – Which? *Industrial and Engineering Chemistry Fundamentals* **18**, 81.
- [11] Zhuang, M. & Atherton, D.P. (1993). Automatic tuning of optimum PID controllers, *IEE Proceedings D* **140**(3), 216–224.
- [12] Luyben, W.L. (2001). Effect of derivative algorithm and tuning selection on the PID control of dead time processes, *Industrial and Engineering Chemistry Research* **40**, 3605.
- [13] Marlin, T.E. (2000). *Process Control*, 2nd Edition, McGraw-Hill, New York.
- [14] Morari, M. & Zafiroiu, E. (1989). *Robust Process Control*, Prentice Hall, Englewood Cliff.
- [15] Vander Weil, S., Tucker, W.T., Faltin, F.W. & Doganaksoy, N. (1992). Algorithmic statistical process control: concepts and an application, *Technometrics* **43**(3), 189–193.

Further Reading

O'Dwyer, A. (2003). *Handbook of PI and PID Controller Tuning Rules*, Imperial College Press, London.

Related Articles

Autocorrelated Data; Control Charts, Overview; Feedback Control.

BIN HE AND BENJAMIN M. ADAMS

Feedback Control

Introduction

A process is a transformation between an input and an output (Figure 1). Engineering **process control** (EPC) is a discipline that deals with controlling the output of a specific process; the goal of which is to reduce variability in a key variable at the output of a process. For example, maintaining the temperature in a room is a process that has the specific outcome to reach a defined temperature and keep it constant over time. In Figure 2, the temperature is the controlled or output variable. At the same time, it is an input or variable since it is measured by a thermometer and used to decide whether to heat or not to heat. The desired temperature is the target or set point. The state of the heater (e.g., the setting of the valve allowing hot water to flow through it) is called the *manipulated variable* since it is subject to control actions. The controller makes adjustments to the variable so that the process stays on target.

Figure 1 represents the simplest, but unlikely structure for the input and output variables. Typically, there is noise in the system as Figure 3 demonstrates. Figure 3 also shows a dynamic system. In a dynamic process, the characterization of a relationship that

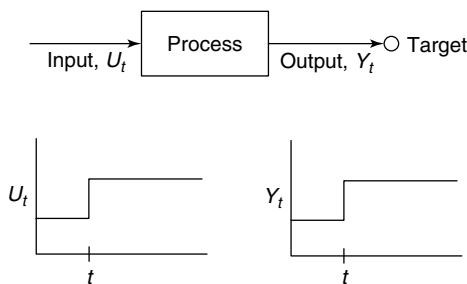


Figure 1 A basic process model

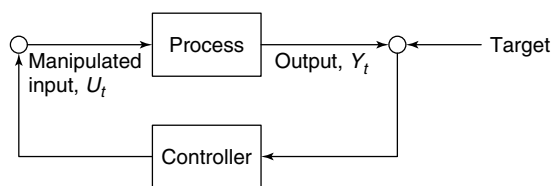


Figure 2 A controlled process

links the input to the output is more difficult than in a steady-state process. A change to a dynamic process may take several time periods to be fully realized, whereas a change to a steady-state process is observed in the next time period.

Dynamic behavior occurs when the output, Y_t , of a process does not respond instantaneously to a change in the manipulated input, U_t . The modeling problem becomes more challenging with the presence of a noise process that affects the output. Noise may occur anywhere in the system; the complete effect of the noise is measured at the output and hence represented as affecting the output. The noise process, y_t , is modeled as a series of random shocks, u_t , that pass through a linear filter. Figure 3 also shows that a controller uses the information available at the output, that is, the sum effect of y_t and Y_t to adjust the process. This type of model has been widely used to represent dynamic industrial processes that are sampled over time [1–3].

Control Strategies

There are many control strategies or algorithms that the controller may use. A control strategy is where a variable(s) is measured and the measurement is used to adjust another process variable which can be manipulated. A more common name for this activity is EPC. A typical control objective is to minimize the variance of the output deviations from a target (or set point) or minimum mean square error (MMSE) control. A criterion in developing an adjustment policy under this objective is cost. If the cost of a process adjustment is zero, then an adjustment is made at every sampling interval. If the cost of adjustment is nonzero, then an adjustment is made only if a significant deviation from the target occurs.

Consider a process in which the output at a specific period of time is Y_t and the process target is T . The goal of proportional-integral-derivative (PID) control is to keep the process output Y_t close to the target T . By adjusting a controllable or manipulative variable the error term, $e_t = T - U_t$, is kept close to zero. There are different methods available in literature, but the most common ones are proportional control, integral control, derivative control, and a combination of these three also known as *proportional-integral-derivative controller*. In EPC, PID controllers are a standard [4, 5].

2 Feedback Control

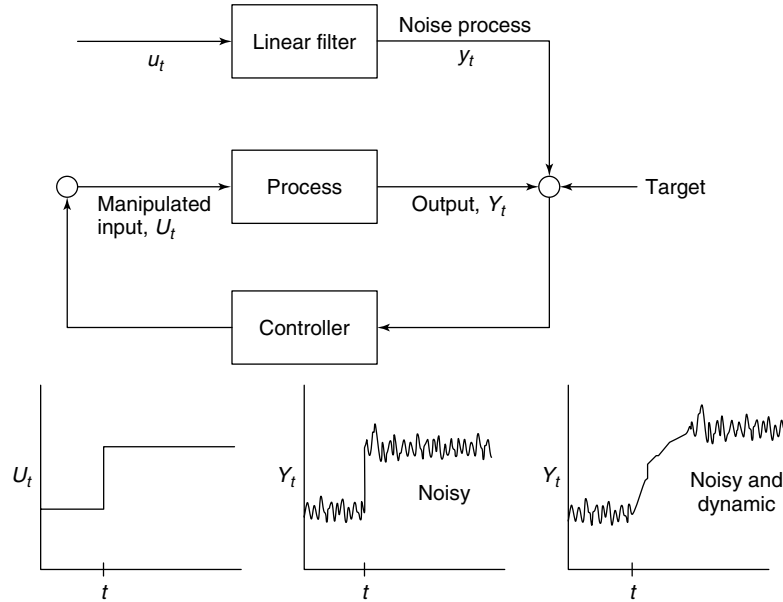


Figure 3 A general model of a process, noise process, and controller

Proportional Control

Proportional control adjusts the controllable variable by a certain proportion called the *gain*: $Y_t = K_p e_t$, where the constant K_p is the proportional gain. The larger the error or offset is, the larger K_p is to bring the process near the target. The drawback to proportional control is that the system may never achieve steady state: the system either oscillates back and forth around the set point because there is nothing to remove the error when it overshoots, or oscillates and/or stabilizes at a too low or too high value.

Integral Control

Integral control sums the current and previous errors over a period of time and added to the manipulated variable: $Y_t = K_i \int_0^t e_t dt$ where the constant K_i is the integral gain. By appropriately selecting an integral control gain term, the steady-state error can be reduced significantly. However, the system may become unstable and produce an overshooting or undershooting effect.

Derivative Control

Derivative control provides a control proportional to the rate at which a process output is deviating with

respect to its previous output: $Y_t = K_d de/dt$ where the constant K_d is the derivative gain. The response to a system change largely depends on the magnitude of K_d : the larger the term, the more quickly the controller responds. One of the main advantages of using derivative control is the minimization of the overshooting or undershooting effect by using information about the rate of change of process output. However, the transient response time is increased by selecting a high value of K_d .

PIDControl

The PID controller combines the above control strategies for control over steady-state and transient errors: $Y_t = K_p \left(e_t + K_i/K_p \int_0^t e_t dt + K_d/K_p de/dt \right)$.

It provides a mechanism to bring an automatic process to target by providing feedback adjustments. It is a very robust tool for process adjustments and enjoys a high esteem in the process control literature. These control strategies may also be combined to form PI and PD controllers. In process control applications, PI control (which is also the exponentially weighted moving average **EWMA**) is responsible for more than 95% of the system control loops. Although these tools are common for making adjustments to

a process, they are usually designed for steady-state behavior [1, 6, 7].

EPC/SPC Integration

There has been a significant amount of discussion in the literature on the difference of statistical process control (SPC) and EPC in terms of their applications context. Traditional SPC emphasizes that any observed difference between a controlled process variable and its desired value may be due to inherent common cause variation (thus no system change is necessary) or systematic special cause variation (where a system change may be necessary). The goal of SPC is to monitor the process so that there is a good chance of detecting the special cause so that it can be removed. The underlying assumption is that there is a disincentive to make an adjustment after every observation period because they are costly. On the other hand, traditional EPC assumes that the process input variables can be adjusted as frequently as desired, that the adjustments are virtually cost free, and that the process output variables can be measured on line without any associated uncertainty [8].

Recently, we have seen an increasing research enthusiasm in combining (or integrating) SPC and EPC techniques despite their differences. One benefit of the combination stems from reducing the mean squared error. By integrating basic EPC and SPC policies, it is possible to reduce the output variance in some systems beyond the capabilities of either tool individually [6, 9–13]. Montgomery *et al.* [14] demonstrate this using a modified funnel experiment in combination with an SPC chart monitoring deviation from target. Another benefit is making an adjustment less often. This may be particularly relevant for systems where there is a distinct cost associated with making each adjustment or where the adjustment must be manually imposed by an operator.

Integrating EPC and SPC has applications to startup periods [15, 16], to run-to-run control [17, 18], and to multivariable control [19, 20] to name just a few.

References

- [1] Box, G.E.P., Jenkins, G.M. & MacGregor J.F. (1974). Some recent advances in forecasting and control part II, *Journal of the Royal Statistical Society. Series C, Applied Statistics* **23**(2), 158–179.
- [2] Box, G.E.P., Jenkins, G.M. & Reinsel, G.C. (1994). *Time Series Analysis: Forecasting and Control*, 3rd Edition, Prentice Hall, Englewood Cliffs.
- [3] Seborg, D.E., Edgar, T.F. & Mellichamp, D.A. (1989). *Process Dynamics and Control*, John Wiley & Sons, New York.
- [4] Dorf, R.C. & Bishop, R.H. (1995). *Modern Control Systems*, 7th Edition, Addison-Wesley, Reading.
- [5] Ogata, K. (1987). *Discrete-Time Control Systems*, Prentice Hall, Englewood Cliffs.
- [6] MacGregor, J.F. (1988). On-line statistical process control, *Chemical Engineering Progress* **84**, 21–31.
- [7] MacGregor, J.F. (1990). A different view of the funnel experiment, *Journal of Quality Technology* **22**(4), 255–259.
- [8] Ogunnaike, B.A. & Ray, W.H. (1994). *Process Dynamics, Modeling and Control*, Oxford University Press, New York.
- [9] VanderWiel, S.A. (1992). Monitoring processes that wander using integrated moving average models, *Technometrics* **38**, 139–151.
- [10] Box, G.E.P. & Luceño, A. (1997). *Statistical Control by Monitoring and Feedback Adjustment*, John Wiley & Sons, New York.
- [11] Del Castillo, E. (2002). *Statistical Process Adjustment for Quality Control*, John Wiley & Sons, New York.
- [12] Box, G.E.P. & Kramer, T. (1992). Statistical process monitoring and feedback adjustment – a discussion, *Technometrics* **34**(3), 251–257.
- [13] Del Castillo, E. (2006). Statistical process adjustment: a brief retrospective, current status and future research, *Statistica Neerlandica* **60**(3), 309–326.
- [14] Montgomery, D.C., Keats, J.B., Runger G.C. & Messina W.S. (1994). Integrating statistical process control and engineering process control, *Journal of Quality Technology* **26**(2), 79–87.
- [15] Nemphard, H.B. & Mastrangelo, C.M. (1998). Integrated process control for startup operations, *Journal of Quality Technology* **30**(3), 201–211.
- [16] Nemphard, H.B., Mastrangelo, C.M. & Kao, M.S. (2001). Statistical monitoring performance for startup operations in a feedback control system, *Quality and Reliability Engineering International* **17**, 379–390.
- [17] Tseng, S.T., Tsung, F. & Liu, P.Y. (2007). Variable EWMA run-to-run controller for drifted processes, *IIE Transactions* **39**(3), 291–301.
- [18] Del Castillo, E. & Hurwitz, A. (1997). Run-to-run process control: literature review and extensions, *Journal of Quality Technology* **29**(2), 184–196.
- [19] Del Castillo, E. & Rajagopal, R. (2002). A multivariate double EWMA process adjustment scheme for drifting processes, *IIE Transactions* **34**(12), 1055–1068.
- [20] Tseng, S., Chou, R. & Lee, S. (2002). A study of multivariate EWMA controller, *IIE Transactions* **34**(6), 541–549.

CHRISTINA MASTRANGELO
AND NAVEEN KUMAR

Change-Point Methods

Introduction

Statistical process control (SPC) methods are intended to detect departures from an in-control state, which are the result of special cause variability. Two broad classes of departures are isolated or intermittent ones, in which the special cause appears and then disappears even in the absence of some corrective action, and sustained or persistent ones that, once introduced, tend to continue until corrective action remedies them. Persistent special causes further subdivide into those that have erratic effects and those that shift the process from its in-control state to a different, stable out-of-control state. These latter special causes provide the motivation for the change-point formulation.

Notation and Approaches

From the statistical viewpoint, the in-control and out-of-control states are distinguished by the distribution of the process measurement stream. Write $X_1, X_2, \dots, X_n, \dots$, for the sequence of readings (which could be uni- or multivariate). A convenient starting point for discussion is that, while in control, the readings are independent and follow some statistical distribution $f(x, \theta)$ with the parameter vector θ having its in-control value θ_0 . Under a persistent change at some time τ , the readings follow this distribution up to time τ , after which their behavior changes to something else. Under the conventional change-point model, this changed behavior has the same distribution with the parameter θ taking the new value θ_1 . If the change point τ were known, we would have a standard two-sample comparison. Sometimes this is the case – we know that something happened at time τ (for example a plant switched suppliers at time τ) and want to test whether it had a visible effect on the process. More commonly, we do not know that changes have intervened until we detect and diagnose their effects. In line with this situation, we will assume τ is unknown.

If both θ_0 and θ_1 are known, then the SPC diagnostic of choice is generally a cumulative sum (**cusum**) chart, with its minimum time to detection

property. The change-point formulation is therefore of interest in settings where at least one of these parameters is unknown. As it is rare for the out-of-control distribution to be known when the in-control one is not, we ignore this possibility, and consider the following situations:

- θ_0 known, θ_1 unknown; and
- both θ_0 and θ_1 unknown.

The approach to testing for a change point is the same; the log-likelihood of the sequence $X_1, \dots, X_n$ under the change-point model is the sum of two terms: one for the in-control and one for the out-of-control portions of the sequence:

$$\log L(\theta_0, \theta_1, \tau) = A(\theta_0, \tau) + B(\theta_1, \tau, n) \quad (1)$$

where $A(\theta_0, \tau) = \sum_{i=1}^{\tau} \log f(X_i, \theta_0)$ is the in-control portion, and $B(\theta_1, \tau, n) = \sum_{i=\tau+1}^n \log f(X_i, \theta_1)$ is the out-of-control one. If any of the parameters are unknown, then they are estimated by maximizing the likelihood. So if, for example θ_1 is unknown, as it commonly is, then it is replaced by the value $\hat{\theta}_1$ maximizing the $B(\theta_1, \tau, n)$ term, and similarly for θ_0 . The change point τ is estimated as the value k maximizing the sum $A(\theta_0, k) + B(\hat{\theta}_1, k, n)$ if θ_0 is known, and $A(\hat{\theta}_0, k) + B(\hat{\theta}_1, k, n)$ if it is not. If there is no change point, then the log-likelihood is $A(\theta_0, n)$. Thus the conventional generalized likelihood ratio (GLR) test for the presence of a change point is $2[\log L(\hat{\theta}_0, \hat{\theta}_1, \hat{\tau}) - A(\hat{\theta}_0, n)]$ if θ_0 is unknown, and $2[\log L(\theta_0, \hat{\theta}_1, \hat{\tau}) - A(\theta_0, n)]$ if it is known.

As with many other likelihood-ratio methods, this test can often be reexpressed in different but functionally equivalent forms.

Normally Distributed Data

The easiest and most familiar example is normally distributed data. Switching to a more standard notation, suppose that the preshift in-control distribution of the data is $N(\mu_0, \sigma_0^2)$ and the postshift out-of-control distribution is $N(\mu_1, \sigma_1^2)$. It is most commonly assumed that the mean changes following the shift, i.e. $\mu_0 \neq \mu_1$. One family of change-point formulations, the homoscedastic models, assumes that $\sigma_0 = \sigma_1 = \sigma$ say, and the other, the heteroscedastic model, allows $\sigma_0 \neq \sigma_1$.

2 Change-Point Methods

Write $\bar{X}_{ik} = \sum_{j=i+1}^k X_j / (k - i)$ for the mean of the sequence $X_{i+1}, \dots, X_k$ and $V_{ik} = \sum_{j=i+1}^k (X_j - \bar{X}_{ik})^2$ for the sum-of-squared deviations of these observations about their mean. Then splitting at a trial value $\tau = k$, gives maximum-likelihood estimators (MLEs) $\hat{\mu}_0 = \bar{X}_{0k}$, $\hat{\mu}_1 = \bar{X}_{kn}$ for the means, and for the unequal variance setting, $\hat{\sigma}_0^2 = V_{0k} / (k - 1)$ and $\hat{\sigma}_1^2 = V_{kn} / (n - k - 1)$ for the separate variances, or in the homoscedastic case, $\hat{\sigma}^2 = (V_{0k} + V_{kn}) / (n - 2)$.

These variance estimators are not, strictly speaking, MLEs, as they have been scaled by their **degrees of freedom**, and not by the sample sizes on which they are based, but this situation is standard.

Homoscedastic Case

In the homoscedastic case, the estimator of the change point τ is the value of k that maximizes, in absolute value,

$$U_{kn} = \sqrt{\frac{k(n-k)}{n}} (\bar{X}_{0k} - \bar{X}_{kn}) \quad (2)$$

or equivalently, minimizing $V_{0k} + V_{kn}$.

Write $U_{\max,n}$ for this maximized absolute value. The test statistic for a change point at an unknown time τ is then $U_{\max,n} / \sigma$ if σ is somehow known, and in the more common situation that it is not known, is $T_{\max,n} = U_{\max,n} / \hat{\sigma}$. This is just the absolute maximum across k of the two-sample t -test comparing the “samples” $X_1, \dots, X_k$ and $X_{k+1}, \dots, X_n$.

The test can also be expressed as the maximum of the F ratio in an analysis of variance comparing these two samples as:

$$\begin{aligned} F_{\max,n} &= \max_k \left[\frac{V_{0n} - V_{0k} - V_{kn}}{V_{0k} + V_{kn}} (n - 2) \right] \\ &= T_{\max,n}^2 \end{aligned} \quad (3)$$

In the setting where μ_0 is already known exactly, U_{kn} is redefined as $U_{kn} = \sqrt{n - k} (\mu_0 - \bar{X}_{kn})$ and there is a corresponding minor adjustment in the pooled standard deviation if σ has to be estimated, but the broad thrust of the procedure is unchanged. See related publications such as [1–4].

Heteroscedastic Case

There are two interesting change-point tests to perform in the heteroscedastic setting. One is to test for a change in the variance without regard to whether the mean changed when the special cause intervened. The other is to test the possibility of a change in the mean, the variance, or both. Likelihood-ratio test (LRT) statistics for these possibilities start out with:

Test variance only [5]:

$$G_{kn} = \frac{\left[(k-1) \log \left(\frac{\hat{\sigma}_0^2}{\hat{\sigma}_1^2} \right) + (n-k-1) \log \left(\frac{\hat{\sigma}_1^2}{\hat{\sigma}_0^2} \right) \right]}{B} \quad (4)$$

with

$$B = 1 + \frac{[(k-1)^{-1} + (n-k-1)^{-1} - (n-2)^{-1}]}{3} \quad (5)$$

Test mean and/or variance shift [4]:

$$G_{kn} = \frac{\left[k \log \left\{ \frac{k V_{0n}}{n V_{0k}} \right\} + (n-k) \log \left\{ \frac{(n-k) V_{0n}}{n V_{kn}} \right\} \right]}{B} \quad (6)$$

with

$$\begin{aligned} B &= 1 + \frac{11}{12} [k^{-1} + (n-k)^{-1} - n^{-1}] \\ &\quad + [k^{-2} + (n-k)^{-2} - n^{-2}] \end{aligned} \quad (7)$$

The constants B are Bartlett correction factors intended to improve the approximation of the distribution of these statistics to χ^2 distributions with degrees of freedom 1 and 2 respectively for fixed k and n . As with the shift-in-mean problem, the unknown time of change τ is estimated by the k value that maximizes the test statistic.

Theoretical Results

For a given k , U_{kn} / σ follows a standard **normal distribution**, and T_{kn} follows a t distribution with $n - 2$ degrees of freedom. Maximization of either statistic over k however takes the distributions out of the realm of these familiar standard distributions. There is a vast literature on the fixed-sample homoscedastic mean change-point problem where the sample size

n is fixed (see for example [6] and [7] where the technical difficulties are discussed in detail). This situation of a fixed sample size is relevant in a quality setting to the analysis of phase I data. Specific applications to SPC are discussed in [8].

Phase II SPC

In phase II SPC, the sample size is not fixed. Rather, while the process is thought to be in control, additional readings are added to the record and incorporated into the **control charts**. As each new reading is obtained, a decision has to be made of whether the process still appears to be in control. If so, then the reading is incorporated into the history and the process is allowed to continue. If, however, the reading gives rise to a suspicion that the process is now out of control, then the process is stopped for diagnosis and remedial action.

Returning to the generic change-point formulation, a purely-change-point-driven procedure can be described as:

- When each new observation accrues, calculate

$$C_n = \max_k 2[\log L(\hat{\theta}_0, \hat{\theta}_1, k) - A(\hat{\theta}_0, n)] \quad (8)$$

- If $C_n < c_n$ then conclude that there is no evidence of a shift.
- If $C_n > c_n$ then conclude that there is evidence of a departure from control.

The sequence c_n is a sequence of control limits chosen to control the false-alarm behavior of the scheme. Finding this sequence c_n does not seem to be amenable to theoretical analysis, and actual implementations of the change-point method have used simulation to find suitable values.

A particular attraction of the change-point method is its postsignal diagnostic capability. Under the model that the process did indeed have a persistent step change from the in-control to the out-of-control parameter value, the k value maximizing the likelihood ratio gives an estimate of τ , the last reading before the shift, and the corresponding $\hat{\theta}_1$ provides a maximum-likelihood estimate of the out-of-control parameter value. Although these estimates maximize likelihood, they do not necessarily have the usual “good” properties of MLEs such as consistency.

The Change-Point Method as Follow-Up

This sketches the use of the change-point method as a self-contained phase II charting methodology. This is not the only way it can be used. For example some other method such as a Shewhart, cusum, or exponentially weighted moving average (EWMA) chart may be used to detect departure from control and following this signal, the change-point calculations may be done in an essentially phase I fashion to estimate the time of the change and the postchange parameters. Relevant articles are [9–11].

In this use, the conventional distributional results of the pure change-point theory are largely inapplicable. Test statistic distributions for example, assume that the data follow some specified (for example common normal) distribution; these distributions do not apply when we condition on the fact that the data look suspicious to some other charting method.

Existing Change-Point Formulations

The change-point framework is general; specific existing applications are as follows:

- normal data, where the mean may undergo a step change [1];
- normal data, where the variance may undergo a step change [5];
- normal data, where the mean and/or variance may undergo a step change [12];
- multivariate normal data, where the mean vector may undergo a step change [13]; and
- regression data, where the dependence of a process variable on some **covariate** may undergo a step change [14].

The framework is however more general than a menu of existing implementations might suggest. There is nothing to prevent, for example the in-control distribution, not one of independent identically distributed readings, but say a **time series**, being modeled by an autoregressive integrated moving average scheme whose coefficients can change when the process goes out of control. This model involves additional computational, but not conceptual, complexity.

Multiple Change Points

The discussion so far has assumed that the process history shows a maximum of one change point. This raises the question of whether attention to multiple change points is needed. There is a case for arguing that it is not – that when the process goes out of control, the process owner should diagnose the situation as soon as the control scheme provides an alert. Following this, the data identified as seeming out of control should be excised from the in-control record, and when charting starts up again, the change-point calculations should be restarted from the last in-control point, and new process readings should be added sequentially.

Comparison with Other Methods

As noted earlier, SPC settings can be divided into those involving transient and persistent special causes. The latter may give rise to either erratic data or some shifted statistical distribution. A further distinction is between settings where the in-control statistical distribution is known exactly (perhaps as a result of an earlier large phase I study) and those in which the in-control parameters are not known exactly.

The standard SPC methods of Shewhart, cusum, and EWMA charting assume that the in-control parameters are known exactly. The more common reality is that they are estimated from finite phase I samples, and recent research has shown that these phase I samples need to be very large if the known-parameter run behavior is to be realistic.

The difficulties caused by these imperfectly known in-control parameters can logically be accommodated in two ways – by adapting the control limits or interpretive rules of the standard charts and by using methodologies that accommodate unknown parameters explicitly. The self-starting approach is an example of the latter, as is the unknown-parameter change-point formulation. They differ in that the self-starting method makes no assumption about the out-of-control distribution of the data, whereas the change-point formulation does. Logically, if the change is persistent and leads to stable behavior, then change-point methods should be the methods of choice; if the change is persistent but leads to erratic behavior, then the unmet assumption of stable behavior should make change-point formulations less effective.

Example: Gold Mine Sampling QC

This example fits the more general scenario where changes in both mean and/or variance are important and can hinder the performance of a process. We refer the reader to [12] for a fuller discussion and details on this example. The objective is to monitor gold mine samplers. Samplers in gold mines chip out rock samples of the “face” where the ore is being extracted and submit them for chemical assay to measure their gold content. As a quality check, supervisors cut out fresh samples at some of the locations already sampled. This gives rise to pairs of samples and of gold content – one by the original sampler and one by the supervisor, an experienced sampler. The log of the ratio of the gold content of the sampler to that of the supervisor has an approximately normal distribution. This distribution should have a zero mean (else the sampler is biased), and a small variance (else the sampler is erratic). Figure 1 shows a sequence of such log ratios for a junior sample.

Change points in both mean and variance are possible and important. A change in mean leads to a bias in mine valuation. An increase in variance may warn of loss of motivation to sample carefully, while a decrease in variance indicates a highly desirable improvement in measuring quality, perhaps as a result of learning skills. We will address these questions using the more general setting of change-point formulation to look for changes in this sampler’s output. Visually, Figure 1 suggests that the variability in the portion up to around observation 40 is higher than that for the remainder of the sequence. The mean, however, appears to center on zero for

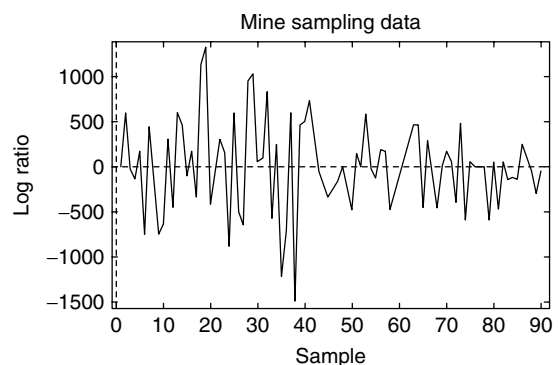


Figure 1 Sample versus log ratio of sampler over supervisor

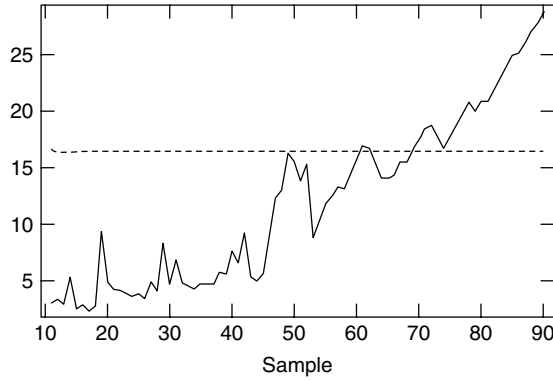


Figure 2 C_n in solid line and $c_{0.002,n}$ in dashed line

Table 1 The gold mine summary statistics

i	$\bar{X}_{\hat{\tau},i}$	$sd_{\hat{\tau},i}$	C_n	t	$\log P$	F	$\log P$
69	-0.035	0.297	16.60	0.85	-0.40	4.74	-4.39
70	-0.027	0.294	17.36	0.79	-0.36	4.83	-4.59
71	-0.024	0.289	18.47	0.77	-0.35	4.99	-4.85
72	-0.037	0.292	18.72	0.88	-0.42	4.90	-4.88
73	-0.020	0.303	17.60	0.72	-0.33	4.57	-4.68
74	-0.037	0.314	16.70	0.87	-0.41	4.24	-4.44
75	-0.034	0.310	17.66	0.86	-0.40	4.36	-4.66
76	-0.033	0.305	18.70	0.85	-0.40	4.50	-4.91
77	-0.032	0.301	19.76	0.85	-0.40	4.63	-5.15
78	-0.032	0.296	20.84	0.85	-0.40	4.77	-5.40
79	-0.047	0.306	19.94	0.98	-0.48	4.47	-5.16
80	-0.044	0.302	20.87	0.96	-0.47	4.58	-5.38
81	-0.055	0.306	20.86	1.06	-0.53	4.47	-5.34
82	-0.052	0.303	21.76	1.03	-0.52	4.57	-5.55
83	-0.054	0.299	22.79	1.06	-0.53	4.67	-5.78
84	-0.056	0.296	23.85	1.08	-0.54	4.79	-6.01
85	-0.057	0.292	24.91	1.10	-0.56	4.89	-6.24
86	-0.050	0.293	25.11	1.03	-0.52	4.88	-6.32
87	-0.047	0.290	25.94	1.01	-0.50	4.97	-6.53
88	-0.047	0.287	27.01	1.01	-0.50	5.08	-6.77
89	-0.052	0.286	27.71	1.06	-0.53	5.11	-6.90
90	-0.052	0.283	28.79	1.06	-0.53	5.22	-7.15

the whole sequence. Figure 2 gives the chart of the GLR along with the control limit $c_{n,0.002}$, where $c_{n,0.002}$ is chosen to yield an in-control average run length (ARL) of 500 for normal data. The chart first crosses the control limit at reading 70, and continues upward except for a brief dip below at reading 74. Right from the first hint of a change point, the GLR gives the left-segment estimates $\hat{\tau} = 42$, $\hat{\mu}_1 =$

0.053, $\hat{\sigma}_1 = 0.42$. Table 1 lists the summary statistics of the right segment parameters, the GLR, the t -, and the F -test statistics for identity of mean and of variance, and the base-10 log of their normal-distribution p values. These confirm that there is no indication of a shift in mean, but that the standard deviation is substantially smaller after the change point than before. By the time of the first signal at sample number 70, the estimated postchange standard deviation has stabilized close to its ultimate value of 0.283 – about two-thirds of its value in the first portion of the series. This reduction of variance is of substantial value to the mine. It means that each sample now produced by this junior sampler is more informative than two of his previous samples were. This has obvious implications for the improvement in mine valuation and selection of which ore to extract.

References

- [1] Hawkins, D.M., Qiu, P. & Kang, C.W. (2003). The change-point model for statistical process control, *Journal of Quality Technology* **35**, 355–365.
- [2] Lai, T.L. (2001). Sequential analysis: some classical problems and new challenges, *Statistica Sinica* **11**, 303–350.
- [3] Siegmund, D. & Venkatraman, E.S. (1994). Using the generalized likelihood ratio statistics for sequential detection of a change point, *Annals of Statistics* **23**, 255–271.
- [4] Worsley, K.J. (1979). On the likelihood ratio test for a shift in location of normal populations, *Journal of American Statistical Association* **74**, 365–367.
- [5] Hawkins, D.M. & Zamba, K.D. (2005). A change point model for a shift in variance, *Journal of Quality Technology* **37**, 21–31.
- [6] Carlstein, E., Mueller, H.G. & Siegmund, D. (1994). *Change-Point Problems*, *IMS Lecture Notes*, Vol. 23, Springer.
- [7] Csörgö, M. & Horváth, L. (1997). *Limit Theorems in Change-Point Analysis*, John Wiley & Sons, New York, Chichester.
- [8] Sullivan, J.H. & Woodall, W.H. (1996). A control chart for preliminary analysis of individual observations, *Journal of Quality Technology* **28**, 265–278.
- [9] Pignatiello Jr, J.J. & Samuel, T.R. (2001). Estimation of the change point of a normal process mean in SPC applications, *Journal of Quality Technology* **33**, 82–95.
- [10] Srivastava, M.S. & Wu, Y. (1999). Quasi-stationary biases of change point and change magnitude estimation after sequential CUSUM test, *Sequential Analysis* **18**, 203–216.

- [11] Wu, Y. (2005). *Inference for Change-Point and Post-Change Means After a CUSUM Test*, *Lecture Notes in Statistics*, Vol. 180, Springer, New York.
- [12] Hawkins, D.M. & Zamba, K.D. (2005). Statistical process control for shifts in mean or variance using a change point formulation, *Technometrics* **47**, 164–173.
- [13] Zamba, K.D. & Hawkins, D.M. (2006). A multivariate change point model for statistical process control, *Technometrics* **48**(4), 539–549.
- [14] Mahmoud, M.A., Parker, P.A., Woodall, W.H. & Hawkins, D.M. (2006). A change point method for linear profile data, *Quality and Reliability Engineering International* **23**, 247–268.

Related Articles

Control Charts for the Mean; Control Charts for the Standard Deviation Control Charts, Overview; Exponentially Weighted Moving Average (EWMA) Control Chart; Cumulative Sum (CUSUM) Chart; Multivariate Control Charts Overview.

DOUGLAS M. HAWKINS, PEIHUA QIU AND
KOKOU D. ZAMBA

Autoregressive Integrated Moving Average (ARIMA) Modeling

Introduction

A time series is the result of observing a magnitude over time (*see Time Series Analysis*). For instance, one can observe the number of defects in a production process every day, the thickness of a coating layer every hour, or the monthly demand for a product. An important class of time series is the stationary one. A time series is stationary if it is stable over time, so that the level of the series is constant, the variability is also constant, and the dependency relationships between successive values of the time series do not change over time.

Many real time series are not stationary but can be transformed to stationarity. Thus a nonstationary series $\{y_t\}$ that wanders over time without a fixed mean may have the property that its growth $x_t := y_t - y_{t-1}$ is a stationary time series. Then we say that the series $\{y_t\}$ is *integrated* and that the series $\{x_t\}$ is the *first difference* of $\{y_t\}$.

ARMA (auto-regressive moving average) models are useful to approximate the dynamics of many stationary time series, whereas ARIMA (autoregressive integrated moving average) models are useful for integrated time series. Figure 1 shows three time series generated by simulation using two ARMA models (a and b) and one ARIMA model (c). Simple inspection of these series shows that they have different dynamics because of their different autocorrelation structures (*see Autocorrelated Data*). The general form of ARMA and ARIMA models will be considered subsequently.

ARMA Models

A process $\{y_t\}$ is called stationary in the weak form if $E(y_t)$ and $\gamma_k := \text{cov}(y_t, y_{t+k})$ for $k = 0, 1, \dots$ are finite and do not depend on t . The sequence $\{\gamma_k\}$ is called the *autocovariance function* and $\rho_k := \gamma_k/\gamma_0$ is the *autocorrelation function* (*see Autocorrelation Function*). Wold [1] proved that any stationary process, $\{y_t\}$, without deterministic components can

be written as

$$\begin{aligned} y_t &= a_t + \psi_1 a_{t-1} + \psi_2 a_{t-2} + \dots \\ &= \sum_{i=0}^{\infty} \psi_i a_{t-i} \quad (\psi_0 = 1) \end{aligned} \quad (1)$$

where $\sum_{i=0}^{\infty} \psi_i^2 < \infty$ (*stationarity condition*) and $\{a_t\}$ is a zero-mean white noise process ($E(a_t a_{t+k}) = 0$ for $k \neq 0$) with variance $\sigma^2 := E(a_t^2)$. This is called the general, or the infinite **moving average** MA(∞), representation of a stationary process. Model (1) leads to $E(y_t) = 0$. If the process $\{y_t\}$ has deterministic mean $\{\mu_t\}$, the general representation applies to the zero-mean process $\{y_t - \mu_t\}$. The autocovariance function is $\gamma_k = E(y_t y_{t+k}) = \sigma^2 \sum_{i=0}^{\infty} \psi_i \psi_{i+k}$, for $k \geq 0$ (see also Appendix 1). Using the lag operator B such that $B y_t = y_{t-1}$, the MA(∞) representation can formally be written as $y_t = \psi(B) a_t$, where $\psi(B) = 1 + \psi_1 B + \psi_2 B^2 + \dots$. The white noise $\{a_t\}$ can often be assumed to be normally distributed.

If model (1) is invertible, it can be written in an alternative way that is called the infinite autoregressive AR(∞) representation of the process and is given by

$$y_t - \pi_1 y_{t-1} - \pi_2 y_{t-2} - \dots = y_t - \sum_{i=1}^{\infty} \pi_i y_{t-i} = a_t \quad (2)$$

where $\sum_{i=1}^{\infty} \pi_i^2 < \infty$ (*invertibility condition*). Model (2) can formally be written as $\pi(B) y_t = a_t$, where $\pi(B) = 1 - \pi_1 B - \pi_2 B^2 - \dots$, so that $\pi(B) \psi(B) = 1$.

These two general representations have an infinite number of parameters. Approximations having a finite number of parameters are needed to model any finite time series $\mathbf{Y}' = (y_1, \dots, y_n)$. The simplest approximation is to assume a finite number q of terms in equation (1), which leads to the q -order moving average MA(q) model $y_t = \theta_q(B) a_t$, where $\theta_q(B) = 1 - \theta_1 B - \dots - \theta_q B^q$ (minus signs are used customarily). For MA(q) models, $\gamma_0 = \sigma^2(1 + \theta_1^2 + \dots + \theta_q^2) < \infty$, so that they are always stationary; moreover $\rho_k = 0$ for $k > q$. A second way to approximate the general form is to use a finite number p of terms in equation (2). This leads to the p -order autoregressive AR(p) model $\phi_p(B) y_t = a_t$, where $\phi_p(B) = 1 - \phi_1 B - \dots - \phi_p B^p$. An AR(p) model is *stationary* if and only if all the roots of

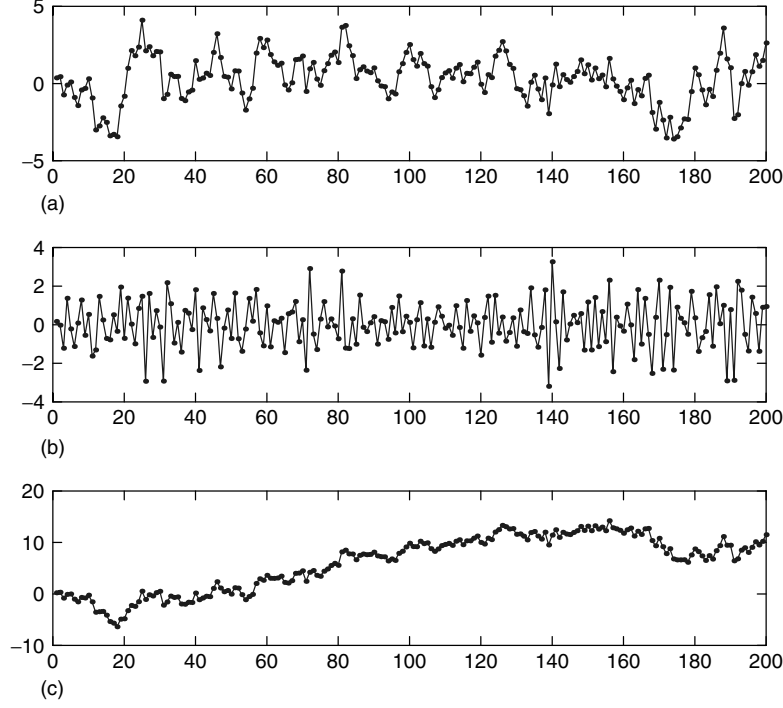


Figure 1 Time series generated using the AR(1) model $y_t - 0.8y_{t-1} = a_t$ (a), the MA(1) model $y_t = a_t - 0.8a_{t-1}$ (b), and the IMA(0, 1, 1) model $y_t - y_{t-1} = a_t - 0.2a_{t-1}$ (c), for a common simulated white noise series $(a_1, \dots, a_{200})$ with zero mean and variance $\sigma^2 = 1$

the polynomial $\phi_p(B)$ are outside the unit circle, so that the model accepts the MA(∞) representation $y_t = \phi_p^{-1}(B)a_t$. Thus the AR(1) model $(1 - \phi B)y_t = a_t$ can be written as $y_t = (1 - \phi B)^{-1}a_t = (1 + \phi B + \phi^2 B^2 + \dots)a_t$ if and only if $|\phi| < 1$, so that $\gamma_0 = \sigma^2(1 - \phi^2)^{-1} < \infty$; the autocorrelation function is $\rho_k = \phi^k$ implying a geometrical decay of $\rho_k \neq 0$, which is a feature common to all stationary AR(p) models.

The AR(p) and MA(q) models can be combined to form the ARMA(p, q) model $\phi_p(B)y_t = \theta_q(B)a_t$. ARMA(p, q) models are *stationary* if and only if all the roots of $\phi_p(B)$ are outside the unit circle, and they are *invertible* if and only if all the roots of $\theta_q(B)$ are outside the unit circle. Only invertible models yield time series *forecasts* directly (see the section titled “Forecasting Time Series”).

ARIMA Models

ARMA(p, q) models can be extended to model the dynamics of integrated time series. A nonstationary

process $\{y_t\}$ whose d th difference $\{(1 - B)^d y_t\}$ is stationary is called an *integrated* process of order $d > 0$. Using the operator $\nabla = 1 - B$, the general ARIMA(p, d, q) equation is

$$\phi_p(B)\nabla^d(y_t - \mu_t) = \theta_q(B)a_t \quad (3)$$

where $\phi_p(B)$ and $\theta_q(B)$ have no common roots and are invertible (i.e., all their roots are outside the unit circle), and $\{\mu_t\}$ is a deterministic trend. In model (3), $\{\nabla^d \mu_t\}$ is the mean of the differenced series $\{\nabla^d y_t\}$ and is estimable, but the d integrating constants in $\{\mu_t\}$ cannot be estimated. Important examples of this type of models are the random walk, $y_t = y_{t-1} + a_t$, also called ARIMA(0, 1, 0) or I(1), and the IMA(0, 1, 1) model $y_t = y_{t-1} + a_t - \theta a_{t-1}$.

Under model (3), $\{\nabla^d(y_t - \mu_t)\}$ is a zero-mean stationary process. The mean and variance of y_t conditioned on given past values $\{y_0, y_{-1}, \dots\}$ are finite. Thus for a random walk $y_t = y_{t-1} + a_t$ and $t > 0$, one has $E(y_t|y_0) = y_0$ whereas $\text{var}(y_t|y_0) = \sigma^2 t$ increases indefinitely as t increases. Also, for $k > 0$,

$\text{cov}(y_t, y_{t+k}) = \sigma^2 t \geq 0$, so that the autocorrelation coefficient $\rho(y_t, y_{t+k}) = (1 + k/t)^{-1/2}$ decreases with k , which is a property of integrated processes. The function $V(k) = \frac{1}{2} E(y_k - y_0)^2 = \gamma_0(1 - \rho_k)$ is called the *variogram*. The quotient $V(k)/V(1)$ is called [2] the *standardized variogram* (SV) and is a dimensionless function of k . The SV of ARMA models is $(1 - \rho_k)/(1 - \rho_1)$ and is always bounded. Unlike the autocovariance function, the SV of ARIMA($p, 1, q$) models always exists and equals $k + 2 \sum_{i=1}^{k-1} (k-i) \text{corr}(\nabla y_t, \nabla y_{t-i})$; for the IMA(0, 1, 1) model, it is $1 + (k-1)(1 - \theta)^2/(1 + \theta^2)$ and increases linearly with k .

Applications

ARIMA models are important for generating forecasts and providing understanding in all kinds of time series problems from economics to health care applications [2–9]. In particular, in quality and reliability, they are important in process monitoring if observations are correlated, designing schemes for process adjustment, monitoring a reliability system over time, forecasting time series, estimating missing values, finding **outliers** and atypical events, understanding the effects of changes in a system, and so on.

Forecasting Time Series

ARIMA models can be used to forecast time series. Forecasting with ARIMA models is very simple when $q = 0$. Using the notation $\phi_p(B)\nabla^d = 1 - \varphi_1 B - \dots - \varphi_{p+d} B^{p+d}$ and assuming, for simplicity, that $\mu_t = 0$ in equation (3), the minimum mean squared error (MMSE) forecast of y_t given $(y_{t-1}, \dots, y_{t-p-d})$ is $\hat{y}_t = \varphi_1 y_{t-1} + \dots + \varphi_{p+d} y_{t-p-d}$. The forecast error is a_t . When $q > 0$, however, the forecasts depend on previously observed forecast errors. Thus for given observations $(y_{t-1}, \dots, y_{t-p-d})$, and computed forecast errors $(a_{t-1}, \dots, a_{t-q})$, the MMSE forecast of y_t is

$$\begin{aligned} \hat{y}_t = & \varphi_1 y_{t-1} + \dots + \varphi_{p+d} y_{t-p-d} \\ & - \theta_1 a_{t-1} - \dots - \theta_q a_{t-q} \end{aligned} \quad (4)$$

and the forecast error is a_t . The quick recursion (4) is numerically stable when used iteratively if and only if $\theta_q(B)$ is *invertible*.

For a finite series $\mathbf{Y}' = (y_1, \dots, y_n)$, equation (4) is exact only asymptotically and has to be initialized. Suppose a series $\mathbf{Y}$ having zero mean and **covariance** matrix $\mathbf{\Omega} := \text{cov}(\mathbf{Y})$ has to be forecast. This problem can be formulated as successively finding uncorrelated (orthogonal) forecast errors $\mathbf{E}' = (e_1, \dots, e_n)$ with zero mean and covariance matrix $\mathbf{D} := \text{cov}(\mathbf{E}) = \text{diag}(D_{11}, \dots, D_{nn})$. This can be solved by applying Gram–Schmidt orthogonalization to $\mathbf{Y}$. Then $\hat{\mathbf{Y}} = \mathbf{Y} - \mathbf{E}$. Starting with $e_1 = y_1$ (so that $D_{11} = \Omega_{11}$ and $\hat{y}_1 = 0$) and using $e_t = y_t - \hat{y}_t = y_t - \sum_{j=1}^{t-1} \lambda_{tj} e_j$, one gets $\text{cov}(y_t, e_i) = \text{cov}\left(e_t + \sum_{j=1}^{t-1} \lambda_{tj} e_j, e_i\right) = \lambda_{ti} \text{var}(e_i)$ and $\text{cov}(y_t, e_t) = \text{var}(e_t) = D_{tt}$. Then $\lambda_{ti} = \text{cov}(y_t, e_i)/\text{cov}(y_i, e_i)$. The required covariances are found using the recursion

$$\begin{aligned} \text{cov}(y_t, e_i) &= \text{cov}\left(y_t, y_i - \sum_{j=1}^{i-1} \lambda_{ij} e_j\right) \\ &= \Omega_{ti} - \sum_{j=1}^{i-1} \lambda_{ij} \text{cov}(y_t, e_j) \end{aligned} \quad (5)$$

for $t = 1, \dots, n$ and $i = 1, \dots, t$. Hence $\mathbf{Y} = \mathbf{L}_{\Omega} \mathbf{E}$, where $\mathbf{L}_{\Omega}$ is the lower triangular matrix with elements λ_{tj} for $t = 1, \dots, n$ and $j = 1, \dots, t$ (note that $\lambda_{tt} = 1$), and $\mathbf{\Omega} = \mathbf{L}_{\Omega} \mathbf{D} \mathbf{L}_{\Omega}'$.

If the series $\mathbf{Y}$ is generated by the ARMA(p, q) model $\phi_p(B)(y_t - \mu_t) = \theta_q(B)a_t$, $\mathbf{\Omega}$ has elements $\Omega_{ij} = \gamma_{|i-j|}$. Then, it is advantageous to apply the recursion (5) to the working series $\mathbf{W}' = (w_1, \dots, w_n)$ defined by $w_t = y_t - \mu_t$ for $t = 1, \dots, p$ and $w_t = \phi_p(B)(y_t - \mu_t)$ for $t = p+1, \dots, n$, because $\mathbf{\Omega}$ is replaced by the *band* matrix $\mathbf{\Theta} := \text{cov}(\mathbf{W})$ (see Appendix 2). Then $\mathbf{W} = \mathbf{L}_{\Theta} \mathbf{E}$; $\mathbf{E} = \mathbf{W} - \hat{\mathbf{W}} = \mathbf{Y} - \hat{\mathbf{Y}}$; and $\mathbf{\Theta} = \mathbf{L}_{\Theta} \mathbf{D} \mathbf{L}_{\Theta}'$, where $\mathbf{L}_{\Theta}$ is a lower triangular band matrix with the same lower bandwidth as $\mathbf{\Theta}$. The forecast errors $\{e_t\}$ converge to the noise $\{a_t\}$ and $\lambda_{t,t-j}$ converges to $-\theta_j$ as t increases, so that one can switch from equation (5) to equation (4) as soon as equation (4) becomes sufficiently accurate. Note also that $\mathbf{W} = \mathbf{\Phi}(\mathbf{Y} - \boldsymbol{\mu})$, where $\mathbf{\Phi}$ is the lower triangular band matrix with bandwidth p having ones in the diagonal; elements $\Phi_{i,i-k} = -\phi_k$ for $i = p+1, \dots, n$ and $k = 1, \dots, p$; and zeros elsewhere. Then $\mathbf{\Theta} = \mathbf{\Phi} \mathbf{\Omega} \mathbf{\Phi}'$ and $\mathbf{L}_{\Theta} = \mathbf{\Phi} \mathbf{L}_{\Omega}$.

This method can also be applied to ARIMA models provided that limiting formulas are used.

For example, for ARMA(p, q) models with constant mean μ , the MMSE forecast of y_2 given y_1 is $\hat{y}_2 = \mu + \rho_1(y_1 - \mu)$ and the variance of the forecast error is $D_{22} = \gamma_0(1 - \rho_1^2)$. Hence, for ARIMA($p, 1, q$) models, $\hat{y}_2 = y_1$ with forecast error variance $\lim_{\rho_1 \rightarrow 1} D_{22} = \text{var}(\nabla y_1)$.

MMSE forecasts $\hat{y}_{t|t-k}$ of y_t at origin $t - k$ for any lead time $k > 1$ are found replacing equation (4) by $\hat{y}_{t|t-k} = \sum_{i=1}^{k-1} \varphi_i \hat{y}_{t-i|t-k} + \sum_{i=k}^{p+d} \varphi_i y_{t-i} - \sum_{i=k}^q \theta_i a_{t-i}$; the forecast error is $e_{t|t-k} = y_t - \hat{y}_{t|t-k}$ and $\text{var}(e_{t|t-k}) = \sigma^2 \sum_{i=0}^{k-1} \psi_i^2$. Equation (5) can also be adapted to forecast y_t at origin $t - k$ by taking $\mathbf{Y}' = (y_1, \dots, y_{t-k}, y_t)$. Thus $\hat{y}_{t|t-k} = \sum_{j=1}^{t-k} \lambda_{tj} e_j$ and $\text{var}(e_{t|t-k}) = \sum_{j=t-k+1}^t \lambda_{tj}^2 D_{jj}$. For ARMA(p, q) models with constant μ , $\hat{y}_{k+1|1} = \mu + \rho_k(y_1 - \mu)$ and $\text{var}(e_{k+1|1}) = \gamma_0(1 - \rho_k^2)$. For ARIMA($p, 1, q$) models, $\hat{y}_{k+1|1} = y_1$ and $\text{var}(e_{k+1|1}) = k\text{var}(\nabla y_1) + 2 \sum_{i=1}^{k-1} (k-i)\text{cov}(\nabla y_1, \nabla y_{1-i})$.

Figure 2 shows the last 11 observations in the time series of Figure 1 together with the MMSE forecasts $\hat{y}_{t+k|t}$ of y_{t+k} at origin $t = 200$ for lead times $k = 1, \dots, 5$ and corresponding probability limits for each individual y_{t+k} at ± 2 standard deviations.

Process Control: Process Monitoring and Process Adjustment

Quality control is often implemented by the complementary use of process monitoring and process adjustment [2, 5]. *Process monitoring* is a part of statistical process control (SPC) and assumes that the process has already been brought to a state of stationary control by using techniques such as the standardization of materials, machines, and methods. The purpose of monitoring is to *guarantee* that the process stays in its stationary state. Hence stationary time series models are used. *Process adjustment* is primarily used when current efforts to bring the process to a state of stationary control have failed. The purpose is to determine when and by how much the process should be adjusted to attain stationarity. This can be implemented by sequentially forecasting the deviation from the target and compensating for this deviation by adjusting an input variable when and as required (see **Feedback Control**). Nonstationary and stationary time series models are used to describe the dynamics of the unadjusted and adjusted processes, respectively.

Process monitoring and process adjustment require that data should be taken routinely at equally spaced

periods of time separated by the so-called unit sampling interval. The determination of the optimal length of this sampling interval is important. It may also be appropriate to change the length of the sampling interval depending on the recent behavior of the process. Hence it is useful that the models used for process control are *invariant* to changes in the length of the sampling interval (only the model parameters may change).

In *process monitoring*, stationary ARMA models that are invariant to the length of the sampling interval are used. The basic type of these models is the white noise $y_t = \mu + a_t$, where the observed process $\{y_t\}$ has mean μ , which should be equal to the target value, and the white noise series $\{a_t\}$ is a disturbance. Alternative invariant models are the ARMA(1, 1) $y_t - \mu = \phi(y_{t-1} - \mu) + a_t - \theta a_{t-1}$ with $0 < \phi < 1$, and the AR(1) with $\theta = 0$.

In *process adjustment*, nonstationary ARIMA models that are invariant to the length of the sampling interval are used. The basic type of these models is the random walk $y_t = y_{t-1} + a_t$, used by Taguchi *et al.* [10] to model drifting observed processes. Box and Luceño [2] consider that the random walk is often too nonstationary for process adjustment applications and use the IMA(0, 1, 1) model instead. Luceño *et al.* [11] and [5] use the IMA(0, 1, 1) with trend $y_t = \beta + y_{t-1} + a_t - \theta a_{t-1}$, which is also invariant and has parameters θ and β to deal with different degrees of nonstationarity and the trend of unadjusted processes.

Fitting ARIMA Models

The approach suggested by Box and Jenkins [3, 12] is often used to fit ARIMA models to finite time series. It has three steps: model identification, model estimation, and model diagnostic.

Model Identification

Suppose that a series $\mathbf{Y}' = (y_1, \dots, y_n)$ has been observed. The first step is to decide whether it is a stationary or an integrated time series. This decision is often made by visual inspection of plots of the time series and its autocorrelation function, as well as similar plots for the first and second difference of the series. It is also possible to use unit root tests of the type introduced in [13]. The following estimators are applied

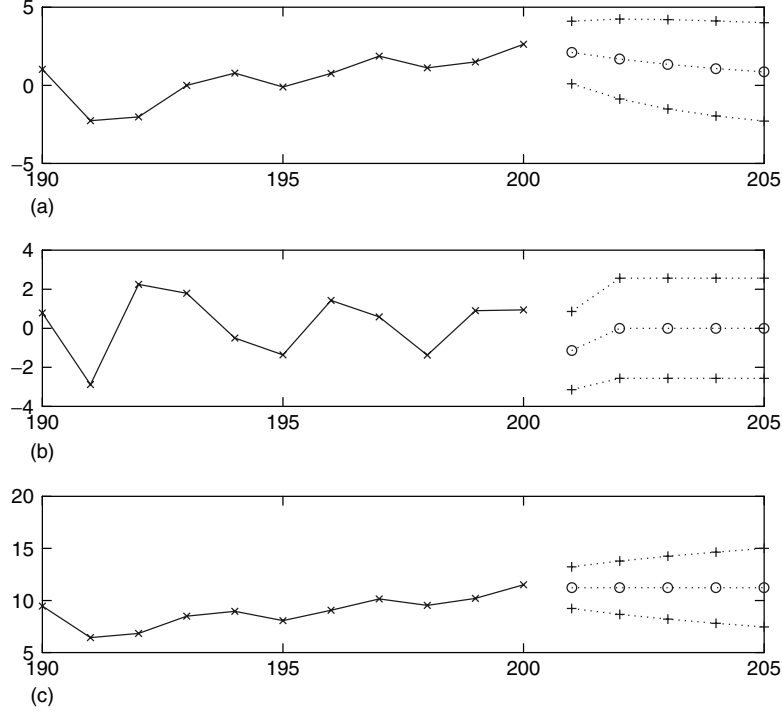


Figure 2 Observations in the time series of Figure 1 for $t = 190, \dots, 200$ ($\times$) along with their MMSE forecasts for $t = 201, \dots, 205$ ($\circ$) with ± 2 standard deviation limits ($+$)

to the series itself and its first differences: $\bar{y} = n^{-1} \sum_{t=1}^n y_t$ for the mean, $\hat{\gamma}_0 = n^{-1} \sum_{t=1}^n (y_t - \bar{y})^2$ for the variance,

$$\hat{\gamma}_k = \frac{1}{n} \sum_{t=1}^{n-k} (y_t - \bar{y})(y_{t+k} - \bar{y}) \quad (k \geq 0) \quad (6)$$

for the autocovariances, and $r_k = \hat{\gamma}_k / \hat{\gamma}_0$ for the autocorrelations. If the sample autocorrelation function of the original or differenced series has only its first q coefficients significantly different from zero, one can tentatively assume an $MA(q)$ model for that series. Otherwise, one should entertain AR, ARMA, or ARIMA models. Figure 3 shows the sample autocorrelation functions of the three time series in Figure 1. The sequence $\{r_k\}$ decreases approximately geometrically for the AR(1) series; only r_1 is clearly significantly different from zero for the MA(1) series; whereas the sequence $\{r_k\}$ decreases very slowly for the IMA(0, 1, 1) series, indicating the need to differentiate this time series.

Many statistical tools have been derived to identify the orders p , d , and q of a model: the partial autocorrelation function [3], the extended autocorrelation function [14], the inverse autocorrelation function [15, 16], and so on. However, because the estimation of ARMA models is nowadays fast and reliable, we recommend that all possible models compatible with the observed autocorrelations be estimated and that a model selection criterion be used to choose from among them (see the sections titled “Model Estimation by Maximum Likelihood”, “Joint Estimation of Missing Values and Model Parameters”, and “Model Diagnostic”).

Model Estimation by Maximum Likelihood

Let $\mathbf{Y}' = (y_1, \dots, y_n)$ be a series generated by a stationary ARMA(p, q) model $\phi_p(B)(y_t - \mu_t) = \theta_q(B)a_t$, so that $\boldsymbol{\mu}' := E(\mathbf{Y}') = (\mu_1, \dots, \mu_n)$ and $\boldsymbol{\Omega} := \text{cov}(\mathbf{Y})$. Let $\mathbf{W} := \Phi(\mathbf{Y} - \boldsymbol{\mu})$ be the working series of the section titled “Forecasting Time Series”, so that $\boldsymbol{\Theta} := \text{cov}(\mathbf{W}) = \Phi \boldsymbol{\Omega} \Phi'$, $\mathbf{W} = \mathbf{L}_{\boldsymbol{\Theta}} \mathbf{E}$,

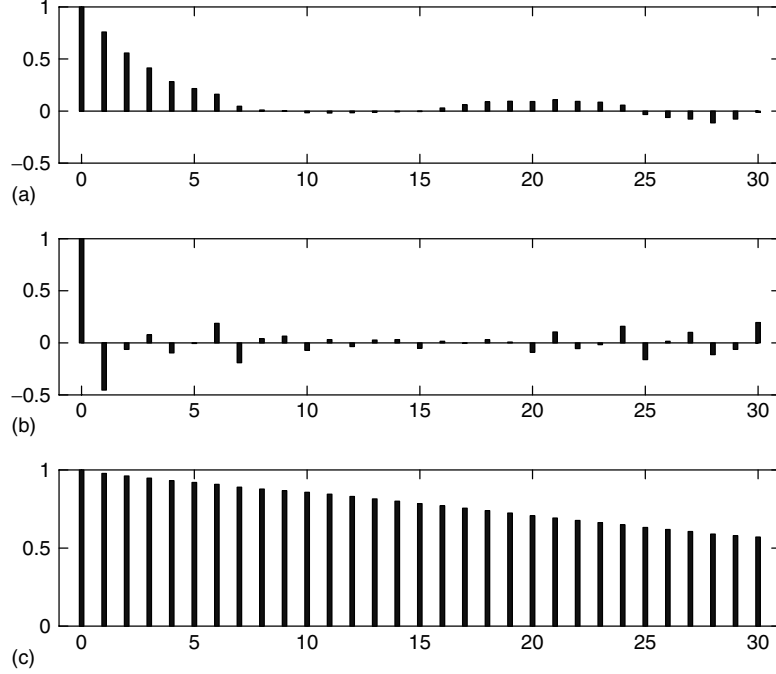


Figure 3 Sample autocorrelation functions for each of the three time series in Figure 1

$\Theta = \mathbf{L}_\Theta \mathbf{D} \mathbf{L}_\Theta'$, $|\Omega| = |\Theta| = |\mathbf{D}|$, Φ , Θ , and $\mathbf{L}_\Theta$ are band matrices, and $\mathbf{D}$ is diagonal. For normal $\{a_t\}$, the likelihood function of the parameters $\phi' = (\phi_1, \dots, \phi_p)$, $\theta' = (\theta_1, \dots, \theta_q)$, μ , and σ^2 given $\mathbf{Y}$ is

$$\begin{aligned}
 L(\phi, \theta, \mu, \sigma^2 | \mathbf{Y}) &= (2\pi)^{-n/2} |\Omega|^{-1/2} \\
 &\quad \times \exp\{-(\mathbf{Y} - \mu)' \Omega^{-1} (\mathbf{Y} - \mu)/2\} \\
 &= (2\pi)^{-n/2} |\Theta|^{-1/2} \\
 &\quad \times \exp(-\mathbf{W}' \Theta^{-1} \mathbf{W}/2) \\
 &= (2\pi)^{-n/2} |\mathbf{D}|^{-1/2} \exp(-\mathbf{E}' \mathbf{D}^{-1} \mathbf{E}/2)
 \end{aligned} \tag{7}$$

where $|\mathbf{D}| = \prod_{i=1}^n D_{ii}$ and $\mathbf{E}' \mathbf{D}^{-1} \mathbf{E} = \sum_{i=1}^n e_i^2 / D_{ii}$. Maximum-likelihood estimates (MLEs) can be found by minimizing $-\ln[L(\phi, \theta, \mu(\beta), \sigma^2 | \mathbf{Y})]$, where μ is parameterized as $\mu(\beta)$ (see **Maximum Likelihood**).

For $\text{AR}(p)$ models, Θ equals the identity matrix excepting its first p rows and columns (and Ω^{-1} has bandwidth p); hence the MLE of ϕ' is close to its minimum sum of squared error estimate. For an $\text{AR}(1)$ model, $|\Theta|^{-1} = \sigma^{-2n} (1 -$

$\phi^2)$ and $\mathbf{W}' \Theta^{-1} \mathbf{W} \sigma^2 = C_0(1 + \phi^2) - 2\phi C_1 - \phi^2 \{(y_1 - \mu_1)^2 + (y_n - \mu_n)^2\}$, where $C_k = \sum_{t=1}^{n-k} (y_t - \mu_t)(y_{t+k} - \mu_{t+k})$; hence disregarding the effects of $|\Theta|$, y_1 , and y_n , the estimate $\hat{\phi} = C_1/C_0$ results. For $\text{MA}(q)$ models, Φ is the identity matrix, $\Omega = \Theta$ has bandwidth q , but the log-likelihood function may be far from quadratic. The MLEs for the parameters of the three time series in Figure 1 are $(\hat{\phi} = 0.765, \hat{\sigma}^2 = 0.927)$, $(\hat{\theta} = 0.838, \hat{\sigma}^2 = 0.925)$, and $(\hat{\theta} = 0.197, \hat{\sigma}^2 = 0.934)$, respectively.

Joint Estimation of Missing Values and Model Parameters

Suppose that series $\mathbf{Y}' = (y_1, \dots, y_n)$ has m missing values (see **Missing Data and Imputation**) $\mathbf{y}'_u = (y_{u_1}, \dots, y_{u_m})$, at positions $1 < u_1 < \dots < u_m < n$, having means $\mu'_u = (\mu_{u_1}, \dots, \mu_{u_m})$. Let $\mathbf{M}'$ be the $m \times n$ matrix with elements $M'_{i, u_i} = 1$, $i = 1, \dots, m$, and zeros elsewhere; and $\mathbf{B}'$ be the $(n - m) \times n$ matrix found eliminating the rows of $\mathbf{M}'$ from the $n \times n$ identity matrix. The MMSE estimator $\hat{\mathbf{y}}_u$ of the missing values is the conditional mean of $\mathbf{y}_u$ given the observations $\mathbf{B}'\mathbf{Y}$,

that is,

$$\begin{aligned}\hat{\mathbf{y}}_u &= \boldsymbol{\mu}_u + (\mathbf{M}'\boldsymbol{\Omega}\mathbf{B})(\mathbf{B}'\boldsymbol{\Omega}\mathbf{B})^{-1}\mathbf{B}'(\mathbf{Y} - \boldsymbol{\mu}) \\ &= \boldsymbol{\mu}_u - (\mathbf{M}'\boldsymbol{\Omega}^{-1}\mathbf{M})^{-1}(\mathbf{M}'\boldsymbol{\Omega}^{-1}\mathbf{B})\mathbf{B}'(\mathbf{Y} - \boldsymbol{\mu})\end{aligned}\quad (8)$$

The two formulas in equation (8) are equivalent, but the latter (using $\boldsymbol{\Omega}^{-1}$) is usually much faster than the former (using $\boldsymbol{\Omega}$) as $m \ll n - m$ [17–19]. Moreover, $\text{cov}(\mathbf{y}_u) = \mathbf{M}'\boldsymbol{\Omega}\mathbf{M}$, $\text{cov}(\mathbf{y}_u - \hat{\mathbf{y}}_u) = (\mathbf{M}'\boldsymbol{\Omega}^{-1}\mathbf{M})^{-1}$, and $\text{cov}(\hat{\mathbf{y}}_u) = \mathbf{M}'\boldsymbol{\Omega}\mathbf{M} - (\mathbf{M}'\boldsymbol{\Omega}^{-1}\mathbf{M})^{-1}$. For AR(p) models with $\boldsymbol{\mu} = 0$, equation (8) yields $\hat{y}_u = -\sum_{k=1}^p \alpha_k(y_{u-k} + y_{u+k})/\alpha_0$, if $m = 1$ and $p < u \leq n - p$, where $\alpha_k = \sum_{j=0}^{p-k} \phi_j \phi_{j+k}$ and $\phi_0 = -1$. The sequence $\{\alpha_k/\alpha_0\}$ is called the *inverse autocorrelation function* of the AR(p) model. Moreover, equation (8) accounts for end effects if $u \leq p$ or $u > n - p$.

Defining $\hat{\mathbf{Y}} = \mathbf{Y} + \mathbf{M}(\hat{\mathbf{y}}_u - \mathbf{y}_u)$ and using $\boldsymbol{\Omega}^{-1} = \boldsymbol{\Phi}'\boldsymbol{\Theta}^{-1}\boldsymbol{\Phi} = \boldsymbol{\Phi}'\mathbf{L}_{\Theta}^{-1}\mathbf{D}^{-1}\mathbf{L}_{\Theta}\boldsymbol{\Phi}$, one can efficiently obtain [17–20] MLEs of $\boldsymbol{\phi}$, $\boldsymbol{\theta}$, $\boldsymbol{\mu}$, and σ^2 maximizing

$$\begin{aligned}L(\boldsymbol{\phi}, \boldsymbol{\theta}, \boldsymbol{\mu}, \sigma^2 | \mathbf{B}'\mathbf{Y}) \\ &= (2\pi)^{-(n-m)/2} (|\boldsymbol{\Omega}| |\mathbf{M}'\boldsymbol{\Omega}^{-1}\mathbf{M}|)^{-1/2} \\ &\quad \times \exp\{-(\hat{\mathbf{Y}} - \boldsymbol{\mu})'\boldsymbol{\Omega}^{-1}(\hat{\mathbf{Y}} - \boldsymbol{\mu})/2\}\end{aligned}\quad (9)$$

This procedure can also be applied to integrated series replacing $(\mathbf{M}'\boldsymbol{\Omega}^{-1}\mathbf{M})^{-1}(\mathbf{M}'\boldsymbol{\Omega}^{-1}\mathbf{B})$ by $(\mathbf{M}'_d\boldsymbol{\Omega}_d^{-1}\mathbf{M}_d)^{-1}(\mathbf{M}'_d\boldsymbol{\Omega}_d^{-1}\mathbf{B}_d)$ in equation (8), where $\boldsymbol{\Omega}_d := \text{cov}(\mathbf{Y}_d)$, and $\mathbf{Y}'_d$, $\mathbf{M}'_d$, and $\mathbf{B}'_d$ are obtained differencing d times the rows of $\mathbf{Y}'$, $\mathbf{M}'$, and $\mathbf{B}'$. Thus, for an ARIMA($p, 1, q$), if y_3 is missing, $\mathbf{Y}_1$ has two unknown values ($y_3 - y_2$ and $y_4 - y_3$) that require only one row of length $n - 1$ in $\mathbf{M}'_1$, namely, $(0, 1, -1, 0, \dots, 0)$, the first difference of $(0, 0, 1, 0, 0, \dots, 0)$. Moreover, $\mathbf{B}_1\mathbf{B}'\mathbf{Y} = (y_2 - y_1, -y_2, y_4, y_5 - y_4, \dots)'$. Thus, for a random walk, $\hat{y}_3 = (y_2 + y_4)/2$.

Model Diagnostic

This stage is designed to check if the **residuals** of the estimated models behave approximately as white noise series and to select the best models among those verifying this condition. If a model is correct, the autocorrelation coefficients $\hat{r}_k$ of its residuals $(\hat{a}_1, \dots, \hat{a}_n)$ should asymptotically be normal variables with zero mean and variance $1/n$. Ljung and Box [21] propose the test statistic $Q(h) = n(n+2)$

$\sum_{k=1}^h \hat{r}_k^2/(n-k)$, which, for large n , follows a χ^2 distribution with $h - \nu$ **degrees of freedom**, where ν is the number of parameters $\{\boldsymbol{\phi}, \boldsymbol{\theta}, \boldsymbol{\mu}\}$ estimated in equation (7).

Peña and Rodriguez [22] propose a test statistic that can be up to 50% more powerful than the one proposed by Ljung and Box. The new statistic is $D_h = -n(h+1)^{-1} \ln |\hat{\mathbf{R}}_h|$, where $\hat{\mathbf{R}}_h$ is an $h \times h$ matrix of residual autocorrelations with elements $\hat{r}_{|i-j|}$. The asymptotic distribution of D_h is gamma with mean $h/2 - \nu$ and variance $h(2h+1)/(3h+3) - 2\nu$.

The selection among the fitted models whose residuals behave approximately as white noise series can be based on *model selection criteria*. Akaike [23, 24] proposes the selection of the model that leads to the smallest out-of-sample forecast error and shows that this is equivalent to minimizing $AIC = n \ln \hat{\sigma}_{MV}^2 + 2\nu$, where $\hat{\sigma}_{MV}^2$ is the MLE of the residual variance. Using a Bayesian approach, Schwarz [25] proposes a criterion that selects the model with highest posterior probability by minimizing $BIC = n \ln \hat{\sigma}_{MV}^2 + \nu \ln n$. See [26] for a comparison of the performance of these criteria for model selection.

Atypical Values

Time series data are often subject to outliers or discordant observations (*see Outliers*). Their study has been approached from two points of view. The first is the *diagnostic approach*, in which outliers are identified by using the residuals of the estimated model and their effect is tested afterward. A model incorporating the identified outliers is proposed, and the outlier effects and model parameters are estimated jointly. The second is the *robust approach*, in which the estimation method is modified so that the estimators are insensitive to the presence of outliers. Outliers can be easily identified and tested using these robust estimators. Both methodologies complement each other, and ideas from one approach can be used to improve the other.

Fox [27] defines additive outliers (AOs) and innovative outliers (IOs) in time series and proposes the use of maximum-likelihood ratio tests for detecting them. An AO corresponds to an external error or exogenous change of the observed value of the time series at a particular time point T ; that is, instead of observing the series $\{y_t\}$, a new series $\{z_t\}$ is observed which differs from $\{y_t\}$ by the magnitude

ω_A of the AO only at the time T of occurrence of the outlier. The model for the observed series is $z_t = \omega_A I_t^{(T)} + \Psi(B)a_t$, where $I_t^{(T)} = 0$ for $t \neq T$, $I_T^{(T)} = 1$, and $\Psi(B) = \nabla^{-d} \phi_p^{-1}(B) \theta_q(B)$ is the ARIMA model for the outlier-free time series.

An AO can be interpreted as a measurement error at time T , $1 \leq T \leq n$, or as an impulse effect due to exogenous causes. Thus if the original series is the output from a system, an AO may correspond to an unexpected event that happens at T , such as a strike, an accident, or a breakdown, which modifies the output of the system at T , without further effects on the future values of the series. AOs can have a very serious effect on the properties of the observed series [9]. In SPC, AOs are monitored using Shewhart or CUMulative SCORE (CUSCORE) charts (*see Control Charts, Overview*).

The second type of outliers introduced in [27], IO, can be generated by some internal change or endogenous effect on the time series which appears as an outlier on the process noise. The model for an IO is built by adding an impulse effect to the noise of the original process, that is, $z_t = \Psi(B)(\omega_I I_t^{(T)} + a_t)$, where ω_I is the outlier size. IOs are expected to produce small effects on the sample autocorrelation function and the parameter estimates.

A third important type of outlier in time series is a level shift, that is, a modification of the local mean or level of the process starting at time T and continuing afterward. The model for the observed series is $z_t = \omega_L S_t^{(T)} + \Psi(B)a_t$, where $S_t^{(T)} = 0$ for $t < T$, $S_t^{(T)} = 1$ for $t \geq T$ (a step function), and ω_L is the magnitude of the shift. In SPC, level shifts are monitored using CUMulative SUM (CUSUM), exponentially weighted moving average (EWMA), or CUSCORE charts (*see Exponentially Weighted Moving Average (EWMA) Control Chart; Cumulative Sum (CUSUM) Chart*).

Several procedures for outlier detection have been proposed [28–33] after the seminal work of Chang *et al.* [34]; see [9] for a revision of some of these procedures.

Model Strengths and Weaknesses

Some strengths of ARIMA models are (a) they provide a very flexible family of models, which is closed to model addition; (b) the properties of the model are well understood and simple to use; (c)

their short-term forecasts are difficult to beat with more complex models; and (d) they can be easily generalized (e.g., they can generate time irreversible series using nonnormal noise).

Some weaknesses of ARIMA models are (a) they provide little help for long term forecasting; (b) exponential smoothing is easier to use and can often provide convenient approximations to reality and forecasts nearly as good as those of full *fitted* ARIMA approximations; and (c) the framework of ARIMA models may be inconvenient to handle (e.g., most models are not invariant to changes in the sampling interval, which is important for process control; or to temporal aggregation, which is important in economy).

Related Topics

ARIMA models can be generalized to incorporate long memory (ARIMA with fractional differencing or auto-regressive fractionally integrated moving average (ARFIMA) models [35, 36]); conditional **heteroscedasticity** (generalized auto-regressive conditional heteroscedasticity (GARCH) models [37]); nonlinearity (e.g., threshold AR models [38, 39]); and information from related simultaneous series (vector ARMA, ARIMA and cointegrated models [8, 40–44]).

The analysis of time series presented here is called time-domain analysis. Time series can also be analyzed using the frequency-domain approach in which Fourier transforms of the series and its autocovariance function are used. Bayesian models can also be used.

An important alternative to ARIMA models are the so-called state space or structural dynamic models [45–48]. In these models, the time series $\{y_t\}$ is represented as a linear function $y_t = h'_t \alpha_t + a_t$ of some unknown state parameters α_t , which are assumed to follow the equation $\alpha_t = A_t \alpha_{t-1} + u_t$, with appropriate noise series $\{a_t\}$ and $\{u_t\}$. The estimation and forecasting of these models are based on the *Kalman filter algorithm* [45, 46, 49].

Appendix 1: Evaluation of Some Useful Covariances

The autocovariance function $\gamma_k := \text{cov}(y_t, y_{t+k})$ may be computed efficiently using the formula

$\sum_{i=0}^p \phi_i \gamma_{|i-j|} = \sum_{i=0}^q \theta_i \gamma_{ay}(j-i)$, where $\phi_0 = \theta_0 = -1$ and $\gamma_{ay}(-k) := \text{cov}(a_t, y_{t+k}) = \psi_k \sigma^2$ for $k \geq 0$. For $j = 0, \dots, p$, the formula provides a linear system having solution $\gamma_0, \dots, \gamma_p$; for $j > p$, it can be used recursively. The covariances $\gamma_{ay}(-k)$ are null for $k < 0$ and satisfy the recursion $\gamma_{ay}(-k) = -\theta_k \sigma^2 + \sum_{i=1}^p \phi_i \gamma_{ay}(-k+i)$ for $k \geq 0$.

Appendix 2: Covariance Matrix of Working Series W

The covariance matrix Θ of the working series $\mathbf{W}' = (w_1, \dots, w_n)$ defined in the section titled “Forecasting Time Series” and used in the section titled “Model Estimation by Maximum Likelihood” is a band matrix having bandwidth $\max(p-1, q)$ in their p first rows and q thereafter, because $w_t = y_t$ for $t = 1, \dots, p$ and $w_t = \theta_q(B)a_t$ for $t = p+1, \dots, n$. Their elements are $\Theta_{ij} = \gamma_{|i-j|}$ for $i, j = 1, \dots, p$; $\Theta_{ij} = \Theta_{ji} = -\sum_{k=1}^q \theta_k \gamma_{ay}(j-i-k)$ for $i = 1, \dots, p$ and $j = p+1, \dots, n$, where γ_{ay} is given in Appendix 1; and $\Theta_{ij} = \sigma^2 \sum_{k=|j-i|}^q \theta_k \theta_{k-|j-i|}$ for $i, j = p+1, \dots, n$, where $\theta_0 = -1$.

Acknowledgments

Luceño and Peña acknowledge the support of the Spanish dirección general de investigación (DGI) under grants MTM2005-00 287 and SEJ2004-03 303, respectively, and thank George Box and two editors for helpful comments.

References

- [1] Wold, H.O. (1938). *A Study in the Analysis of Stationary Time Series*, Almqvist and Wicksell, Uppsala.
- [2] Box, G.E.P. & Luceño, A. (1997). *Statistical Control by Monitoring and Feedback Adjustment*, John Wiley & Sons, New York.
- [3] Box, G.E.P. & Jenkins, G.M. (1970). *Time Series Analysis: Forecasting and Control*, Holden Day, San Francisco.
- [4] Box, G.E.P., Jenkins, G.M. & Reinsel, G.C. (1994). *Time Series Analysis: Forecasting and Control*, 3rd Edition, Prentice-Hall, Englewood Cliffs.
- [5] Luceño, A. (2003). Dead-band adjustment schemes for on-line feedback quality control, in *Handbook of Statistics-Statistics in Industry*, R., Khattree & C.R., Rao, eds, Elsevier Science B. V., Vol. 22, pp. 695–727.
- [6] Pandit, S.M. & Wu, S.M. (1983). *Time Series and System Analysis with Applications*, John Wiley & Sons, New York.
- [7] Pankratz, A. (1983). *Forecasting with Univariate Box-Jenkins Models*, John Wiley & Sons, New York.
- [8] Shumway, R.H. & Stoffer, D.S. (2000). *Time Series Analysis and its Applications*, Springer-Verlag, New York.
- [9] Peña, D., Tiao, G.C. & Tsay, R.S. (2001). *A Course in Time Series Analysis*, John Wiley & Sons, New York.
- [10] Taguchi, G., Elsayed, E.A. & Hsiang, T. (1989). *Quality Engineering in Production Systems*, McGraw-Hill, New York.
- [11] Luceño, A., González, F.J. & Puig-Pey, J. (1996). Computing optimal adjustment schemes for the general tool-wear problem, *Journal of Statistical Computation and Simulation* **54**, 87–113.
- [12] Newbold, P. (1981). Some recent developments in time series analysis, *International Statistical Review* **49**, 53–66.
- [13] Dickey, D.A. & Fuller, W.A. (1981). Likelihood ratio statistics for autoregressive processes with a unit root, *Econometrica* **49**, 1057–1072.
- [14] Tsay, R.S. & Tiao, G.C. (1984). Consistent estimates of autoregressive parameters and extended sample autocorrelation function of stationary and nonstationary ARMA models, *Journal of the American Statistical Association* **79**, 84–96.
- [15] Cleveland, W.S. (1972). The inverse autocorrelations of a time series and their applications, *Technometrics* **14**, 277–293.
- [16] Chatfield, C. (1980). Inverse autocorrelations, *Journal of the Royal Statistical Society. Series A* **142**, 363–377.
- [17] Ljung, G.M. (1989). A note on the estimation of missing values in time series, *Communications in Statistics Part B* **18**, 459–465.
- [18] Peña, D. & Maravall, A. (1991). Interpolation, outliers and inverse autocorrelations, *Communications in Statistics Part A* **20**, 3175–3186.
- [19] Luceño, A. (1997). Estimation of missing values in possibly partially nonstationary vector time series, *Biometrika* **84**, 495–499.
- [20] Jones, R.H. (1980). Maximum likelihood fitting of ARMA models to time series with missing observations, *Technometrics* **22**, 389–396.
- [21] Ljung, G.M. & Box, G.E.P. (1978). On a measure of lack of fit in time series models, *Biometrika* **65**, 297–303.
- [22] Peña, D. & Rodríguez, J. (2006). The log of the determinant of the autocorrelation matrix for testing goodness of fit in time series, *Journal of Statistical Planning and Inference* **136**, 2706–2708.
- [23] Akaike, H. (1973). Information theory and an extension of the maximum likelihood principle, in *Proceedings of the 2nd International Symposium on Information Theory*, Akademiai Kiadó, Budapest, pp. 267–281.
- [24] Akaike, H. (1974). A new look at the statistical model identification, *IEEE Transactions on Automatic Control* **19**, 203–217.

-
- [25] Schwarz, G. (1978). Estimating the dimension of a model, *The Annals of Statistics* **6**, 461–464.
 - [26] McQuarrie, A.D.R. & Tsai, C.L. (1998). *Regression and Time Series Model Selection*, World Scientific, Singapore.
 - [27] Fox, A.J. (1972). Outliers in time series, *Journal of the Royal Statistical Society. Series B* **34**, 350–363.
 - [28] Tsay, R.S. (1988). Outliers, level shifts, and variance changes in time series, *Journal of Forecasting* **7**, 1–20.
 - [29] Wu, L.S.-Y., Ravishanker, N. & Hosking, J.R.M. (1993). Reallocation outliers in time series, *Applied Statistics* **42**, 301–313.
 - [30] Barnett, V. & Lewis, T. (1994). *Outliers in Statistical Data*, 3rd Edition, John Wiley & Sons, Chichester.
 - [31] Chen, C. & Liu, L.M. (1993). Forecasting time series with outliers, *Journal of Forecasting* **12**, 13–35.
 - [32] Luceño, A. (1998). Detecting possibly non-consecutive outliers in industrial time series, *Journal of the Royal Statistical Society B* **60**, 295–310.
 - [33] Tsay, R.S., Peña, D. & Pankratz, A. (2000). Outliers in multivariate time series, *Biometrika* **87**, 789–804.
 - [34] Chang, I., Tiao, G.C. & Chen, C. (1988). Estimation of time series parameters in the presence of outliers, *Technometrics* **30**, 193–204.
 - [35] Hosking, J.R.M. (1981). Fractional differencing, *Biometrika* **68**, 165–176.
 - [36] Luceño, A. (1996). A fast likelihood approximation for vector general linear processes with long series: Application to fractional differencing, *Biometrika* **83**, 603–614.
 - [37] Engle, R.F. (1982). Autoregressive conditional heteroscedasticity with estimates of the variance of United Kingdom inflations, *Econometrica* **50**, 987–1007.
 - [38] Tong, H. (1983). *Threshold Models in Nonlinear Time Series Analysis*, Springer-Verlag, New York.
 - [39] Tong, H. (1990). *Non-linear Time Series: A Dynamical System Approach*, Oxford University Press, Oxford.
 - [40] Hannan, E.J. (1970). *Multiple Time Series*, John Wiley & Sons, New York.
 - [41] Wei, W.W.S. (1990). *Time Series Analysis*, Addison-Wesley, Redwood.
 - [42] Lütkepohl, H. (1993). *Introduction to Multiple Time Series Analysis*, Springer-Verlag, New York.
 - [43] Reinsel, G.C. (1993). *Elements of Multivariate Time Series Analysis*, Springer-Verlag, New York.
 - [44] Engle, R.F. & Granger, C.W.J. (1987). Co-integration and error correction: representation, estimation and testing, *Econometrica* **55**, 251–276.
 - [45] Kalman, R.E. (1960). A New approach to linear filtering and prediction problems, *Transactions of the ASME, Series D, Journal of Basic Engineering* **82**, 35–45.
 - [46] Kalman, R.E. & Bucy, R.S. (1961). New results in linear filtering and prediction theory, *Transactions of the ASME, Series D, Journal of Basic Engineering* **83**, 95–107.
 - [47] Harvey, A.C. (1989). *Forecasting Structural Time Series Models and the Kalman Filter*, Cambridge University Press.
 - [48] West, M. & Harrison, J. (1989). *Bayesian Forecasting and Dynamic Models*, Springer-Verlag, New York.
 - [49] Brockwell, P.J. & Davis, R.A. (1991). *Time Series: Theory and Methods*, 2nd Edition, Springer-Verlag, New York.

Related Article

Brownian Motion; Correlation; Covariance; Exponentially Weighted Moving Average (EWMA); Markov Processes; Moving Averages; Normal Distribution.

ALBERTO LUCEÑO AND DANIEL PEÑA

Process Capability, Error in Estimation of

Introduction

The inevitable variations in process measurements come from two sources: the manufacturing process and the gauge. Gauge capability reflects the gauge's precision, or lack of variation, but is not the same as **calibration**, which assures the gauge's accuracy. Process capability measures the ability of a manufacturing process to meet preassigned specifications. Nowadays, many customers use process capability to judge a supplier's ability to deliver quality products. Suppliers ought to be aware of how gauges affect various process capability estimates.

The gauge capability includes two components of **measurement systems** variability: repeatability and reproducibility. Repeatability represents the variability from the gauge or measurement instrument when it is used to measure the same specimen (with the same operator or setup or in the same time period). Reproducibility reflects the variability arising from different operators, setups, time periods, or in general, different conditions. These studies are often referred to as *gauge repeatability and reproducibility* (GR&R) studies. To summarize we have

$$\begin{aligned}\sigma_{\text{measurement error}}^2 &= \sigma_{\text{gauge}}^2 \\ &= \sigma_{\text{repeatability}}^2 + \sigma_{\text{reproducibility}}^2\end{aligned}\quad (1)$$

Estimates for $\sigma_{\text{repeatability}}^2$ and $\sigma_{\text{reproducibility}}^2$ come from a gauge study, or a GR&R study. Barrentine [1], Levinson [2], and Montgomery [3, 4], among others, provided procedures for GR&R studies.

Gauge capability is a gauge's ability to repeat and reproduce measurements. Its measurement is the percentage of tolerance consumed by (gauge) capability. Montgomery [4] referred to it as the *precision-to-tolerance* (or P/T) ratio. It is the ratio of the gauge's variation to the specification width; its smaller numbers are of course preferable. Denoting the gauge's standard deviation by σ_{gauge} , we have

$$P/T = \frac{6\sigma_{\text{gauge}}}{USL - LSL} \times 100\% \quad (2)$$

where USL and LSL represent the upper and lower specification limits, respectively. Some authors and

practitioners use the coefficient 5.15 instead of 6 (see, e.g. Barrentine [1] and Levinson [2, 5]). This formula uses 6σ as the natural tolerance width for the gauge, based on and motivated by the **normal distribution** assumptions. Values of the estimated ratio P/T of 0.1 or less often are taken to imply adequate gauge capability. This is based on the generally used rule that requires a measurement device to be calibrated in units one-tenth as large as the accuracy required in the final measurement.

GR&R studies focus on quantifying the measurement errors. Common approaches to GR&R studies, such as the range method [6] and the **ANOVA** (analysis of variance) method [7, 8], assume that the distribution of the measurement errors is normal with a mean error of zero. Let the measurement errors be described by a random variable $M \sim N(0, \sigma_M^2)$, Montgomery and Runger [6] presented the gauge capability λ by means of the formula:

$$\lambda = \frac{6\sigma_M}{USL - LSL} \times 100\% \quad (3)$$

For the measurement system to be deemed acceptable, the variability in the measurements due to the measurement system ought to be less than a predetermined percentage of the engineering tolerance. The recommended guidelines for gauge acceptance are presented in the Automotive Industry Action Group (AIAG) Measurement Systems Analysis (**MSA**) manual. The purpose of a GR&R study is to determine if the variability of the measurement system is small relative to the variability of the monitored process. Several commonly reported ratios (functions of parameters of interest) in GR&R studies are discussed in Burdick *et al.* [9]. Burdick *et al.* [9] also provided a review of methods for constructing measurement systems capability studies with emphasis on the ANOVA approach to analyze the results. More recently, Burdick *et al.* [10] applied the generalized inference methods to construct **confidence intervals** for assessing the ability of a measurement system to correctly classify parts in a GR&R study. The reported simulation suggests that, for appropriately sized experiments, the proposed method provides two-sided intervals and upper bounds that generally maintain the stated confidence level.

Process capability indices (PCIs) are convenient and powerful tools for measuring process performance. In recent years, PCIs have received considerable research attention in the quality assurance and statistical literature. Those indices quantify process performance by taking into consideration process location, process variation, and manufacturing specifications, which reflect process consistency, process accuracy, **process yield**, and process loss. A substantial majority of capability research works that appeared in the literature do not take into account gauge measurement errors. However, the gauge measurement errors have a significant effect on estimating and assessing process capability since measurement systems are not 100% precise. An inaccurate measurement system can thwart all the benefits of improvement endeavors resulting in poor quality. Analyzing process capability without considering gauge capability may often lead to unreliable decisions. It could result in a serious loss to producers if gauge capability is not being considered in process capability **estimation** and testing. On the other hand, improving the gauge measurements and having properly trained operators can reduce the measurement errors. Since measurement errors unfortunately cannot be avoided, using appropriate confidence coefficients and power becomes an essential task. However, the reality is that no measurement is free from error or uncertainty even if it is carried out with the aid of highly sophisticated and precise measuring instruments. Montgomery and Runger [6] pointed out that the quality of the data related to the process characteristics relies substantially on the gauge. Any variation in the measurement process has a direct impact on the ability to execute sound judgment about the manufacturing process. Conclusions about capability of a process based only on the single numerical value of the index are not reliable. Analyzing the effects of measurement errors on PCIs, Mittag [11, 12] and Levinson [2] developed very definitive techniques for quantifying the percentage error in PCIs estimation in the pence of measurement errors. Bordignon and Scagliarini [13] studied the inferential properties of the estimators of C_p and C_{pk} when observations are contaminated by measurement errors, while Scagliarini [14] analyzed the properties of the estimator of C_p for autocorrelated data in the presence of measurement errors. Bordignon and Scagliarini [15] studied the properties of the estimator of C_{pm} when observations

are affected by measurement errors. They compared the performances of the estimator on the error case with those of the estimator in the error-free case. The results indicated that the presence of measurement errors in the data leads to different behavior of the estimator according to the entity of the error variability.

Estimating Process Capability with Measurement Errors

Most research works related to PCIs are carried out under the assumption of no gauge measurement errors. Unfortunately, such an assumption does not adequately accommodate real-world situations even with modern, highly sophisticated measuring instruments and devices. In this section, we have emphasized constructing confidence intervals and testing procedures on the most widely used index C_{pk} in presence of gauge measurement errors.

Let $X \sim N(\mu, \sigma^2)$ represent the relevant quality characteristic of a manufacturing process and C_{pk} measure the “true” process capability. However, in practice, we deal with the observed variable Y rather than with the “true” variable X . Assume that X and M (the measurement error) are stochastically independent. In this case, we have $Y \sim N(\mu_Y = \mu, \sigma_Y^2 = \sigma^2 + \sigma_M^2)$, and the empirical process capability index C_{pk}^Y will be obtained after substituting σ_Y for σ . The relationship between the true process capability $C_{pk} = \min\{(USL - \mu)/3\sigma, (\mu - LSL)/3\sigma\} = d - |\mu - m|/(3\sigma)$ and the empirical process capability C_{pk}^Y can be expressed as follows:

$$\frac{C_{pk}^Y}{C_{pk}} = \frac{1}{\sqrt{1 + \lambda^2 C_p^2}} \quad (4)$$

where $d = (USL - LSL)/2$ refers to the half-length of the allowable tolerance of the process, $m = (USL + LSL)/2$ is the midpoint of the specification interval, $C_p = d/3\sigma$, and λ is the gauge capability index.

Since the variation of the observed data is larger than the variation of the original one, the denominator of the index C_{pk} becomes larger, and the true capability of the process will be understated if calculations of PCIs are based on the empirical data Y (the observed measurement contaminated with errors). Suppose that the empirical data

$\{Y_i, i = 1, 2, \dots, n\}$ is available, then the natural estimator

$$\hat{C}_{pk}^Y = \frac{d - |\bar{Y} - m|}{3S_Y} \quad (5)$$

obtained by replacing the process mean μ and the process standard deviation σ by their conventional estimators $\bar{Y} = \sum_{i=1}^n Y_i/n$ and $S_Y = \left[\frac{\sum_{i=1}^n (Y_i - \bar{Y})^2}{(n-1)} \right]^{1/2}$, from a stable process. Accordingly, applying the technique used in Kotz and Johnson [16] and Pearn and Lin [17], the cumulative distribution function (CDF) of $\hat{C}_{pk}^Y$ is obtained as follows:

$$\begin{aligned} F_{\hat{C}_{pk}^Y}(x) &= 1 - \int_0^{3C_p^Y \sqrt{n}} G \left(\frac{(n-1)(3C_p^Y \sqrt{n} - t)^2}{9nx^2} \right) \\ &\quad \times \phi[t + 3(C_p^Y - C_{pk}^Y)\sqrt{n}] \\ &\quad + \phi[t - 3(C_p^Y - C_{pk}^Y)\sqrt{n}] dt \end{aligned} \quad (6)$$

where $C_p^Y = \frac{C_p}{\sqrt{1+\lambda^2 C_p^2}}$, $C_{pk}^Y = \frac{C_{pk}}{\sqrt{1+\lambda^2 C_p^2}}$, $\phi(\cdot)$ is the probability density function (PDF) of the standard normal distribution, and, as above, $\lambda = 6\sigma_M/(USL - LSL) \times 100\%$.

To determine whether a given process meets the preset capability requirement, we could consider the statistical testing with the null hypothesis $H_0: C_{pk} \leq c$ (the process is not capable) and the alternative hypothesis $H_1: C_{pk} > c$ (the process is capable), where c is the required capability level. If the calculated process capability is greater than the corresponding critical value, we reject the null hypothesis and conclude that the process is capable.

Discussions in Pearn and Liao [18] indicated that the true process capability would be severely underestimated if $\hat{C}_{pk}^Y$ is used. The probability that $\hat{C}_{pk}^Y$ is greater than c_0 would be less than when using $\hat{C}_{pk}$. Thus, the α risk of using $\hat{C}_{pk}^Y$ is less than the α risk of using $\hat{C}_{pk}$ when estimating C_{pk} . The power of the test based on $\hat{C}_{pk}^Y$ is also less than that of the one based on $\hat{C}_{pk}$; that is, the α risk and the power of the test decrease when the gauge measurement error increases. Since the lower confidence bound is severely underestimated and the power becomes small, producers cannot firmly state that their processes meet the capability requirement even if their processes are indeed sufficiently

capable. Adequate and even superior product units could be incorrectly rejected in this case. To remedy the situation, Pearn and Liao [18] considered the adjustment of the confidence bounds and of the critical values to provide a superior capability assessment. Suppose that the desired confidence coefficient is θ , the adjusted confidence interval of C_{pk} with the lower confidence bound L^* , can be shown to be

$$\begin{aligned} \theta &= P(C_{pk} > L^*) \\ &= 1 - \int_0^{b^* \sqrt{n}} G \left(\frac{(n-1)(b^* \sqrt{n} - t)^2}{9n(\hat{C}_{pk}^Y)^2} \right) \\ &\quad \times [\phi(t + \xi^Y \sqrt{n}) + \Phi(t - \xi^Y \sqrt{n})] dt \end{aligned} \quad (7)$$

where $b^* = \frac{3L^*}{\sqrt{1+\lambda^2 C_p^2}} + |\xi^Y|$ and $\xi^Y = 3(C_p^Y - C_{pk}^Y)$. To eliminate the need for estimating ξ^Y , Pearn and Liao [18] followed the technique of Pearn and Shu [19] (by setting $\xi^Y = 1.00$) to find the adjusted lower confidence bound L^* , where C_p can be obtained from the equation $\frac{3(C_p - L^*)}{\sqrt{1+\lambda^2 C_p^2}} = 1.00$ as

$$\begin{aligned} C_p &= \frac{18L^* + \sqrt{324L^{*2} - 4(9 - \lambda^2)(9L^{*2} - 1)}}{2(9 - \lambda^2)} \\ &= C_{p1} \end{aligned} \quad (8)$$

To improve the power of the test, revised critical values c_0^* can be determined from

$$\begin{aligned} \alpha^* &= P(\hat{C}_{pk}^Y \geq c_0^* | C_{pk} = c) \\ &= \int_0^{3C_p^Y \sqrt{n}} G \left(\frac{(n-1)(3C_p^Y \sqrt{n} - t)^2}{9n(c_0^*)^2} \right) \\ &\quad \times \phi[t + 3(C_p^Y - C_{pk}^Y)\sqrt{n}] \\ &\quad + \phi[t - 3(C_p^Y - C_{pk}^Y)\sqrt{n}] dt \end{aligned} \quad (9)$$

where α^* is the α risk corresponding to the test using $\hat{C}_{pk}^Y$ and c_0^* , $C_p^Y = \frac{C_p}{\sqrt{1+\lambda^2 C_p^2}}$, $C_{pk}^Y = \frac{c}{\sqrt{1+\lambda^2 C_p^2}}$, and c is the capability requirement. To eliminate the need for further estimation of the characteristic parameter C_p , we set $C_p^Y = C_{pk}^Y + 1/3$ as suggested by Pearn and Lin [17], and find a reliable adjusted critical value c_0^* , where C_p can be obtained using the equation

4 Process Capability, Error in Estimation of

$\frac{C_p}{\sqrt{1+\lambda^2 C_p^2}} = \frac{c}{\sqrt{1+\lambda^2 C_p^2}} + 1/3$ as

$$C_p = \frac{18c + \sqrt{324c^2 - 4(9 - \lambda^2)(9c^2 - 1)}}{2(9 - \lambda^2)} = C_{p2} \quad (10)$$

To ensure that the α risk is within the preset magnitude, let $\alpha^* = \alpha$ and solve the above equation to obtain c_0^* . The revised power of the test (denoted by π^*) can be calculated as

$$\begin{aligned} \pi^*(C_{pk}) &= P(\hat{C}_{pk}^Y \geq c_0 | C_{pk}, C_p = C_{p2}) \\ &= \int_0^{3C_p^Y \sqrt{n}} G\left(\frac{(n-1)(3C_p^Y \sqrt{n} - t)^2}{9n(c_0^*)^2}\right) \\ &\quad \times [\phi(t + \sqrt{n}) + \Phi(t - \sqrt{n})] dt \quad (11) \end{aligned}$$

where $C_p^Y = C_p^Y + 1/3$ and $C_{pk}^Y = \frac{C_{pk}}{\sqrt{1+\lambda^2 C_{p2}^2}}$. Pearn and Liao [18] adjusted the confidence intervals and the critical values in order to ensure that the intervals possess the desired confidence coefficients and will improve the power of the test with an appropriate α risk. When estimating the capability, the estimator $\hat{C}_{pk}^Y$ severely underestimates the true capability in the presence of measurement errors. Consequently, if a statistical test is used to determine whether the process meets the capability requirement, the power of the test would decrease substantially. Consequently, lower confidence bounds and critical values ought to be adjusted to improve the accuracy of capability assessment.

An Application Example

Consider the following case taken from a manufacturing factory that makes a type of low-power, fast warm-up, and excellent stability precision voltage reference (PVR) of 15 V. The output voltage is extremely insensitive to both line and load variations and can be externally adjusted with minimal effect on drift and stability. They are ideal for communications equipment, data acquisition systems, instrumentation and process control, high-precision power supplies, and portable battery-powered equipment. Portable battery-powered equipment (Notebook computers, PDAs, DVMs, GPS,

etc.). Initial accuracy is one critical quality characteristic of this PVR which has a significant impact on the PVR quality/reliability. This characteristic is usually for room temperature only, and it provides a starting point for most of the other specifications. The output voltage tolerance of a reference is measured without a load applied after the device is turned on and warmed up. Manufacturers specify a reference with a tight initial error so that they do not have to perform room-temperature systems calibration after assembly.

For this particular model of PVR product, the specification limits are $T = m = 15$ V, $USL = 15.025$ V, and $LSL = 14.975$ V. A total of 80 observations are collected from the manufacturing process in the factory. Histograms and normal **probability plots** show that the collected data follows the normal distribution. The Shapiro–Wilk test is applied to further justify the assumption. The knowledge of the gauge capability λ is important in order to determine the effects of measurement errors on the performance of the estimator for C_{pk} . In practice, we do not know the value of λ , however, it is possible to obtain a point estimate and confidence interval through a GR&R study performed with the ANOVA method. Burdick and Larsen [20] reported the upper and lower bounds for approximate $100(1 - \alpha)\%$ confidence intervals for the parameters σ_Y^2 , σ^2 , and σ_M^2 based on modified large-sample (MLS) methods. To determine whether the process is “excellent” ($C_{pk} > 1.33$) with unavoidable measurement errors $\lambda = 0.20$ (provided by the gauge manufacturing factory), we first determine that $c = 1.33$ and $\alpha = 0.05$. Then, based on the sample data of 80 observations, we obtain the sample mean $\bar{y} = 15.0049$, the sample standard deviation $s_y = 0.0045$ ($\bar{y}$ and s_y are the realized sample values for $\bar{Y}$ and S_Y), and the point estimator $\hat{C}_{pk}^Y = 1.489$. Thus, the 95% lower confidence bound of the true process capability can be obtained as $L^* = 1.356$. Moreover, similar to the adjusted lower confidence bound, we obtain the adjusted critical value 1.464 based on $\alpha = 0.05$, $\lambda = 0.20$, and $n = 80$. Since $\hat{C}_{pk}^Y > 1.464$, we conclude that the process is “excellent”. We also see that if we ignore the measurement errors and evaluate the critical value without any correction, the critical value may be calculated as $c_0 = 1.543$. In this case we would reject that the process is “excellent” since $\hat{C}_{pk}^Y$

is not greater than the uncorrected critical value 1.543.

Conclusions

To summarize our discussion: when estimating capability the confidence bounds are substantially underestimated in the presence of measurement errors. When we use statistical testing to determine if the process meets the capability requirements, we find that the power of the test decreases with the measurement errors. Since measurement errors are, unfortunately, unavoidable in many branches of the manufacturing industry, to obtain a more accurate confidence bound and to improve the power with an appropriate α risk, adjusted confidence bounds and the critical values should definitely be applied. If the producers do not take into account the effects of the gauge capability on process capability estimation and testing, it may cause substantial and serious loss. In that case, producers are unable to affirm anymore that their processes will meet the capability requirements even if they are indeed sufficiently capable. The producers may incur substantial costs as (large) quantities of acceptable product units will be incorrectly rejected. Improving the gauge measurement accuracy and training the operators to properly execute the production procedures at each step are essential for reducing the measurement errors.

References

- [1] Barrentine, L.B. (1991). *Concepts for R&R Studies*, ASQC Quality Press, Milwaukee.
- [2] Levinson, W.A. (1995). How good is your gage? *Semiconductor International*, **October**, 165–168.
- [3] Montgomery, D.C. (2001). *Introduction to Statistical Quality Control*, 4th Edition, John Wiley & Sons, New York.
- [4] Montgomery, D.C. (2005). *Introduction to Statistical Quality Control*, 5th Edition, John Wiley & Sons, New York.
- [5] Levinson, W.A. (1996). Do you need a new gage? *Semiconductor International*, **February**, 113–116.
- [6] Montgomery, D.C. & Runger, G.C. (1993). Gauge capability and designed experiments. Part I: basic methods, *Quality Engineering* **6**(1), 115–135.
- [7] Mandel, J. (1972). Repeatability and reproducibility, *Journal of Quality Technology* **4**(2), 74–85.
- [8] Montgomery, D.C. & Runger, G.C. (1993). Gauge capability analysis and designed experiments. Part II: experimental design models and variance component estimation, *Quality Engineering* **6**(2), 289–305.
- [9] Burdick, R.K., Borror, C.M. & Montgomery, D.C. (2003). A review of methods for measurement systems capability analysis, *Journal of Quality Technology* **35**(4), 342–354.
- [10] Burdick, R.K., Park, Y.J., Montgomery, D.C. & Borror, C.M. (2005). Confidence intervals for misclassification rates in a gauge R&R study, *Journal of Quality Technology* **37**(4), 294–303.
- [11] Mittag, H.J. (1994). Measurement error effects on the performance of process capability indices. in Paper presented at the *5th International Workshop on Intelligent Statistical Quality Control*, Osaka.
- [12] Mittag, H.J. (1997). Measurement error effects on the performance of process capability indices, *Frontiers in Statistical Quality Control* **5**, 195–206.
- [13] Bordignon, S. & Scagliarini, M. (2002). Statistical analysis of process capability indices with measurement errors, *Quality and Reliability Engineering International* **18**, 321–332.
- [14] Scagliarini, M. (2002). Estimation of C_p for autocorrelated data and measurement errors, *Communications in Statistics: Theory and Methods* **31**, 1647–1664.
- [15] Bordignon, S. & Scagliarini, M. (2006). Estimation of C_{pm} when measurement error is present, *Quality and Reliability Engineering International*, **22**(7), 787–801.
- [16] Kotz, S. & Johnson, N.L. (1993). *Process Capability Indices*, Chapman & Hall, London.
- [17] Pearn, W.L. & Lin, P.C. (2004). Testing process performance based on the capability index C_{pk} with critical values, *Computers and Industrial Engineering* **47**, 351–369.
- [18] Pearn, W.L. & Liao, M.Y. (2005). Measuring process capability based on C_{pk} with gauge measurement errors, *Microelectronics Reliability* **45**, 739–751.
- [19] Pearn, W.L. & Shu, M.H. (2003). Manufacturing capability control for multiple power distribution switch processes based on modified C_{pk} MPPAC, *Microelectronics Reliability* **43**, 963–975.
- [20] Burdick, R.K. & Larsen, G.A. (1997). Confidence intervals for a ratio in a R&R study, *Journal of Quality Technology* **29**(3), 261–273.

Further Reading

- Hamada, M. & Weerahandi, S. (2000). Measurement system assessment via generalized inference, *Journal of Quality Technology* **32**(3), 241–253.

WEN LEA PEARN AND CHIEN-WEI WU

Process Capability Indices, Skewed

Introduction

The conventional process capability indices (PCIs) such as C_p , C_{pk} , C_{pm} , or C_{pmk} are usually determined under the assumption that quality characteristics follow a **normal distribution**. However, the normality assumption is usually difficult to justify and is often not appropriate in practice. For example, the measurements from drilling processes, coating processes, chemical processes, and semiconductor processes are often skewed.

In general, a process is considered to be capable if the process consistently produces items within given specifications. Therefore, many practitioners have utilized the PCIs, especially C_p or C_{pk} , in connection with the **process yield**. If a process does not follow a normal distribution, however, none of the conventional PCIs provide valid measures of the number of parts nonconforming. Table 1 gives the values of C_p as skewness (see **Skewness**) α_3 increases for Weibull, lognormal, and gamma distributions. The table also gives expected number of nonconforming

items per million (NPM) and “equivalent $C_p(EC_p)$ ” calculated by $-(1/3)\Phi^{-1}((NPM/2) \times 10^{-6})$ since $NPM \times 10^{-6} = 2\Phi(-3C_p)$ under normality, where $\Phi(\cdot)$ is the standard normal distribution function. If the value of a PCI is close to that of the equivalent C_p , the PCI describes the process capability very well. In each case in the table, it is assumed that the lower and upper specification limits (LSL and USL) are -3 and 3 , respectively, and the distribution is shifted and scaled to produce the same value of the process mean $\mu = 0$ and the standard deviation $\sigma = 1$, so that $C_p = 1$ for all cases. The table shows that the NPM increases and the EC_p decreases as skewness increases. The C_p , however, does not reflect this phenomenon and overestimates the process capability. Therefore, the index could not be used to evaluate process capability if the distribution is skewed. For more detailed discussions on the effects of nonnormality (see Somerville and Montgomery [1] and Kotz and Lovelace [2]).

Skewed PCIs to remedy the shortcoming of the conventional PCIs can be divided into five main lines: (a) To use normalizing transformations, (b) to replace the unknown distribution by an empirical distribution or by a known three- or four-parameter distribution, (c) to modify the standard definition of PCIs in order to increase their robustness, (d) to construct PCIs with the estimate of the process yield, and (e) to develop heuristic methods adequate for skewed populations.

Normalizing Transformation Method

The normalizing transformation method is the simplest way to deal with nonnormal data, which transforms original process data to normal or close to normal with some mathematical functions so that the conventional PCIs can be used. For example, if a skewed distribution becomes normal by a square root transformation, the transformed data and specification limits can be used to estimate capability indices. Such transformations include Box–Cox power transformation (see **Box and Cox Transformation**) [1] or Johnson transformation [3]. The method is easy to understand, but it may be difficult to translate the results with transformed data back to the original scale and it is likely to require considerable work due to its trial and error nature for finding a proper transformation function.

Table 1 C_p , NPM, and EC_p versus α_3

Distribution	α_3	NPM	EC_p	C_p
Normal	0.0	2700	1.00	1.00
Weibull	0.5	4227	0.95	1.00
	1.0	9870	0.86	1.00
	1.5	14 915	0.81	1.00
	2.0	18 316	0.79	1.00
	2.5	20 280	0.77	1.00
	3.0	21 256	0.76	1.00
Lognormal	0.5	5639	0.92	1.00
	1.0	10 461	0.85	1.00
	1.5	14 087	0.82	1.00
	2.0	16 358	0.80	1.00
	2.5	17 653	0.79	1.00
	3.0	18 325	0.79	1.00
Gamma	0.5	5431	0.93	1.00
	1.0	10 336	0.86	1.00
	1.5	14 782	0.81	1.00
	2.0	18 316	0.79	1.00
	2.5	20 856	0.77	1.00
	3.0	22 528	0.76	1.00

Quantile Estimation Method

Since $\mu = X_{0.5}$, $\mu - 3\sigma = X_{0.00135}$, and $\mu + 3\sigma = X_{0.99865}$ for a normal distribution, where X_p is the p th quantile (see **Quantiles**), C_p and C_{pk} can be redefined as

$$C_p = \frac{USL - LSL}{X_{0.99865} - X_{0.00135}} \quad (1)$$

$$C_{pk} = \min \left\{ \frac{USL - X_{0.5}}{X_{0.99865} - X_{0.5}}, \frac{X_{0.5} - LSL}{X_{0.5} - X_{0.00135}} \right\} \quad (2)$$

If a proper distribution can be specified, accurate measures of the three quantiles can be obtained. This involves modeling the process data with a probability model such as Weibull, lognormal, or gamma. If the variability pattern of the data cannot be described by a particular distribution, a flexible distribution approach or a nonparametric approach can be considered. The flexible distribution approach is to fit a family of distributions such as Burr or Pearson distribution to the data and use the quantiles of the fitted distribution. The nonparametric approach is to perform a nonparametric **estimation** of the quantiles. See Clements [4], Castagliola [5], and Ding [6] for the flexible distribution approach and Chang and Lu [7], Polansky [8], and **Process Capability Indices, Nonparametric**.

Clements [4] proposed to use Pearson distributions in estimating the quantiles, $X_{0.5}$, $X_{0.00135}$, and $X_{0.99865}$. These quantiles are calculated by

$$X_{0.5} = \bar{x} + z_{0.5} \cdot s \quad (3)$$

$$X_{0.00135} = \bar{x} + z_{0.00135} \cdot s \quad (4)$$

$$X_{0.99865} = \bar{x} + z_{0.99865} \cdot s \quad (5)$$

where z_p is the standardized quantile of the Pearson curve expressed as a function of skewness and **kurtosis** and can be obtained from the tables in Clements [4] and $\bar{x}$ and s are the sample mean and the sample standard deviation, respectively. The process capability can then be obtained by equations (1) and (2). Pearn and Kotz [9] applied the Clements' method to calculating second and third generation PCIs, C_{pm} and C_{pmk} , for nonnormal populations. The Clements' PCIs can be easily calculated since no complicated distribution fitting is required. Clements' approach, however, requires estimates of the third and fourth **moments** of the data and it is not possible to establish a functional relationship between Clements'

C_p and process yield that can be applied to various distributions.

Shore [10] proposed a PCI with the family of distributions introduced by Shore [11] which can be fitted with the estimated first and second moments. Note that the third and fourth moments are needed to obtain Clements' PCIs. Based on the distribution, Shore [10] suggested to calculate the three quantiles by

$$X_{0.5} = M \quad (6)$$

$$X_{0.00135} = A_1 \cdot 0.001352^{B_1} \quad (7)$$

$$X_{0.99865} = A_2 \cdot \log(739.7) + B_2 \quad (8)$$

where M is the sample median and $\{A_i, B_i\} (i = 1, 2)$ are parameters of the assumed distribution determined by the modified two-moment procedure. The **standard error** of the estimator of the Shore's PCI is fairly small, but it does not perform as well as the Clements' PCI, partly because Shore uses only the first and second moments for distribution fitting.

The reliability of the capability estimates can only be assured when the data follow at least approximately the parametric distribution being used even if a flexible distribution is assumed. To avoid this weakness of the flexible distribution method, Polansky [8] used a nonparametric approach based on the kernel technique.

The quantile estimation method can be applied to a broad range of distributions, but the method is somewhat complicated and needs considerably large samples.

Robust PCI

Pearn *et al.* [12] introduced C_θ defined as

$$C_\theta = \frac{USL - LSL}{\theta\sigma} \quad (9)$$

where θ is chosen so that the probability

$$\Pr \left\{ \mu - \frac{1}{2}\theta\sigma < X < \mu + \frac{1}{2}\theta\sigma \right\} \quad (10)$$

would be close to one and as independent as possible of the distribution. They recommended using $\theta = 5.15$ because the approximate relation

$$\Pr\{\mu - 2.575\sigma \leq X \leq \mu + 2.575\sigma\} \simeq 0.99 \quad (11)$$

varies only slightly for a number of distributions including χ^2 distributions with various **degrees of freedom**. This robustness is accomplished with the trade-off giving up the 2700 NPM for the normal population and replacing it by 10 000.

Rodriguez [13] proposed a PCI using the robust estimators of the process mean and dispersion as follows.

$$\tilde{C}_{pk} = \min \left\{ \frac{USL - \tilde{X}}{4.45MAD}, \frac{\tilde{X} - LSL}{4.45MAD} \right\} \quad (12)$$

where $\tilde{X}$ is the median of a process and MAD is the median of $(|X_1 - \tilde{X}|, \dots, |X_n - \tilde{X}|)$.

The robust method may be useful if there are insufficient data to provide accurate distribution fitting. The value of robust PCIs, however, may not coincide with that of conventional PCIs with “6 σ ” even for a normal distribution.

Process Yield Estimation Method

If $USL = \mu + k\sigma$ and $LSL = \mu - k\sigma$, a simple way to assess the process capability, is to obtain a nonconforming proportion like $p_{NC} = \Pr\{X < \mu - k\sigma \text{ or } X > \mu + k\sigma\}$, and transform the probability into the PCI using $C_p = -(1/3)\Phi^{-1}(p_{NC}/2)$. To obtain the limit of the probability for an arbitrary distribution, the Chebyshev inequality could be used, however, the limit is so broad that it is of little practical value. Flaig [14] proposed to tighten the probability bounds and estimate PCIs using an extension of the Camp–Meidell theorem. The Camp–Meidell inequality

$$\Pr\{X < \mu - k\sigma \text{ or } X > \mu + k\sigma\} \leq \frac{1}{2.25k^2} \quad (13)$$

is applicable only if the underlying distribution is unimodal and symmetrical, but he found a way to overcome this limitation. The Camp–Meidell inequality leads to the calculation of the probability that the process produces individuals outside the specification limits. It cannot be viewed as a PCI in the strict sense [2].

Chao and Lin [15] proposed to construct a PCI with an estimate of the process yield using the relation between the process yield and the C_p used under normality;

$$C_p = \frac{1}{3}\Phi^{-1} \left[\frac{1}{2} \{F(USL) - F(LSL) + 1\} \right] \quad (14)$$

where $F(\cdot)$ is the distribution function of the process. They also suggested use of a nonparametric method for estimating $F(USL)$ and $F(LSL)$ like

$$\hat{F}(x) = \frac{1}{n} \sum_{i=1}^n \Phi \left(\frac{x - X_i}{1.06 \cdot s \cdot n^{-1/5}} \right) \quad (15)$$

The nonparametric method for estimating the tail probabilities, however, does not provide accurate estimates when the sample size is small.

Heuristic Method

Wright [16] proposed a modification of C_{pmk} which incorporates a skewness term in the denominator to reduce the index value when nonsymmetry is present. Wright’s index C_s is defined as

$$\begin{aligned} C_s &= \frac{\min\{USL - \mu, \mu - LSL\}}{3\sqrt{\sigma^2 + (\mu - T)^2 + |\mu_3/\sigma|}} \\ &= \frac{(d - |\mu - T|)\sigma}{3\sqrt{1 + [(\mu - T)/\sigma]^2 + |\alpha_3|}} \end{aligned} \quad (16)$$

where T is the target, $d = (USL - LSL)/2$, and α_3 is skewness. C_s decreases as skewness increases. The index, however, cannot be related to process yield.

Chang *et al.* [17] proposed PCIs for skewed populations based on a weighted standard deviation (WSD) approach, which improves the performance of the weighted variance (WV) PCIs suggested by Choi and Bai [18]. The WSD method utilizes different divisors at the upper and lower limits of the specification interval, and the PCIs are defined as

$$C_p^{WSD} = \min \left\{ \frac{USL - LSL}{6\sigma \cdot 2P}, \frac{USL - LSL}{6\sigma \cdot 2(1 - P)} \right\} \quad (17)$$

$$C_{pk}^{WSD} = \min \left\{ \frac{USL - \mu}{3\sigma \cdot 2P}, \frac{\mu - LSL}{3\sigma \cdot 2(1 - P)} \right\} \quad (18)$$

where $P = \Pr\{X \leq \mu\}$. Wu *et al.* [19] proposed a new WV-based PCI calculated as

$$C_p^{WV} = \frac{USL - LSL}{3(S_1 + S_2)} \quad (19)$$

where $S_j^2 = 2 \sum_{i=1}^{n_j} (X_i - \bar{X})^2 / (2n_j - 1)$, n_1 is the number of observations less than or equal to $\bar{X}$, and n_2 is that greater than $\bar{X}$. The WV and WSD methods are based on the idea that the asymmetric distribution

can be divided into upper and lower parts from the process mean and the standard deviations above and below the mean can be calculated separately with the divided distributions.

The heuristic approach may not describe the capability of a process accurately. Therefore, if the type of a distribution is known or sufficient process data are available, experienced practitioners can estimate the capability fairly accurately with the data transformation or the curve fitting method. Earlier in the product life cycle, however, quality data are not sufficient. In such a case, the heuristic method may be adequate.

Multivariate PCI

Capability analyses involving more than one quality characteristic are sometimes of interests, and multivariate statistical techniques can be used to analyze several quality characteristics simultaneously. A difficulty of defining multivariate PCIs (*see Process Capability Indices, Multivariate*) in that there is no consensus on the methodology for assessing capability, which arises since the multivariate relationship among quality characteristics may or may not be reflected in the engineering specifications. In general, the upper and lower specification limits may be given for each quality characteristic and these specification ranges, taken together, form a rectangle or hypercube. Only a few methods applied to the multivariate skewed populations as well as multivariate normal populations have been suggested. For example, Polansky [20] proposed assessing the capability of a process using a nonparametric estimator based on a kernel estimate, and Chang and Bai [21] applied the WSD method to the modified region approach by Hubele *et al.* [22]. Further studies on multivariate PCIs for nonnormal populations are needed.

References

- [1] Somerville, S.E. & Montgomery, D.C. (1996–1997). Process capability indices and non-normal distributions, *Quality Engineering* **9**, 305–316.
- [2] Kotz, S. & Lovelace, C.R. (1998). *Process Capability Indices in Theory and Practice*, Arnold, New York.
- [3] Polansky, A.M., Chou, Y.M. & Mason, R.L. (1999). An algorithm for fitting Johnson transformations to non-normal data, *Journal of Quality Technology* **31**, 345–350.
- [4] Clements, J.A. (1989). Process capability calculations for non-normal distributions, *Quality Progress* **22**, 95–100.
- [5] Castagliola, P. (1996). Evaluation of non-normal process capability indices using Burr's distribution, *Quality Engineering* **8**, 587–593.
- [6] Ding, J. (2004). A method of estimating the process capability index from the first four moments of non-normal data, *Quality and Reliability Engineering International* **20**, 787–805.
- [7] Chang, P. & Lu, K. (1994). PCI calculations for any shape of distribution with percentile, *Quality World, Technical Supplement* **20**, 110–114.
- [8] Polansky, A.M. (1998). A smooth nonparametric approach to process capability, *Quality and Reliability Engineering International* **14**, 43–48.
- [9] Pearn, W.L. & Kotz, S. (1994). Application of Clements' method for calculating second and third generation process capability indices for non-normal Pearsonian populations, *Quality Engineering* **7**, 139–145.
- [10] Shore, H. (1998). A new approach to analyzing non-normal quality data with application to process capability analysis, *International Journal of Production Research* **36**, 1917–1933.
- [11] Shore, H. (1995). Fitting a distribution by the first two sample moments (partial and complete), *Computational Statistics and Data Analysis* **19**, 563–577.
- [12] Pearn, W.L., Kotz, S. & Johnson, N.L. (1992). Distributional and inferential properties of process capability indices, *Journal of Quality Technology* **24**, 216–231.
- [13] Rodriguez, R.N. (1992). Recent developments in process capability analysis, *Journal of Quality Technology* **24**, 176–187.
- [14] Flaig, J.J. (1996). A new approach to process capability analysis, *Quality Engineering* **9**, 205–211.
- [15] Chao, M.T. & Lin, D.K.J. (2006). Another look at the process capability index, *Quality and Reliability Engineering International* **22**, 153–163.
- [16] Wright, P.A. (1995). A process capability index sensitive to skewness, *Journal of Statistical Computation and Simulation* **52**, 195–203.
- [17] Chang, Y.S., Choi, I.S. & Bai, D.S. (2002). Process capability indices for skewed populations, *Quality and Reliability Engineering International* **18**, 383–393.
- [18] Choi, I.S. & Bai, D.S. (1996). Process capability indices for skewed populations, *Proceedings 20th International Conference on Computer and Industrial Engineering*, Kyongju, Korea, pp. 1211–1214.
- [19] Wu, H.H., Swain, J.J., Farrington, P.A. & Messimer, S.L. (1999). A weighted variance capability index for general non-normal processes, *Quality and Reliability Engineering International* **15**, 397–402.
- [20] Polansky, A.M. (2001). A smooth nonparametric approach to multivariate process capability, *Technometrics* **43**, 199–211.
- [21] Chang, Y.S. & Bai, D.S. (2003). Multivariate process capability indices for skewed populations with weighted

- standard deviations, *Journal of the Korean Institute of Industrial Engineers* **29**, 114–125.
- [22] Hubele, N.F., Shahriari, H. & Cheng, C.S. (1992). A bivariate process capability vector, in *Statistical Process Control in Manufacturing*, J.B. Keats & D.C. Montgomery, eds, Marcel Dekker, New York, pp. 299–310.

YOUNG SOON CHANG AND DO SUN BAI

Process Capability Indices, Nonnormal

Introduction

Process capability indices measure the ability of a process to manufacture product that meets specifications when the process is in statistical control. Most process capability studies assume that the process quality characteristic under study, X , is normally distributed. However, nonnormal process distributions are common in semiconductor and other manufacturing industries. For example, the particle counts on wafers and chemical impurity levels in semiconductor manufacturing often have nonnormal distributions. A one-sided specification is often an indication that the underlying distribution may be nonnormal. Normal distributions generally do not fit constrained data. Process capability indices for nonnormal distributions have received increasing attention in recent years. Books that devote some chapters to covering process capability indices for nonnormal distributions include Kotz and Johnson [1], Kotz and Lovelace [2], Bothe [3], and Pearn and Kotz [4]. A recent bibliography of process capability papers is given in Spiring *et al.* [5].

Many authors study the effects of nonnormality on process capability indices and conclude that using normal-based process capability indices such as C_p , C_{pu} , C_{pl} , and C_{pk} , or the equivalent superstructure $C_p(u, v)$ of Vännman [6], to analyze nonnormal data may lead to erroneous results. Somerville and Montgomery [7] investigate the errors incurred in using normal-based analysis to make inference about expected process fallout rate based on nonnormal data. Discussions of the effects of nonnormality are in Kotz and Johnson [8], Lu and Rudy [9], and Spiring *et al.* [10].

Some approaches to deal with nonnormality in process capability analysis include: (a) transforming nonnormal data to normality so that the normal-based process capability indices can be applied to the transformed data, (b) fitting data by a distribution so that quantile-related process capability indices can be evaluated, and (c) developing other non-quantile-based indices that are applicable to nonnormal distributions. Discussions of various approaches include [2] and [8].

Data Transformation

One approach to the nonnormality problem is to transform the data using some well-known transformations such as the **Box**–Cox power transformations and the Johnson system of transformations (see Box and Cox [11], Page [12], Farnum [13], Chou *et al.* [14] and Polansky *et al.* [15, 16]). A test of normality must be performed to check the normality of the transformed data. If the test supports normality, the normal-based process capability indices are appropriate for the transformed data. A powerful test of normality is given by Chen and Shapiro [17].

In a comparative study, Tang and Than [18] use the normal-based index C_{pu} as the standard and compare the performance of seven nonnormal-based indices when applied to lognormal and Weibull data. Their results show that the Box–Cox and Johnson transformation methods perform very well.

Quantile-Based Process Capability Indices

The second approach is to fit the data by an empirical distribution function or by a three or four-parameter distribution in order to obtain the **quantiles** of the distribution of X and define quantile-based process capability indices. The p th quantile of the distribution of X is defined to be a value x_p such that $\Pr(X \leq x_p) = p$. If X has a normal $N(\mu, \sigma^2)$ distribution, then $\mu = x_{0.5}$, $\mu - 3\sigma = x_{0.00135}$, and $\mu + 3\sigma = x_{0.99865}$.

The idea is to modify the usual (μ, σ) -related indices so that they are applicable to both normal and nonnormal distributions. In this approach, normal-based process capability indices are modified in order to increase their robustness (see [2]). For example, the modified index of C_p is given by

$$C'_p = \frac{USL - LSL}{x_{0.99865} - x_{0.00135}} \quad (1)$$

where LSL and USL are the lower and upper specification limits, respectively. Veevers [19] introduces the viability index which is equal to $1 - \frac{1}{C'_p}$ for any symmetric unimodal distribution.

Clements [20] suggests using a Pearson system distribution to fit the underlying distribution in order to determine the quantiles for calculating the process capability indices. The Pearson system has been used by Pearn and Kotz [21] and Bittanti *et al.* [22].

2 Process Capability Indices, Nonnormal

Castagliola [23] fits nonnormal data by a Burr distribution. Polansky applies a kernel smoothing method to fit data (see **Process Capability Indices, Nonparametric**). Krishnamoorthi and Khatwani [24] use a Weibull distribution to fit data. Ding [25] uses only the first four sample **moments** to fit the underlying distribution by a Chebyshev–Hermite polynomial up to the 10th order, where the original process data need not be available. Pal [26] uses the generalized lambda distribution to fit data. In fitting data by a distribution, a goodness-of-fit test must be performed to check if the distribution fits the data well.

Pearn and Chen [27] generalize the superstructure $C_p(u, v)$ and introduce quantile-based indices, $C_{Np}(u, v)$ and $C'_{Np}(u, v)$, for nonnormal distributions. The index $C_{Np}(u, v)$ is defined as

$$C_{Np}(u, v) = \frac{d - u|M - m|}{3\sqrt{\left(\frac{x_{0.99865} - x_{0.00135}}{6}\right)^2 + v(M - T)^2}} \quad (2)$$

where $M = x_{0.5}$, $m = (USL + LSL)/2$, $d = (USL - LSL)/2$, T is the target, and $u, v \geq 0$. $C'_{Np}(u, v)$ has exactly the same form as $C_{Np}(u, v)$ except that the median M is replaced by μ . Their study shows that the generalizations $C_{Np}(u, v)$ are superior to $C'_{Np}(u, v)$ in assessing process capability. Both indices are applicable to normal distributions since they reduce to $C_p(u, v)$. More details about the indices are in [4].

McCormack *et al.* [28] find that the empirical distribution function fits their process data very well for samples of size 100 and conclude that their index $\hat{C}_{Npk}$ is a viable estimator of process capability for nonnormal distributions, where

$$\hat{C}_{Npk} = \min\left(\frac{USL - x_{0.5}}{x_{0.995} - x_{0.5}}, \frac{x_{0.5} - LSL}{x_{0.5} - x_{0.005}}\right) \quad (3)$$

Other Process Capability Indices

The third approach to deal with nonnormality is to develop indices that are not quantile-based. The process capability index, C_{pm} , developed by Chan *et al.* [29], can be applied to assess process capability for both normal and nonnormal distributions since C_{pm} is not distribution sensitive (see [10]).

Luceno [30] introduces the index C_{pc} to handle the nonnormal data and estimate the index by an approximate **confidence interval**, where C_{pc} is defined as

$$C_{pc} = \frac{USL - LSL}{6\sqrt{\frac{\pi}{2}E|X - T|}} \quad (4)$$

A point estimator of C_{pc} is given by

$$\hat{C}_{pc} = \frac{USL - LSL}{6\sqrt{\frac{\pi}{2n} \sum_{i=1}^n |X_i - T|}} \quad (5)$$

Wright [31] proposes an index, C_s , to account for skewness (see **Process Capability Indices, Skewed**). C_s is defined as

$$C_s = \frac{\min(USL - \mu, \mu - LSL)}{3\sqrt{\sigma^2 + (\mu - T)^2 + |\mu_3/\sigma|}} \quad (6)$$

Bootstrap confidence intervals for C_s are given in Han *et al.* [32]. Nahar *et al.* [33] assess the performance of Wright's index C_s and a modified index C_{sm} , for positively skewed distribution, where C_{sm} is obtained from C_s by replacing $|\mu_3/\sigma|$ by $(-\mu_3/\sigma)$. They find that the modified Wright's index C_{sm} has some useful attributes in reflecting process performance with respect to the percent nonconforming and the accuracy for relatively large samples.

Wu and Swain [34] propose process capability indices based on the weighted variance method. Chang *et al.* [35] propose indices based on the weighted standard deviation method. The idea for both methods is to divide a skewed or asymmetric distribution into two normal distributions that have the same mean but different variances or standard deviations. Simulation studies indicate these methods are fairly consistent in estimating process fallout rate for some nonnormal distributions.

Perakis and Xekalaki [36, 37] propose an index $(1 - p_0)/(1 - p)$, where p and p_0 denote the proportion of conformance (yield) and the minimum allowable proportion of conformance of the process under study. They state that the index can be used for both continuous and discrete process distributions.

Chao and Lin [38] believe a process capability index must represent the true yield and propose a

universal index, C_y , defined as

$$C_y = \frac{1}{3} \Phi^{-1} \left[\frac{1}{2} F(USL) - F(LSL) + 1 \right] \quad (7)$$

where Φ is the standard normal distribution function and F is the distribution function of X . The index can be extended to a multivariate process capability index (see **Process Capability Indices, Multivariate**) when F is a multivariate distribution.

Future Research

Many process capability indices have been proposed for nonnormal distributions or data. Various methods of estimating a process capability index can lead to different estimates of the rate of nonconforming product. Kaminsky *et al.* [39] show that it is difficult to demonstrate that a good process is indeed a good process by using a point estimator. Information about confidence interval estimation of the process capability indices is generally lacking.

As pointed out in [10], there has been little research in the area of loss and loss functions as methods for assessing process capability. The potential problem of nonnormality leads researchers to develop new indices or investigate the properties of particular process capability indices by evaluating their estimators using a variety of nonnormal process data. A comprehensive study of the indices and situations where they are applicable seems necessary.

References

- [1] Kotz, S. & Johnson, N.L. (1993). *Process Capability Indices*, Chapman & Hall, London.
- [2] Kotz, S. & Lovelace, C.R. (1998). *Process Capability Indices in Theory and Practice*, Arnold, London.
- [3] Bothe, D.R. (2001). *Measuring Process Capability*, Landmark Publishing, Cedarburg.
- [4] Pearn, W.L. & Kotz, S. (2006). *Encyclopedia and Handbook of Process Capability Indices*, World Scientific Publishing.
- [5] Spiring, F., Leung, B., Cheng, S.W. & Yeung, A. (2003). A bibliography of process capability papers, *Quality and Reliability Engineering International* **19**, 445–460.
- [6] Vännman, K. (1995). A unified approach to capability indices, *Statistica Sinica* **5**, 805–820.
- [7] Somerville, S.E. & Montgomery, D.C. (1996–1997). Process capability indices and non-normal distributions, *Quality Engineering* **9**, 305–316.
- [8] Kotz, S. & Johnson, N.L. (2002). Process capability indices – a review, 1992–2000, *Journal of Quality Technology* **34**, 2–19.
- [9] Lu, M.W. & Rudy, R.J. (2002). “Discussion on process capability indices – a review, 1992–2000”, *Journal of Quality Technology* **34**, 38–39.
- [10] Spiring, F., Cheng, S.W., Yeung, A. & Leung, B. (2002). “Discussion on Process capability indices – a review, 1992–2000”, *Journal of Quality Technology* **34**, 23–27.
- [11] Box, G.E.P. & Cox, D.R. (1964). An analysis of transformations, *Journal of the Royal Statistical Society, Series B* **26**(2), 211–252.
- [12] Page, M. (1994). Analysis of non-normal process distributions, *Semiconductor International* **17**, 88–96.
- [13] Farnum, N.R. (1996–1997). Using Johnson curves to describe non-normal process data, *Quality Engineering* **9**(2), 329–336.
- [14] Chou, Y.M., Polansky, A.M. & Mason, R.L. (1998). Transforming non-normal data to normality in statistical process control, *Journal of Quality Technology* **30**, 133–141.
- [15] Polansky, A.M., Chou, Y.M. & Mason, R.L. (1999). An algorithm for fitting Johnson transformations to non-normal data, *Journal of Quality Technology* **31**, 345–350.
- [16] Polansky, A.M., Chou, Y.M. & Mason, R.L. (1998–1999). Estimating process capability indices for a truncated normal distribution, *Quality Engineering* **11**(2), 257–265.
- [17] Chen, L. & Shapiro, S.S. (1995). An alternative test for normality based on normalized spacings, *Journal of Statistical Computation and Simulation* **53**, 269–287.
- [18] Tang, L.C. & Than, S.E. (1999). Computing process capability indices for non-normal data: a review and comparative study, *Quality and Reliability Engineering International* **15**, 339–353.
- [19] Veevers, A. (1998). Viability and capability indices for multi-response processes, *Journal of Applied Statistics* **25**, 545–558.
- [20] Clements, J.A. (1989). Process capability calculations for non-normal distributions, *Quality Progress* **22**(9), 95–100.
- [21] Pearn, W.L. & Kotz, S. (1994–1995). Application of Clements’ method for calculating second- and third-generation process capability indices for non-normal Pearsonian populations, *Quality Engineering* **7**(1), 139–145.
- [22] Bittanti, S., Lovera, M. & Moiraghi, L. (1998). Application of non-normal process capability indices to semiconductor quality control, *IEEE Transactions on Semiconductor Manufacturing* **11**(2), 296–303.
- [23] Castagliola, P. (1996). Evaluation of non-normal process capability indices using Burr’s distributions, *Quality Engineering* **8**, 587–593.

- [24] Krishnamoorthi, K.S. & Khatwani, S. (2000). A capability index for all occasions, *Annual Quality Congress*, Indianapolis, Vol. 54, pp. 77–81.
- [25] Ding, J. (2004). A method of estimating the process capability index from the first four moments of non-normal data, *Quality and Reliability Engineering International* **20**, 787–805.
- [26] Pal, S. (2005). Evaluation of nonnormal process capability indices using generalized lambda distribution, *Quality Engineering* **17**(1), 77–85.
- [27] Pearn, W.L. & Chen, K.S. (1997). Capability indices for non-normal distributions with an application in electrolytic capacitor manufacturing, *Microelectronics and Reliability* **37**(12), 1853–1858.
- [28] McCormack, D.W., Harris, I.R., Hurwitz, A.M. & Spagon, P.D. (2000). Capability indices for non-normal data, *Quality Engineering* **12**(4), 489–495.
- [29] Chan, L.K., Cheng, S.W. & Spiring, F.A. (1988). A new measure of process capability: Cpm, *Journal of Quality Technology* **20**(3), 162–175.
- [30] Luceno, A. (1996). A process capability index with reliable confidence intervals, *Communications in Statistics: Simulation and Computation* **25**, 235–245.
- [31] Wright, P.A. (1995). A process capability index sensitive to skewness, *Journal of Statistical Computation and Simulation* **52**, 195–203.
- [32] Han, J., Cho, J. & Leem, C.S. (2000). Bootstrap confidence limits for Wright's Cs, *Communications in Statistics: Theory and Methods* **29**(3), 485–505.
- [33] Nahar, P.C., Hubele, N.F. & Zimmer, L.S. (2001). Assessment of a capability index sensitive to skewness, *Quality and Reliability Engineering International* **17**, 233–241.
- [34] Wu, H.H. & Swain, J.J. (2001). A Monte Carlo comparison of capability indices when processes are non-normally distributed, *Quality and Reliability Engineering International* **17**, 219–231.
- [35] Chang, Y.S., Choi, I.S. & Bai, D.S. (2002). Process capability indices for skewed populations, *Quality and Reliability Engineering International* **18**, 383–393.
- [36] Perakis, M. & Xekalaki, E. (2002). A process capability index that is based on the proportion of conformance, *Journal of Statistical Computation and Simulation* **72**(9), 707–718.
- [37] Perakis, M. & Xekalaki, E. (2005). A process capability index for discrete processes, *Journal of Statistical Computation and Simulation* **75**(3), 175–187.
- [38] Chao, M.T. & Lin, D.K.J. (2006). Another look at the process capability index, *Quality and Reliability Engineering International* **22**(2), 153–163.
- [39] Kaminsky, F.C., Dovich, R.A. & Burke, R.J. (1998). Process capability indices: now and in the future, *Quality Engineering* **10**(3), 445–453.

Related Articles

Process Capability Indices, Multivariate; Process Capability Indices, Nonparametric; Process Capability Indices, Skewed.

YOUN-MIN CHOU AND ALAN M. POLANSKY

Process Capability Indices, Nonparametric

Introduction

A manufacturing process is considered to be capable if it consistently produces items within specifications. The consistency of a process, usually assessed using control chart techniques, is typically the first concern when studying a process (*see **Control Charts, Overview***). Once consistency has been established, a capability study then focuses on the ability of the process to produce items within specifications. While the primary focus of such a study is usually on the **process yield**, process capability indices are generally the measure used to study the capability of a process. These indices are typically unitless ratios that compare the size of the specification set to the natural variability of the process. Interpretation of the indices and their relationship to the process yield are usually reliant on specific assumptions about the distribution of the quality characteristic. The most common assumption is that the distribution of the quality characteristic follows a **normal distribution**, though indices related to other parametric distributions have been developed.

Many researchers have carefully studied examples of process data and have generally concluded that normal process data may be rare in practice [1–5]. It has also been shown that violations of the normal assumption can have a large effect on the ability of the process capability indices to properly indicate the actual capability of the process [6–8]. To address this problem, many indices have been developed to measure the capability of processes whose quality characteristics follow specific nonnormal distributions, such as the exponential and Weibull. Other indices have been developed that use more flexible families of distributions such as the Burr, Johnson, and Pearson systems (*see **Process Capability Indices, Nonnormal***). These indices are not strictly nonparametric indices in the sense that one is tacitly assuming that the distribution of the quality characteristic is within the family of distributions. Still other indices have implemented penalties for **skewness** (*see **Process Capability Indices, Skewed***).

A nonparametric process capability index is one that properly indicates the capability of a process without making specific assumptions about the form of the underlying distribution of the quality characteristic. Some relatively minor assumptions are usually made. For example, it is usually assumed that the distribution of the quality characteristic has some smoothness properties and has a certain number of finite **moments**. This article will focus on these types of methods for assessing the capability of a process. The first part of the article discusses nonparametric process capability indices. The second part of the article discusses nonparametric methods for evaluating the process yield. Alternative types of process capability indices that are a function of the process yield have been suggested by several researchers. The final part of this article addresses issues related to formal statistical inference using these methods.

Capability Indices

Consider a manufacturing process that produces units with a univariate quality characteristic X that follows an unknown distribution F with $\mu = E(X)$ and $\sigma^2 = V(X)$. Assume that the specification set for X is an interval of the form $S = (L, U) \subset \mathbb{R}$. In the usual case, where F is unknown, we consider observing a sample $X_1, \dots, X_n$ from F by randomly sampling items from the process. A typical process capability index compares the allowable variability of the process, given by the size of the specification set, to the actual process variability as a ratio. A simple measure of the allowable process variability is given by $U - L$, regardless of the form of F . The actual process variability is related to the variability of F , which is usually measured as $\theta\sigma$, where θ is chosen so that $P(X - \theta\sigma < X < \mu + \theta\sigma) \simeq 1$ [3]. Therefore, the general form of the process capability index is $C_\theta = (U - L)/\theta\sigma$, where σ is often estimated using the sample standard deviation s . When F is normal, the usual choice is $\theta = 6$. The more universal choice of $\theta = 5.15$ has also been suggested by [3]. A nonparametric index can be formed by replacing $\theta\sigma$ with the width of a 99.73% distribution-free tolerance interval for F [9]. Note that these approaches are direct analogs to the C_p (process capability index), which assumes that the process mean is centered in the specification set.

2 Process Capability Indices, Nonparametric

Otherwise, these indices merely indicate the best possible capability of the process. See Section 2.2 of [10].

The standard deviation may not be an appropriate measure of variability for all process distributions. A measure based on the **percentiles** of F has been suggested by [11]. The analog of the C_p process capability index they propose has the form

$$\hat{C}_{np} = \frac{U - L}{\hat{\xi}_{99.5} - \hat{\xi}_{0.5}} \quad (1)$$

where $\hat{\xi}_\eta$ is the η th sample percentile of the observed process data. The same authors suggest an analog of the C_{pk} index, which accounts for the location of the process distribution. This index is given by $\hat{C}_{npk} = \min\{\hat{C}_{npl}, \hat{C}_{npu}\}$, where

$$\hat{C}_{npl} = \frac{\hat{\xi}_{50.0} - L}{\hat{\xi}_{50.0} - \hat{\xi}_{0.5}} \quad (2)$$

and

$$\hat{C}_{npu} = \frac{U - \hat{\xi}_{50.0}}{\hat{\xi}_{99.5} - \hat{\xi}_{50.0}} \quad (3)$$

An empirical study presented in [11] compared the performance of the usual $\hat{C}_{pk}$ process capability index and the proposed $\hat{C}_{npk}$ index. These simulations showed that the **bias** of the $\hat{C}_{npk}$ index was lower than that of the $\hat{C}_{pk}$ index for nonnormal distributions. Similar proposals have been suggested by [12–14].

Another modification of the C_p index is suggested by [15]. The sample form of this index is

$$\hat{C}_{pc} = \frac{U - L}{6\bar{c}\sqrt{\pi/2}} \quad (4)$$

where

$$\bar{c} = n^{-1} \sum_{i=1}^n |X_i - (U + L)/2| \quad (5)$$

The choice of the denominator of this index takes advantage of the quick convergence properties of the central limit theorem so that reliable **confidence intervals** for this index are simple to compute, regardless of the form of F . While the process yield associated with this index does depend on F , it is shown by [15] that the relationship is more stable than the relationship between the C_p process capability index and the process yield.

Indices Based on the Process Yield

The process yield is defined to be $p = P(L \leq X \leq U)$, the expected proportion of items in a sample that are within specification. Several authors have suggested that process capability indices that are based on estimates of the process yield may be more useful, particularly for nonnormal process data [16–20]. The advantages of this approach are that the interpretation of the process yield does not depend on the process distribution, the indices easily generalize to the case of a multivariate quality characteristic, and that the index can be applied to discrete processes [21]. For example [20] suggests a process capability index of the form

$$C_{pc} = \frac{1 - p_0}{1 - p} \quad (6)$$

where p_0 is the minimum allowable process yield. The process fallout $q = 1 - p$ can be used as an equivalent basis for these indices. Let $q_L = P(X < L)$ and $q_U = P(X > U)$. Then [22] suggests a process capability index of the form

$$C_f = \min \left\{ \frac{q_L^0}{q_L}, \frac{q_U^0}{q_U} \right\} \quad (7)$$

where q_L^0 and q_U^0 are the maximum allowable process fallout rates below and above the specifications, respectively.

In order to implement these indices in the nonparametric setting, a reliable nonparametric method for estimating the process yield p must be used. A simple nonparametric approach consists of estimating the process yield with the proportion of items in the sample that are within specification, that is,

$$\hat{p} = n^{-1} \sum_{i=1}^n \delta(L \leq X_i \leq U) \quad (8)$$

where δ is the indicator function. Confidence limits for $\hat{p}$ are easily computed using the binomial distribution [23]. Such an interval is usually quite wide unless the sample size is very large, as the method does not take advantage of information in the sample about the form of F . Methods for computing upper bounds on p based on a modification of Tchebysh-eff's theorem have been considered by [24]. The use of extreme value theory has been suggested by [22]. Note that if F is continuous then $p = F(U) - F(L)$.

This implies that the process yield can also be estimated by replacing F with a nonparametric estimate, such as a kernel estimate [25, 26]. Bandwidth selection for this estimate can be achieved using the plug-in method of [27] or the cross validation technique of [28]. Another approach suggested by [29] takes advantage of the nearly pivotal transformation $(x - \bar{x})/s$ and the **bootstrap** to develop a nonparametric approach to estimating the process fallout rate. Using the fact that a density can be written as a series of Chebyshev–Hermite polynomials [30] derive a method for estimating process yield based only on the first four sample moments.

The Multivariate Case

The case of a multivariate quality characteristic presents challenges for any measure of process capability. This is particularly true for nonparametric process capability indices. Let $\mathbf{X}$ be a d -dimensional quality characteristic and let $S \subset \mathbb{R}^d$ be the specification set for the characteristic. It will be assumed that $\mathbf{X}$ follows a d -dimensional distribution F . In the usual sample case, $\mathbf{X}_1, \dots, \mathbf{X}_n$ will denote the quality characteristics observed from items sampled from the process.

A multivariate process capability index has been developed by [31] that is directly related to the process fallout rate. This approach requires a specification set of the form

$$S = \{\mathbf{x} \in \mathbb{R}^d : h(\mathbf{x} - \mathbf{T}) \leq r\} \quad (9)$$

where $h : \mathbb{R}^d \rightarrow \mathbb{R}^+$, $\mathbf{T}$ is a target vector, and $r > 0$ is a constant. Such a specification set is general enough to include ellipsoidal and rectangular specification sets. In this case, the process fallout rate is $q = P[h(\mathbf{X} - \mathbf{T}) > r]$, and [31] defines the multivariate process capability index as

$$MC_p = \frac{q_0}{q} \quad (10)$$

where q_0 is the maximum acceptable fallout rate. This index is generalized to the nonparametric setting by [32], who use extreme value theory to develop a nonparametric method for estimating q based on $\mathbf{X}_1, \dots, \mathbf{X}_n$. The process fallout rate can also be estimated in the nonparametric setting using a multivariate kernel density estimate. With this method, one must assume that F has an absolutely

continuous d -dimensional density f so that the process fallout rate can be written as

$$q = 1 - \int_S f(\mathbf{x}) d\mathbf{x} \quad (11)$$

The fallout rate is then estimated by replacing f with a multivariate kernel density estimate. The bandwidth for this estimate must be specifically tailored to the specific shape of S . A bootstrap method for selecting the bandwidth matrix for the kernel estimate based on observations $\mathbf{X}_1, \dots, \mathbf{X}_n$ is given by [33].

Statistical Inference

A complete plan for studying the capability of a process should include methods for reliable statistical inference, including confidence intervals and tests of hypotheses. Most of the nonparametric methods developed so far have dealt exclusively with **estimation**. The form of the estimates and the lack of assumptions of the distribution of the quality characteristic generally exclude the possibility of simple closed form inferential methods for the case of finite sample sizes. Asymptotic inference for the indices may also be problematic due to the complexity of the indices. One exception is the index proposed by [15] which was specifically designed for simple inference. In general, most of the methods discussed here may require the application of the nonparametric approaches to inference afforded by such methods as the jackknife and the bootstrap.

References

- [1] Dovich, R.A. (1987). CPI, C_{pk} , capability ratio – measures of performance or contributors to confusion? *Machine and Tool Bluebook*, Hitchcock Publishing Company, Wheaton, Illinois, USA, pp. 10–14.
- [2] Gunter, B.H. (1989). The use and abuse of C_{pk} , *Quality Progress* **22**(1), 72–73.
- [3] Gunter, B.H. (1989). The use and abuse of C_{pk} , part 2, *Quality Progress* **22**(3), 108–109.
- [4] McCoy, P.F. (1991). Using performance indices to monitor production processes, *Quality Progress* **24**(2), 49–55.
- [5] Pyzdek, T. (1995). Why normal distributions aren't [all that normal], *Quality Engineering* **7**, 769–777.
- [6] Chan, L.K., Cheng, S.W. & Spiring, F.A. (1988). The robustness of the process capability index, C_p , to departures from normality, *Statistical Theory and Data Analysis II*, Elsevier Science Publishers.

- [7] Somerville, S.E. & Montgomery, D.C. (1996). Process capability indices and non-normal distributions, *Quality Engineering* **9**, 305–316.
- [8] Pearn, W.L., Kotz, S. & Johnson, N.L. (1992). Distributional and inferential properties of process capability indices, *Journal of Quality Technology* **24**, 216–231.
- [9] Chan, L.K., Cheng, S.W. & Spiring, F.A. (1988). A graphical technique for process capability, *ASQC Quality Congress Transactions - Dallas* **42**, 268–278.
- [10] Kotz, S. & Johnson, N.L. (1993). *Process Capability Indices*, Chapman & Hall, London.
- [11] McCormack, D.W., Harris, I.R., Hurwitz, A.M. & Spagon, P.D. (2000). Capability indices for non-normal data, *Quality Engineering* **12**, 489–495.
- [12] Schneider, H.J., Pruett, J. & Lagrange, C. (1995). Uses of process capability indices in the supplier certification process, *Quality Engineering* **8**, 225–235.
- [13] Pearn, W.L. & Chen, K.S. (1997). Capability indices for non-normal distributions with an application in electrolytic capacitor manufacturing, *Microelectronics and Reliability* **37**(12), 1853–1858.
- [14] Tang, L.C. & Than, S.E. (1999). Computing process capability indices for non-normal data: a review and comparative study, *Quality and Reliability Engineering International* **15**, 339–353.
- [15] Luceño, A. (1996). A process capability index with reliable confidence intervals, *Communications in Statistics: Simulation and Computation* **25**, 235–245.
- [16] Carr, W.E. (1991). A new process capability index: parts per million, *Quality Progress* **24**, 152.
- [17] Chao, M.T. & Lin, D.K.J. (2006). Another look at the process capability index, *Quality and Reliability Engineering International* **22**, 153–163.
- [18] Constable, G.K. & Hobbs, J.R. (1992). Small samples and non-normal capability, *ASQC Quality Congress Transactions - Nashville* **46**, 37–43.
- [19] Kaminsky, F.C., Dovich, R.A. & Burke, R.J. (1998). Process capability indices: now and in the future, *Quality Engineering* **10**, 445–453.
- [20] Perakis, M. & Xekalaki, E. (2002). A process capability index that is based on the proportion of conformance, *Journal of Statistical Computation and Simulation* **72**, 707–718.
- [21] Perakis, M. & Xekalaki, E. (2005). A process capability index for discrete processes, *Journal of Statistical Computation and Simulation* **75**, 175–187.
- [22] Yeh, A.B. & Bhattacharya, S. (1998). A robust process capability index, *Communications in Statistics: Simulation and Computation* **27**, 565–589.
- [23] Hollander, M. & Wolfe, D.A. (1999). *Nonparametric Statistical Methods*, 2nd Edition, John Wiley & Sons, New York, pp. 31–33.
- [24] Flaig, J.J. (1996). A new approach to process capability analysis, *Quality Engineering* **9**, 205–211.
- [25] Polansky, A.M. (1998). A smooth nonparametric approach to process capability, *Quality and Reliability Engineering International* **14**, 38–43.
- [26] Rodriguez, R.N. (1992). Recent developments in process capability analysis, *Journal of Quality Technology* **24**, 176–187.
- [27] Polansky, A.M. & Baker, E.R. (2000). Multistage plug-in bandwidth selection for kernel distribution function estimates, *Journal of Statistical Computation and Simulation* **65**, 63–80.
- [28] Bowman, A., Hall, P. & Prvan, T. (1998). Bandwidth selection for the smoothing of distribution functions, *Biometrika* **85**, 799–808.
- [29] Ciarlini, P., Gigli, A. & Regoliosi, G. (1999). The computation of accuracy of quality parameters by means of Monte Carlo simulation, *Communications in Statistics: Simulation and Computation* **28**, 821–848.
- [30] Ding, J. (2004). A method of estimating the process capability index from the first four moments of non-normal data, *Quality and Reliability Engineering International* **20**, 787–805.
- [31] Chen, H. (1994). A multivariate process capability index over a rectangular solid tolerance zone, *Statistica Sinica* **4**, 749–758.
- [32] Yeh, A.B. & Chen, H. (2001). A non-parametric multivariate process capability index, *International Journal of Modelling and Simulation* **21**, 218–224.
- [33] Polansky, A.M. (2001). A smooth nonparametric approach to multivariate process capability, *Technometrics* **43**, 199–211.

Related Articles

Process Capability Indices, Comparison of; Process Capability Indices, Multivariate; Process Capability Indices, Nonnormal; Process Capability Indices, Skewed; Process Yield.

ALAN M. POLANSKY

Process Capability Indices, Multivariate

Basic Capability Indices

Process capability indices are becoming a common approach to report process performance, both internally and in managing customer–supplier relationships. In the univariate case, *capability indices* are defined as ratios of spread between the process specification limits to variability of the process. For the univariate case, C_p , C_{pk} , C_{pm} , and C_{pmk} are the most commonly known *process capability indices* [1–3]. Most process capability indices assume that the process performance data are normally distributed. Let μ and σ be the mean and standard deviation, respectively, of a process performance characteristic such as weight, assay, particle size, or voltage levels. Furthermore let USL, LSL, and T be the upper and lower specification limits and the target value, respectively. The above mentioned population capability indices are defined as follows:

$$C_p = \frac{USL - LSL}{6\sigma} \quad (1)$$

$$C_{pk} = \min \left[\frac{USL - \mu}{3\sigma}, \frac{\mu - LSL}{3\sigma} \right] \quad (2)$$

$$C_{pm} = \frac{USL - LSL}{6\sqrt{\sigma^2 + (\mu - T)^2}} \quad (3)$$

$$C_{pmk} = \min \left[\frac{USL - \mu}{3\sqrt{\sigma^2 + (\mu - T)^2}}, \frac{\mu - LSL}{3\sqrt{\sigma^2 + (\mu - T)^2}} \right] \quad (4)$$

with the sample estimates being obtained by replacing the population μ and σ by their sample estimates, $\bar{x}$ and s , respectively.

General Multivariate Capability Indices

However, many industrial processes have more than one performance or quality characteristic so that multivariate evaluation of process performance becomes more and more important. Although univariate process capability indices have been extensively studied in the literature, their multivariate counterparts have received relatively little attention.

Taam *et al.* [4] and later Pal [5] proposed the index $C_{pb} = S_R/A_P$, where S_R represents the surface defined by the rectangle of the upper and lower tolerances and A_P is the process surface including 99.73% of the process data. Wang and Chen [6] proposed multivariate equivalents for C_p , C_{pk} , C_{pm} , and C_{pmk} based on principal component analysis decomposition. Wang and Du [7] proposed the same indices with extension to the nonnormal multivariate case. Wang *et al.* [8] proposed a process capability multivariate vector to evaluate the process performance. Such a vector has the following components: a multivariate C_p index, the p value of a T^2 Hotelling test, and a location index. Chen [9] also investigated a method to estimate the multivariate C_p using a nonconforming proportion approach. Castagliola and Castellanos [10] propose a method for estimating both C_p and C_{pk} indices for two quality variables labeled BC_p . The proposed bivariate C_{pk} index is calculated by first finding four probabilities corresponding to the theoretical fraction of values that exceed the surface formed by the process tolerances (nonconforming proportion). They then use a numerical quadrature method based on Green's theorem for computing the probabilities associated with each convex polygon. The capability index BC_p is found by maximizing the capability index BC_{pk} , i.e. by maximizing the theoretical proportion of conforming items.

A general formulation of multivariate capability indices is based on the region $L \leq g(\mathbf{X}) \leq U$ where $g(\mathbf{X})$ is a monotonic function of the joint **probability density function** of $\mathbf{X}$ and often, L is zero. If $\mathbf{X}$ is assumed to have a multivariate normal $N_v(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ distribution a natural $g(\mathbf{X})$ function is $g(\mathbf{X}) = (\mathbf{X} - \boldsymbol{\mu})\boldsymbol{\Sigma}^{-1}(\mathbf{X} - \boldsymbol{\mu})$ and one may regard an item as nonconforming if $(\mathbf{X} - \boldsymbol{\mu})\boldsymbol{\Sigma}^{-1}(\mathbf{X} - \boldsymbol{\mu}) > U$. We thus obtain the ellipsoidal specification region $(\mathbf{X} - \boldsymbol{\mu})\boldsymbol{\Sigma}^{-1}(\mathbf{X} - \boldsymbol{\mu}) \leq U$.

An analog of C_p is

$$\frac{\text{Volume of } \{(\mathbf{X} - \boldsymbol{\mu})\boldsymbol{\Sigma}^{-1}(\mathbf{X} - \boldsymbol{\mu}) \leq U\}}{\text{Volume of } \{(\mathbf{X} - \boldsymbol{\mu})\boldsymbol{\Sigma}^{-1}(\mathbf{X} - \boldsymbol{\mu}) \leq R\}} \quad (5)$$

where $\Pr[(\mathbf{X} - \boldsymbol{\mu})\boldsymbol{\Sigma}^{-1}(\mathbf{X} - \boldsymbol{\mu}) \leq R] = 1 - p$.

Again, if $\mathbf{X}$ has a multivariate **normal distribution** $(\mathbf{X} - \boldsymbol{\mu})\boldsymbol{\Sigma}^{-1}(\mathbf{X} - \boldsymbol{\mu})$ has a χ^2 distribution with ν **degrees of freedom** and $R = \chi_{\nu, 1-p}^2$ is the corresponding upper 100(1 - p)% point of the χ^2 distribution with ν degrees of freedom. For more examples of such indices see Wierda [11], Chan *et al.*

[12], Kotz and Johnson [2], Foster *et al.* [13], Pearn *et al.* [14], Pearn [15], and Pearn and Kotz [16].

Capability Indices Based on Tolerance Regions

Tolerance regions provide an alternative to the $L \leq g(\mathbf{X}) \leq U$ region, that is based on estimates of **percentiles** from a multivariate distribution with parameters either known or estimated from the data.

Fuchs and Kenett [17] suggest the establishment of a process control schema based on tolerance regions by estimating the level set $\{f \geq c\}$ with a prespecified probability content $1 - \alpha$, of the density f which generates the data. A reasonable criterion would be to reject an additional observation $\mathbf{X}_{n+1}$ whenever $\mathbf{X}_{n+1} \in \{f < c\}$, which would give an exact false alarm probability α . Since f is usually unknown, the population tolerance region $\{f \geq c\}$ can be estimated by plugging in an estimator function of f . For the multivariate normal case, Krishnamoorthy and Mathew [18] provide various estimates and clarify an inconsistency in notation adopted by John [19] and Fuchs and Kenett [17]. Baillo and Cuevas [20] develop a nonparametric solution of the procedure suggested by Fuchs and Kenett [17]. Di Bucchianico *et al.* [21] develop a nonparametric procedure for constructing multivariate tolerance regions that relies on the choice of a class of sets to which the tolerance region belongs (for example, ellipsoids or hyperrectangles). Polansky [22] also uses the tolerance regions approach to evaluate the capability of a manufacturing process.

Thus, a natural extension of process capability measures to the multivariate case builds on the concept of tolerance regions. This area is an active research area and will probably remain so in the next few years. Applying tolerance regions to the general formulation of **multivariate process capability indices** offers new definitions of process capability indicators.

References

- [1] Kotz, S. & Johnson, N.L. (1993). *Process Capability Indices*, Chapman & Hall, London.
- [2] Kotz, S. & Johnson, N.L. (2002). Process capability indices – a review 1992–2000 with discussion, *Journal of Quality Technology* **34**, 2–53.
- [3] Kenett, R. & Zacks, S. (1998). *Modern Industrial Statistics: Design and Control of Quality and Reliability*, 2nd Edition, Chinese edition 2004, Duxbury Press, San Francisco, 2003.
- [4] Taam, W., Subbaiah, P. & Liddy, W.L. (1993). A note on multivariate capability indices, *Journal of Applied Statistics* **20**(3), 339–351.
- [5] Pal, S. (1999). Performance evaluation of a bivariate normal process, *Quality Engineering* **11**(3), 379–386.
- [6] Wang, F.K. & Chen, J.C. (1998). Capability index using principal components analysis, *Quality Engineering* **11**(1), 21–27.
- [7] Wang, F.K. & Du, T.C.T. (2000). Using principal component analysis in process performance for multivariate data, *Omega: The International Journal of Management Science* **28**, 185–194.
- [8] Wang, F.K., Hubele, N.F., Lawrence, F.P., Miskulin, J.D. & Shahriari, H. (2000). Comparison of three multivariate process capability indices, *Journal of Quality Technology* **32**(3), 263–275.
- [9] Chen, H. (1994). A multivariate process capability index over a rectangular solid zone, *Statistica Sinica* **4**, 749–758.
- [10] Castagliola, P. & Castellanos, J.V. (2005). Capability indices dedicated to the two quality characteristics case, *Quality Technology and Quantitative Management* **2**(2), 201–220.
- [11] Wierda, S. (1994). *Multivariate Statistical Process Control*, Groningen Theses in Economics, Management and Organization, Wolters-Noordhoff, Groningen.
- [12] Chan, L.K., Cheng, S.W. & Spiring, F.A. (1991). A multivariate measure of process capability, *International Journal of Modelling and Simulation* **11**, 1–6.
- [13] Foster, E.J., Barton, R.R., Gautam, N., Truss, L.T. & Tew, J.D. (2005). The process-oriented multivariate capability index, *International Journal of Production Research* **43**(10), 2135–2148.
- [14] Pearn, W.L., Kotz, S. & Johnson, K.L. (1992). Distribution and inferential properties of process capability indices, *Journal of Quality Technology* **24**, 216–231.
- [15] Pearn, W.L. (1998). New generalization of process capability index Cpk, *Journal of Applied Statistics* **25**(6), 801–810.
- [16] Pearn, W.L. & Kotz, S. (2006). *Encyclopedia and Handbook of Process Capability Indices: A Comprehensive Exposition of Quality Control Measures, Series on Quality, Reliability and Engineering Statistics*, World Scientific Publishing Company.
- [17] Fuchs, C. & Kenett, R.S. (1987). Multivariate tolerance regions and F-tests, *Journal of Quality Technology* **19**, 122–131.
- [18] Krishnamoorthy, K. & Mathew, T. (1999). Comparison of approximation methods for computing tolerance factors for a multivariate normal population, *Technometrics* **41**(3), 234–249.

- [19] John, S. (1963). A tolerance region for multivariate normal distributions, *Sankhya, Series A* **25**, 363–368.
- [20] Baillo, A. & Cuevas, A. (2006). Parametric versus nonparametric tolerance regions in detection problems, *Computational Statistics* **21**, 523–536.
- [21] Di Buccianico, A., Einmahl, J. & Mushkudiani, N. (2001). Smallest nonparametric tolerance regions, *Annals of Statistics* **29**, 1320–1343.
- [22] Polansky, A.M. (2001). A smooth nonparametric approach to multivariate process capability, *Technometrics* **43**, 199–211.
- Kane, V.E. (1986). Process capability indices, *Journal of Quality Technology* **18**(1), 41–52.
- Kotz, S. & Lovelace, C.R. (1998). *Introduction to Process Capability Indices: Theory and Practice*, Arnold, London.
- Mason, R.L. & Young, J.C. (2002). *Multivariate Statistical Process Control with Industrial Applications*, ASA-SIAM Series on Statistics and Applied Probability, ASA-SIAM Philadelphia.
- Palmer, K. & Tsui, K.-L. (1999). A review and interpretations of process capability indices. *Annals of Operations Research* **87**(1), 31–47.
- Vannman, K. & Kotz, S. (1995). A superstructure of capability indices-distributional properties and implications, *Scandinavian Journal of Statistics* **22**, 477–491.
- Yeh, A.B. & Chen, H. (2001). A nonparametric multivariate process capability index, *International Journal of Modelling and Simulation* **21**(3), 218–223.

Further Reading

- Fuchs, C. and Kenett, R. (1998). *Multivariate Quality Control: Theory and Application*, *Quality and Reliability Series*, Vol. 54, Marcel Dekker, New York.
- Hubele, N.F., Shahriari, H. & Cheng, C.S. (1991). A bivariate process capability vector, in *Statistical Process Control in Manufacturing*, J.B. Keats & D.C. Montgomery, eds, Marcel Dekker, New York, pp. 299–310.

CAMIL FUCHS AND RON S. KENETT

Multivariate Process Capability Indices, Comparison of

Introduction

Quality measures can be used to evaluate a process's performance. In general, there are several issues that need to be discussed in assessing product quality. One important issue is the process capability index (PCI). A PCI is a numerical summary that compares the behavior of a product or process characteristic to engineering specifications. The capability indices, C_p , C_{pk} , and C_{pm} , are widely used to evaluate process performance based on a single engineering specification. However, products usually have two or more quality characteristics in processes. Thus, these univariate capability indices do not satisfy this situation, and multivariate methods for assessing process capability are needed. Process capability and related analysis have been studied for well over 20 years. Kotz and Johnson [1] provided a compact survey and brief interpretations and comments on some 170 publications on PCIs from 1992 to 2000.

In general, multivariate PCIs can be divided into two categories: the ratio of the tolerance region to the process region and the **process yield** index. In the following section, several multivariate PCIs which were proposed by Chan *et al.* [2], Pearn *et al.* [3], Taam *et al.* [4], Chen [5], Shahriari *et al.* [6], Wang and Du [7], Walter *et al.* [8], and Pearn *et al.* [9] will be reviewed. In the third section, an example is used to compare the performance of different multivariate PCIs. The conclusion is made in the final section.

Multiple Characteristics Processes and Multivariate Process Capability Indices

Multiple Characteristics Processes

In most processes, products usually have multiple quality characteristics. Each of these characteristics must satisfy its specification. So, the assessed quality of a product depends on the combined effects of these characteristics, rather than on their individual values. With respect to the tolerance region of multiple characteristics, it can be a rectangular tolerance

region for two-dimensional cases. In higher dimensions, they form a hypercube. For more complex engineering specifications, the tolerance region will be more complicated.

In most **multivariate capability indices**, it is usually assumed that the observations X have a multivariate **normal distribution**, $N_v(\mu, \Sigma)$, where v is the number of variables, μ is the mean vector, and Σ represents the variance–**covariance** matrix of X . Also, T is the target vector, upper specification limits (USL) and lower specification limits (LSL) are the specification limits, $\bar{X}$ is the sample mean vector, and S is the sample covariance matrix. For convenience, the assumption and notations are used in the following subsection.

Multivariate Process Capability Indices

Chan *et al.* [2] introduced a new version of the multivariate index C_{pm} , which is defined by

$$C_{pm} = \frac{\sqrt{\frac{nv}{n}}}{\sqrt{\sum_{i=1}^n (X_i - T)' A^{-1} (X_i - T)}} \quad (1)$$

This value indicates the **degrees of freedom** associated, while the denominator represents the sum of the observed Mahalanobis distances from the target [10]. Smaller values of C_{pm} suggest that the process is unable to meet the specifications, while larger values of C_{pm} suggest that the process is indeed capable of meeting its specifications.

Pearn *et al.* [3] introduced two multivariate PCIs ${}_vC_p$ and ${}_vC_{pm}$, which are a generalization of C_p and C_{pm} , respectively. They are defined as

$${}_vC_p = \frac{K}{\theta \chi_{v,0.9973}} \quad (2)$$

and

$${}_vC_{pm} = \frac{{}_vC_p}{\left(1 + \frac{(\mu - T)' A^{-1} (\mu - T)}{v}\right)^{1/2}} \quad (3)$$

where the region R is defined by $(X - T)' A^{-1} (X - T) \leq K^2$ and $V_0 = \theta^2 A$.

Taam *et al.* [4] proposed a multivariate capability index defined as a ratio of two volumes: $MC_{pm} =$

2 Multivariate Process Capability Indices, Comparison of

$\frac{\text{vol}(R_1)}{\text{vol}(R_2)} = \frac{\text{vol}(\text{modified tolerance region})}{\text{vol}((X - \mu)' \Sigma_T^{-1} (X - \mu) \leq k(q))}$, where R_1 is a modified tolerance region and R_2 is a scaled 99.73% process region from a multivariate normal assumption. The modified tolerance region is the largest ellipsoid that is centered at the target completely within the original tolerance region. The estimate for MC_{pm} can be expressed as

$$\begin{aligned} \widehat{MC}_{pm} &= \frac{\widehat{C}_p}{\widehat{D}} \\ &= \frac{\text{vol}(\text{modified tolerance region})}{|S|^{1/2} (\pi k)^{v/2} [\Gamma(v/2 + 1)^{-1}]} \left[1 + \frac{n}{n-1} (\bar{X} - T)' S^{-1} (\bar{X} - T) \right]^{1/2} \quad (4) \end{aligned}$$

where K is the 99.73% percent quantile of a χ^2 distribution and $|\cdot|$ denotes the determinant. The numerator $\widehat{C}_p$ is analogous to C_p , that is, a value greater than 1 implies that the process has a smaller variation than allowed by specification limits with a certain confidence level; a value less than 1 implies more variation. Similarly, $0 < 1/\widehat{D} < 1$ measures the closeness between the process mean and the target; a larger $1/\widehat{D}$ implies the mean is closer to target.

Chen [5] proposed a multivariate process index over a tolerance zone. A tolerance zone is expressed as $V = \{X \in R^v : h(X - \mu_0) \leq r_0\}$, where r_0 is a positive number and $h(X - \mu_0)$ is a cumulative distribution function. A process is capable if $P(X \in V) \geq 1 - \alpha$, where α is the allowable expected proportion of nonconforming production from a process. Let $r = \min\{c : P(h(X - \mu_0) \leq c) \geq 1 - \alpha\}$. If the cumulative distribution function of $h(X - \mu_0)$ is increasing in a neighborhood of r , then r is simply the unique root of the equation $P(h(X - \mu_0) \leq r) = 1 - \alpha$. The process is deemed capable if $r \leq r_0$, that is, $(r_0/r) \geq 1$. The multivariate process capability is defined as

$$MC_p = \frac{r_0}{r} \quad (5)$$

Shahriari *et al.* [6] proposed a multivariate capability vector, which consists of three components (C_{pM} , PV , LI). The first component of the vector is a ratio of areas or volumes, analogous to the ratio of lengths of the univariate C_p index. The denominator is the area (two-dimensional case) or the volume (three or more dimensions) defined by the engineering tolerance region, which is defined

as $C_{pM} = \left[\frac{\text{volume of engineering tolerance}}{\text{volume of modified process region}} \right]^{1/v}$. The approach forms a “modified process region” by drawing the smallest rectangle around the ellipse. The edges of the rectangle are determined by solving the equations of the first derivatives, with respect to each x_i , of the quadratic form $(X - \mu_0)' \Sigma (X - \mu_0) = \chi_{v,\alpha}^2$. Two solutions to this equation are given in the following:

$$\begin{aligned} UPL_i &= \mu_i + \sqrt{\frac{\chi_{v,\alpha}^2 \det(\Sigma_i^{-1})}{\det(\Sigma^{-1})}}, \\ LPL_i &= \mu_i - \sqrt{\frac{\chi_{v,\alpha}^2 \det(\Sigma_i^{-1})}{\det(\Sigma^{-1})}}, \quad \forall \dots i = 1, 2, \dots, v \end{aligned} \quad (6)$$

where $\chi_{v,\alpha}^2$ is the upper 100 $\alpha\%$ of a χ^2 distribution with v degrees of freedom associated with the probability contour, and $\det(\Sigma_i^{-1})$ is the determinant of Σ_i^{-1} , which is a matrix obtained from Σ^{-1} by deleting the i th row and column. As a result, the equation can be expressed as

$$C_{pM} = \left[\frac{\prod_{i=1}^v (USL_i - LSL_i)}{\prod_{i=1}^v (UPL_i - LPL_i)} \right]^{1/v} \quad (7)$$

If the value is greater than 1, then it shows that the circumscribed modified process region is smaller than the engineering specified region, that it is “acceptable”. Necessarily, the modified process region is influenced by the shape of the elliptical contour and the size of the contour. The second component is the **significance level** of the Hotelling’s T^2 statistic ($T^2 = n(\bar{X} - \mu)' S^{-1} (\bar{X} - \mu)$) which is expressed as

$$PV = Pr \left(F_{v,n-v} > \frac{n-v}{v(n-1)} T^2 \right) \quad (8)$$

where $F_{v,n-v}$ denotes a variable having the F distribution with $v, n-v$ degrees of freedom. The PV value will never exceed 1, and it indicates that the center of process is far from the engineering target value when the value is close to 0. The third component summarizes a comparison of the location of the modified process region and the tolerance region. It

has a value 1 if the entire modified process region is contained within the tolerance region and, otherwise, a value of 0. That is,

$$LI = \begin{cases} 1 & \text{if modified process region is contained} \\ & \text{within the tolerance} \\ 0 & \text{otherwise} \end{cases} \quad (9)$$

Wang and Du [7] proposed a very useful method using principal component analysis (PCA) in process performance for multivariate data. By using the spectral decomposition, we can get a matrix $D = U^T S U$, where D is a diagonal matrix. The elements of D , $\lambda_1, \lambda_2, \dots, \lambda_v$, are the eigenvalues of S and the columns of U , $u_1, u_2, \dots, u_v$ are the eigenvectors of S . So the i th principal component is expressed as $PC_i = u_i^T x$, $i = 1, 2, \dots, v$, where x s are $v \times 1$ vectors of original observations on variables. The engineering specifications and target values of PC_i s are $\begin{cases} LSL_{PC_i} = u_i^T LSL \\ USL_{PC_i} = u_i^T USL \end{cases} \forall i = 1, 2, \dots, v$. Similarly, $\begin{cases} T_{PC_i} = u_i^T T \\ S^2_{PC_i} = \lambda_i \\ \bar{X}_{PC_i} = u_i^T \bar{X} \end{cases} \forall i = 1, 2, \dots, v$. Now,

where

$$\hat{C}_{p;PC_i} = \frac{(USL_{PC_i} - LSL_{PC_i})}{6\sqrt{S^2_{PC_i}}} \quad (12)$$

and

$$\hat{C}_{pk;PC_i} = \frac{\min(USL_{PC_i} - \bar{X}_{PC_i}, \bar{X}_{PC_i} - LSL_{PC_i})}{3\sqrt{S^2_{PC_i}}} \quad (13)$$

The approximate $100(1 - \alpha)\%$ **confidence intervals** for MC_p and MC_{pk} are

$$\left[\left(\prod_{i=1}^v \hat{C}_{p;PC_i} \sqrt{\frac{\chi^2_{1-\alpha/2, n-1}}{n-1}} \right)^{1/v}, \left(\prod_{i=1}^v \hat{C}_{p;PC_i} \sqrt{\frac{\chi^2_{\alpha/2, n-1}}{n-1}} \right)^{1/v} \right] \quad (14)$$

and

$$\left[\left(\prod_{i=1}^v \hat{C}_{pk;PC_i} \left[1 - Z_{\alpha/2} \sqrt{\frac{1}{9n\hat{C}_{pk;PC_i}} + \frac{1}{2(n-1)}} \right] \right)^{1/v}, \left(\prod_{i=1}^v \hat{C}_{pk;PC_i} \left[1 + Z_{\alpha/2} \sqrt{\frac{1}{9n\hat{C}_{pk;PC_i}} + \frac{1}{2(n-1)}} \right] \right)^{1/v} \right] \quad (15)$$

only few principal components can explain most of the total variability (about 70 ~ 90%). Jackson [11] proposed a χ^2 test for identifying the significant components. Referring to this method, we can choose the suitable number of PC_i s easily. Now, the related formulas of multivariate PCIs for the multivariate normal data and non-multivariate normal data can be displayed for the following: (a) two multivariate PCIs for multivariate normal data are defined as

$$\widehat{MC}_p = \left(\prod_{i=1}^v \hat{C}_{p;PC_i} \right)^{1/v} \quad (10)$$

and

$$\widehat{MC}_{pk} = \left(\prod_{i=1}^v \hat{C}_{pk;PC_i} \right)^{1/v} \quad (11)$$

where $Z \sim N(0, 1)$, respectively. (b) The multivariate PCI for non-multivariate normal data is based on a PCI by Luceño [12], which is designed to consider both the process location and spread, as well as to provide for confidence bounds that are insensitive to departures from normality. The multivariate PCI is defined as

$$\widehat{MC}_{pc} = \left(\prod_{i=1}^v \hat{C}_{pc;PC_i} \right)^{1/v} \quad (16)$$

where $\hat{C}_{pc;PC_i} = (USL_{PC_i} - LSL_{PC_i}) / (6\sqrt{\pi/2\bar{c}})$. The approximate $100(1 - \alpha)\%$ confidence interval for MC_{pc} is

$$\left[\left(\prod_{i=1}^v \frac{\hat{C}_{pc;PC_i}}{1 + t_{\alpha, n-1} S_{C_i} / (\bar{c}_i \sqrt{n})} \right)^{1/v}, \right]$$

4 Multivariate Process Capability Indices, Comparison of

$$\left(\prod_{i=1}^v \frac{\hat{C}_{pc;PC_i}}{1 - t_{\alpha,n-1} S_{C_i} / (\bar{c}_i \sqrt{n})} \right)^{1/v} \quad (17)$$

where

$$\bar{c}_i = \frac{1}{n} \sum_{j=1}^n \left| PC_{ij} - \frac{(USL_{PC_i} + LSL_{PC_i})}{2} \right| \quad (18)$$

$$S_{C_i} = \frac{1}{n-1} \left(\sum_{j=1}^n \left| PC_{ij} - \frac{(USL_{PC_i} + LSL_{PC_i})}{2} \right|^2 - n \bar{c}_i^2 \right) \quad (19)$$

and t is the Student's t distribution.

Walter *et al.* [8] proposed the multivariate PCIs based on the eigenvalue problem and quadratic form. First, we can generalize the formula of the univariate PCI $C_p = (USL - LSL)/6s$ using the expression $\text{diag}((USL - LSL)/6) S_{X'X}^{-1} \text{diag}((USL - LSL)/6)$. Then, we can define the uncorrected multivariate PCI by

$$MC_p^{(B1)} = \min_{j=1,\dots,v} \left\{ \text{eigenvalues} \times \left[\text{diag} \left(\frac{USL - LSL}{6} \right) S_{X'X}^{-1} \times \text{diag} \left(\frac{USL - LSL}{6} \right) \right]^{1/2} \right\} \quad (20)$$

This formula compares the diagonal of a standardized tolerance region with the length of the longest space diagonal of the standardized variance matrix. Furthermore, a quadratic correction of the equation (20) is given by

$$MC_{pk}^{(B1)} = \frac{MC_p^{(B1)}}{D_B} \quad (21)$$

where $D_B = \sqrt{1 + (\bar{X} - T)^T (\sum \text{th})^{-1} (\bar{X} - T)}$ and $\sum \text{th} = \text{diag}((USL_j - LSL_j)/6) \times ((USL_j - LSL_j)/6)$.

The quadratic correction above does not consider the dependence between the product variables. That is, the tolerance region can be a circle or a sphere. The interpretation of this index is similar to the univariate index C_{pk} . They also proposed the multivariate PCIs

based on the eigenvalue problem (see the details in Walter *et al.* [8]).

Pearn *et al.* [9] provided the probability of producing a good product satisfying all its specifications. The process yield is defined as

$$p = \int_{[LSL, USL]} N_v(X|\mu, \Sigma) dX \quad (22)$$

Furthermore, they proposed a generalized process yield index $TS_{pk,PC}$, based on the index S_{pk} [13] and the PCA method which is similar to Wang and Du [7]. They also provided a lower confidence bound for the true process yield.

Discussion

These methods proposed by Chan *et al.* [2], Pearn *et al.* [3], Taam *et al.* [4], Shahriari *et al.* [6], Walter *et al.* [8], and Pearn *et al.* [9], respectively, must use the assumption of multivariate normality, while those proposed by Chen [5] and Wang and Du [7] do not require this assumption of multivariate normality. The tolerance regions of these methods using the assumption of multivariate normality are ellipsoidal except for Shahriari *et al.* [6]. On one hand, Chen [5] and Wang and Du [7] provided more flexible methods in assessing the capability for multivariate data. Wang *et al.* [14] presented a comparison of three methods by Tamm *et al.* [4], Shahriari *et al.* [6], and Chen [5]. In general, these proposed approaches all have their own reasonable points of view, but we still cannot obtain a consistent method of assessing the multivariate PCIs. Apparently, we should make an effort to study the relevant issues in the multivariate PCIs. Here, we give a simple summary for the multivariate PCIs (see Table 1).

Example

Let us consider an example of a device on an integrated circuit. The width (x_1) and thickness (x_2) are two important characteristics for a device on an integrated circuit to ensure good conductivity. The specifications for these two variables are (5, 6) (μm) and (0.6, 1.2) (μm), respectively. The targets of the specifications are (5.5, 0.9). From a sample of 30 observations, we have the sample mean vector and the sample covariance matrix which are $\bar{X} = \begin{bmatrix} 5.512 \\ 0.901 \end{bmatrix}$

Table 1 The summary of multivariate process capability indices

Authors	PCI	Distribution application	Tolerance form	Category
Chan <i>et al.</i> [2]	C_{pm}	Multivariate Normal	Elliptical	Ratio value
Pearn <i>et al.</i> [3]	${}_vC_{pm}, {}_vC_p$	Multivariate Normal	Elliptical	Ratio value
Taam <i>et al.</i> [4]	MC_{pm}, MC_p	Multivariate Normal	Elliptical	Ratio value
Chen [5]	MC_p	No specific	No specific	Process yield
Shahriari <i>et al.</i> [6]	$[C_{pM}, PV, LI]$	Multivariate Normal	Elliptical	Ratio value
Wang and Du [7]	MC_{pm}, MC_{pk}, MC_p	No specific	No specific	Ratio value
Walter <i>et al.</i> [8]	$MC_p^{(B1)}, MC_{pk}^{(B1)}$	Multivariate Normal	Elliptical	Ratio value
Pearn <i>et al.</i> [9]	$2\Phi(3TS_{pk;PC}) - 1$	Multivariate Normal	Elliptical	Process yield

Table 2 The results for example

Authors	Estimated values
Chan <i>et al.</i> [2]	$C_{pm} = 1.0157$
Taam <i>et al.</i> [4]	$MC_p = 2.9099, MC_{pm} = 2.9099/1.0166 = 2.8765$
Shahriari <i>et al.</i> [6]	$(C_{pM}, PV, LI) = (1.1899, 0.7250, 1)$
Wang and Du [7]	$MC_p = 1.1133, MC_{pk} = 1.091^{(a)}, MC_{pm} = 1.1110$
Walter <i>et al.</i> [8]	$MC_p^{(B1)} = 0.8733, MC_{pk}^{(B1)} = 0.8733/1.0063 = 0.8678$
Pearn <i>et al.</i> [9]	The estimated process yield is 0.9984

^(a)95% lower confidence bound: $MC_p = 0.8699, MC_{pk} = 0.8388$

and $S = \begin{bmatrix} 0.0257 & 0.0097 \\ 0.0097 & 0.0044 \end{bmatrix}$, respectively. Now we calculate the capability values for this case depending on several methods we introduce and we give a summary in Table 2. From Table 2, all capability indices are larger than 1 except the indices of Walter *et al.* [8]. This means that the variation of the process is small, the mean of the process is close to the center of specification and the mean of the process is around the target. Based on this result, we can conclude that this process is capable. By observing the spread of the 30 observations in the above figures, it is clear that this conclusion is reasonable. Here we get the estimated process yield 0.9984 by using the maple program. This means that the process produces high percentage of conforming products and implies the process is capable. Hence, the process should have high process capability values. According to the above process capability values calculated, we can infer that the estimated process yield is rational.

Conclusion

Often, a manufactured product is described in multiple characteristics. That is, manufactured items require values of several different characteristics

for adequate description of their quality. Each of those characteristics must satisfy certain specifications. The assessed quality of a product depends on the combined effects of these characteristics, rather than on their individual values.

Currently, the research of multivariate PCIs is still very sparse in comparison to the research of the univariate PCIs. Also, there is no consistency regarding a methodology for evaluating the process capability with multiple quality characteristics. In addition, it is very difficult to obtain the relevant statistical properties to make a more detailed inference for the multivariate PCIs. Thus, we realize that there are still critical difficulties in trying to assess the value of multivariate systems in terms of a single value. Clearly, further investigations in this field are needed.

References

- [1] Kotz, S. & Johnson, N.L. (2002). Process capability indices – a review, 1992–2000, *Journal of Quality Technology* **34**, 2–19.
- [2] Chan, L.K., Cheng, S.W. & Spiring, F.A. (1991). A multivariate measure of process capability, *Journal of Modeling and Simulation* **11**, 1–6.

6 Multivariate Process Capability Indices, Comparison of

- [3] Pearn, W.L., Kotz, S. & Johnson, N.L. (1992). Distributional and inferential properties of process capability indices, *Journal of Quality Technology* **24**, 216–231.
- [4] Taam, W., Subbaiah, P. & Liddy, J.W. (1993). A note on multivariate capability indices, *Journal of Applied Statistics* **20**, 339–351.
- [5] Chen, H. (1994). A multivariate process capability index over a rectangular solid tolerance zone, *Statistica Sinica* **4**, 749–758.
- [6] Shahriari, H., Hubele, N.F. & Lawrence, F.P. (1995). A multivariate process capability vector, *Proceedings of the 4th Industrial Engineering Research Conference*, Institute of Industrial Engineers, Nashville, pp. 304–309.
- [7] Wang, F.K. & Du, T.C.T. (2000). Using principal component analysis in process performance for multivariate data, *Omega* **28**, 185–194.
- [8] Walter, J., Anghel, C. & Braun, L. (2004). Multivariate process capability indices: new approaches and comparisons with known indices, working paper.
- [9] Pearn, W.L., Wang, F.K. & Yen, C.H. (2006). Measuring production yield for processes with multiple quality characteristics, *International Journal of Production Research* **44**, 4649–4661.
- [10] Mahalanobis, P.C. (1936). On the generalized distance in statistics, *Proceedings of the National Institute of Science of India* **12**, 49–55.
- [11] Jackson, J.E. (1980). Principal component and factor analysis: part I – principal components, *Journal of Quality Technology* **12**, 201–213.
- [12] Luceño, A. (1996). A process capability index with reliable confidence intervals, *Communications in Statistics – Simulation and Computation* **25**, 235–245.
- [13] Boyles, R.A. (1994). Process capability with asymmetric tolerances, *Communications in Statistics – Simulation* **23**, 615–643.
- [14] Wang, F.K., Hubele, N.F., Lawrence, F.P., Miskulin, J.D. & Shahriari, H. (2000). Comparison of three multivariate process capability indices, *Journal of Quality Technology* **32**, 263–275.

FU-KWUN WANG

Statistical Process Control, Multivariate

Introduction

Multivariate statistical process control (SPC) combines process control and process capability analysis. Although multivariate SPC was developed in the 1940s [1], its intensive use in industry started in the 1980s [2–5]. The emergence of computer-automated production systems and automated data collection systems equipped with data analysis technology had a significant impact on the use of multivariate SPC in industry. Such systems collect multivariate measurements rapidly, in a cost-effective way, and provide effective data analysis capabilities. Some examples include integration in production systems of process analytical techniques (PAT), three-dimensional coordinate measurement machines and vision systems based on cameras, and laser technology [6].

The oldest and fundamental multivariate SPC chart is the Hotelling T^2 chart that monitors the process mean and provides an out-of-control signal when the process mean shifts significantly from its normal behavior (see **Hotelling's T^2 ; Multivariate Control Charts Overview**). The Hotelling T^2 chart, plots the T^2 statistic also known as the *Mahalanobis distance*. The chart has an upper control limit (UCL) and points exceeding that limit are regarded as an out-of-control signal. The charted T^2 statistic is a function that reduces multivariate observations into a single value. This function includes the **covariance** matrix among its variables. Although the chart is powerful in detecting shifts in mean values, the lack of practical meaning of the T^2 statistic hinder its interpretation by engineers.

Recently, various methods have been developed to facilitate the proper interpretation of the T^2 signals [5, 7]. Some of these methods attempt to detect variables responsible for out-of-control signals on the chart. These methods are effective in uncovering the problem source, whenever the problem is caused by a subset of variables contributing significantly to the signal. The assumption that the problem sources affect only certain variables is not always valid and explains only a certain class of out-of-control signals. For the determination of patterns created by specific problem sources, one may use methods based on

in-depth process expertise. Once such patterns are identified, special charts can be designed for detecting their occurrences in the process. These charts track deviations in the variable means in specific directions or in the form of the T^2 statistic [8, 9].

The setup of a multivariate SPC chart consists of a process capability study [7, 10]. Following a process capability study, control limits and control procedures are determined. The process capability study period is sometimes referred to as *phase I*. The ongoing control using control limits determined in phase I is then called *phase II* [5].

The distinction between these two phases is important. Phase I analysis estimates parameters that will be used subsequently for ongoing monitoring of the process or phase II [11]. Phase I should use a relatively large sample to estimate the parameters and control limits used in phase II. For multivariate data analysis techniques and **multivariate capability indices**, see **Multivariate Analysis, Process Yield, and Multi-Vari Charts**.

In setting multivariate SPC charts, one meets several alternative scenarios derived from the characteristics of the reference sample and the appropriate control procedure. These include

1. internally derived targets
2. using an external reference sample
3. externally assigned targets
4. measurements units considered as batches.

In this section we cover techniques for multivariate SPC under these four scenarios. We then proceed to cover the step down and variable decomposition procedures and applications of multivariate statistical tolerance regions to SPC.

Internally Derived Targets

Internally derived targets are a typical scenario for process capability studies. This phase is further divided into two stages [12]. The first stage of phase I is *retrospective*, and historical data are used to estimate the relevant parameters. The second stage is *prospective* and the parameters estimated in the first stage are applied to a group of values and tested for accuracy. The parameters that must be estimated include the process mean vector, i.e., the *internally derived targets*, the process covariance matrix, and the control limit for the control chart.

2 Statistical Process Control, Multivariate

Consider a process capability study with a the base sample of size n of p -dimensional observations, $X_1, X_2, \dots, X_n$. When the data are grouped and k subgroups of observations of size m are being monitored, $n = km$, the covariance matrix estimator, S_p can be calculated as the pooled covariances of the subgroups [13]. In that case, for the j th subgroup, the Hotelling T^2 statistic is given by:

$$T^2 = m(\bar{X}_j - \bar{\bar{X}})' S_p^{-1} (\bar{X}_j - \bar{\bar{X}}) \quad (1)$$

where $\bar{X}_j$ is the mean of the j th subgroup, $\bar{\bar{X}}$ is the overall mean, and S_p^{-1} is the inverse of the pooled estimated covariance matrix. The UCL for this case is

$$UCL = \frac{p(k-1)(m-1)}{k(m-1) - p + 1} F_{p, k(m-1) - p + 1}^\alpha \quad (2)$$

When the data are ungrouped, and individual observations are analyzed, the estimation of the proper covariance matrix and control limits requires further consideration. Typically in this case, the covariance matrix is estimated from the pooled individual observations as $S = (1/n - 1) \sum_{i=1}^n (X_i - \bar{X})(X_i - \bar{X})'$, where $\bar{X}$ is the mean of the n observations. The corresponding T^2 statistic for the i th observation, $i = 1, \dots, n$ is given by $T^2 = (X_i - \bar{X})' S^{-1} (X_i - \bar{X})$.

In this case $(X_i - \bar{X})$ and S are not independently distributed and the appropriate upper control for T^2 is based on the beta distribution with

$$UCL = \frac{(n-1)^2}{n} B_{p/2, (n-p-1)/2}^{\alpha/2} \quad (3)$$

and

$$LCL = \frac{(n-1)^2}{n} B_{p/2, (n-p-1)/2}^{1-\alpha/2} \quad (4)$$

where $B_{p/2, (n-p-1)/2}^\alpha$ is the $100(1 - \alpha)$ th percentile of the beta distribution with $p/2$ and $(n - p - 1)$ as parameters. While theoretically the lower control limit (LCL) can be calculated as above, in most circumstances LCL is set to zero [14].

While control charts based on the T^2 and control values are commonly used, Sullivan and Woodall [12] found that the pooled individual observation estimator was relatively inaccurate in most of the tests. Using simulation, they evaluated several techniques to estimate the covariance matrix for monitoring individual data. Two techniques that incorporated a subgroup breakdown of the individual observations and an overlapping of the subgroups were found to perform particularly well. The overlapping technique,

in which observations were included in more than one subgroup, created a dependence of subgroups on each other and strengthened their statistical power to estimate the process covariance matrix and mean vector [12]. The use of such techniques is particularly relevant when the assumption of independence of the observations in the base sample is not necessarily valid. The resulting process covariance matrix and mean vector can then be used for setting the control limits. Once a UCL is set, the data of the reference sample can be evaluated for in-control characteristics. Points outside the UCL can be due to outliers or process instability. Gnanadesikan [15] and Fuchs and Kenett [7] present several techniques for identifying multivariate outliers.

Using an External Reference Sample

Consider again a process yielding assumed independent observations $X_1, X_2, \dots$ of a p -dimensional random variable $\mathbf{X}$, such as quality characteristics of a manufactured item or process measurements. Initially, when the process is “in control”, the observations follow a distribution F , with density f . We now assume that we have a “reference” (prerun) sample $X_1, \dots, X_n$ of F from an in-control period. To control the quality of the produced items or of the process, multivariate data is monitored for potential change in the distribution of $\mathbf{X}$, by sequentially collecting and analyzing the observations X_i . At some time $t = n + k$, k time units after n , the process may run out of control and the distribution of the X_i 's changes to G . Our aim is to detect, in phase II, the change in the distribution of subsequent observations X_{n+k} , $k \geq 1$, as quickly as possible, subject to a bound $\alpha \in (0, 1)$ on the probability of raising a false alarm at each time point $t = n + k$ (that is, the probability of erroneously deciding that the distribution of X_{n+k} is not F). The reference sample $X_1, \dots, X_n$ does not incorporate the observations X_{n+k} taken after the prerun stage, even if no alarm is raised for it, so the rule will be conditional only on the reference sample. As described below the structure of the reference sample can vary and we will later indicate how to handle other alternatives.

As an example, consider the measurements of the aluminum pins case study presented in [7]. The data consists of six recorded dimensions, per part. Data

Table 1 Data on six dimensions recorded on 30 parts used in process capability study (Phase I)

ID	Dimension 1	Dimension 2	Dimension 3	Dimension 4	Dimension 5	Dimension 6
1	9.99	9.97	9.96	14.97	49.89	60.02
2	9.96	9.96	9.95	14.94	49.84	60.02
3	9.97	9.96	9.95	14.95	49.85	60
4	10	9.99	9.99	14.99	49.89	60.06
5	10	9.99	9.99	14.99	49.91	60.09
6	9.99	9.99	9.98	14.99	49.91	60.08
7	10	9.99	9.99	14.98	49.91	60.08
8	10	9.99	9.99	14.99	49.89	60.09
9	9.96	9.95	9.95	14.95	50	60.15
10	9.99	9.98	9.98	14.99	49.86	60.06
11	10	9.99	9.98	14.99	49.94	60.08
12	10	9.99	9.99	14.99	49.92	60.05
13	9.97	9.96	9.96	14.96	49.9	60.02
14	9.97	9.96	9.96	14.96	49.91	60.02
15	9.97	9.97	9.96	14.97	49.9	60.01
16	9.97	9.97	9.96	14.97	49.89	60.04
17	9.98	9.97	9.96	14.96	50.01	60.13
18	9.98	9.97	9.97	14.96	49.93	60.06
19	9.98	9.98	9.97	14.98	49.93	60.02
20	9.98	9.97	9.97	14.97	49.94	60.06
21	9.98	9.97	9.97	14.97	49.93	60.06
22	9.98	9.97	9.97	14.97	49.91	60.02
23	9.98	9.97	9.96	14.98	49.92	60.06
24	10	9.99	9.98	14.98	49.88	60
25	9.99	9.99	9.99	14.98	49.91	60.04
26	10	9.99	9.99	14.99	49.85	60.01
27	10	10	9.99	14.99	49.91	60.05
28	10	9.99	9.99	15	49.92	60.04
29	10	9.99	9.99	14.99	49.89	60.01
30	10	10	9.99	14.99	49.88	60

from 30 parts manufactured in an “in-control” base sample is presented in Table 1.

The mean vector of the base sample is $\bar{X}$ and the estimated covariance matrix is S . For the data above:

$$\bar{X}' = (9.99, 9.98, 9.97, 14.98, 49.91, 60.05) \quad (5)$$

The inverse covariance matrix is

$$S^{-1} = \begin{vmatrix} 53192.3 & -22792.3 & -17079.7 & -9343.4 & 145 & -545.3 \\ & 66325.1 & -28341.1 & -10877.9 & 182 & 1522.8 \\ & & 50553.9 & -6467.9 & 853 & -1465.1 \\ & & & 25745.6 & -527.5 & 148.6 \\ & & & & 1622.3 & -1156 \\ & & & & & 1577.6 \end{vmatrix} \quad (6)$$

For each future observation $\mathbf{Y}$ of dimension $p = 6$, we compute the T^2 statistic as

$$T^2 = (\mathbf{Y} - \bar{\mathbf{X}})' S^{-1} (\mathbf{Y} - \bar{\mathbf{X}}) \quad (7)$$

Under the null hypothesis, the T^2 statistic has a $((n-1)p/(n-p))F_{p,n-p}$ distribution with $n = 30$ and $p = 6$.

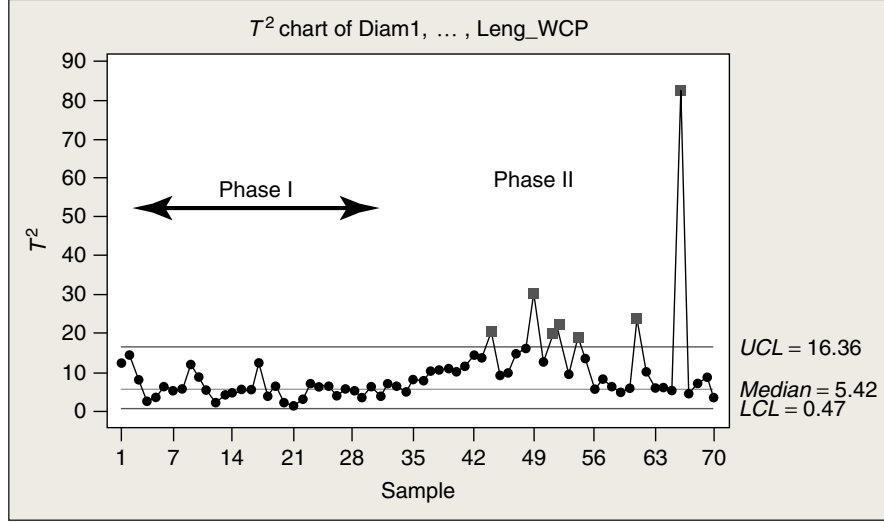


Figure 1 T^2 value chart with setup values determined by the first 30 observations. Observation 44 is the first out-of-control signal

Accordingly, the 0.95 UCL for T^2 is $UCL = ((n-1)p/(n-p))F_{p,n-p}^{.95} = 16.35$.

Figure 1 presents the values of T^2 using the mean and covariance of the first 30 samples. The base sample is indeed identified as “in control” with no T^2 value exceeding the UCL. As the production process continued producing aluminum pins some products signaled an out-of-control condition. For example, the T^2 value for the 44th observation (10, 10, 10, 15, 49.93, 59.98) is 20.167, and exceeds the UCL.

When the data in phase II are grouped, and the reference sample from historical data includes k subgroups of observations of size m , $n = km$, with the covariance matrix estimator S_p calculated as the pooled covariances of the subgroups, the T^2 for a new subgroup of size m with mean $\bar{Y}$ is given by $T^2 = m(\bar{Y} - \bar{\bar{X}})'S_p^{-1}(\bar{Y} - \bar{\bar{X}})$, and the UCL is given by $UCL = (p(k+1)(m-1)/k(m-1) - p + 1)F_{p,k(m-1)-p+1}^\alpha$. Furthermore, if in phase I, l subgroups were outside the control limits and assignable causes were determined, those subgroups are omitted from the computation of $\bar{\bar{X}}$ and S_p^{-1} , and the control limits for this case are $UCL = (p(k-l+1)(m-1)/(k-l)(m-1) - p + 1)F_{p,(k-l)(m-1)-p+1}^\alpha$. The T^2 control charts constructed in phase I and used both in phase I and in phase II are the multivariate equivalent of the Shewhart control chart. Those charts as well as some

more advanced ones simplify the calculations down to single-number criteria and produce a desired Type I error or in-control run length. More advanced charts like the multivariate cumulative sum (MCUSUM) and the multivariate exponentially weighted moving average (MEWMA) produce charts, which, similar to their univariate equivalent, are more sensitive to small deviations from control.

While we focused on the reference sample provided by phase I of the multivariate process control, other possibilities can occur as well. In principle, the reference in-control sample can also originate from historical data which are not part of the current process control. In this case, the statistical analysis will be the same but this situation has to be treated with precaution since both the control limits and the possible correlations between observations may be shifted.

Externally Assigned Targets

If all parameters of the underlying multivariate distribution are known and externally assigned, the T^2 value for a single multivariate observation of dimension p is computed as:

$$T^2 = (\mathbf{Y} - \boldsymbol{\mu})' \boldsymbol{\Sigma}^{-1} (\mathbf{Y} - \boldsymbol{\mu}) \quad (8)$$

where $\boldsymbol{\mu}$ and $\boldsymbol{\Sigma}$ are the expected value and covariance matrix, respectively.

The probability distribution of the T^2 statistic is a χ^2 distribution with p **degrees of freedom**. Accordingly, the 0.95 UCL for T^2 is $UCL = \chi^2(0.05, p)$.

When the data are grouped in subgroups of size m , and both μ and Σ are known, the T^2 value of the mean vector $\bar{Y}$ is $T^2 = m(\bar{Y} - \mu)' \Sigma^{-1} (\bar{Y} - \mu)$ with the same UCL as above.

If only the expected value of the underlying multivariate distribution, μ , is known and externally assigned, the covariance matrix has to be estimated from the tested sample. The T^2 value for a single multivariate observation of dimension p is computed as

$T^2 = (\mathbf{Y} - \mu)' S^{-1} (\mathbf{Y} - \mu)$, where μ is the expected value and S is the estimate of the covariance matrix Σ , estimated either as the pooled contribution of the individual observations, i.e. $S = (1/n - 1) \sum_{i=1}^n (X_i - \bar{X})(X_i - \bar{X})'$, or by one of the methods suggested by Sullivan and Woodall [12], which account for possible lack of independence between observations.

In this case, the 0.95 UCL for T^2 is $UCL = (p(n - 1)/n(n - p)) F_{p, n-p}^{0.95}$.

When the tested observations are grouped, the mean vector of a subgroups with m observations (a rational sample) will have the same expected value as the individual observations, μ , and a covariance matrix Σ/m . The covariance matrix Σ can be estimated by S or as S_p obtained by pooling the covariances of the k subgroups.

When Σ is estimated by S , the T^2 value of the mean vector $\bar{Y}$ of m tested observation is $T^2 = m(\bar{Y} - \mu)' S^{-1} (\bar{Y} - \mu)$ and the 0.95 UCL is $UCL = (p(m - 1)/m - p) F_{p, m-p}^{0.95}$. When Σ is estimated by S_p , the 0.95 UCL of $T^2 = m(\bar{Y} - \mu)' S_p^{-1} (\bar{Y} - \mu)$ is $UCL = (pk(m - 1)/k(m - 1) - p + 1) F_{p, k(m-1)-p+1}^{0.95}$.

Measurements Units Considered as Batches

In many situations, the production is organized in batches. In those cases, the quality-control process can be performed in one of the two modes, either at the completion of the batch or sequentially, in a curtailed inspection, aiming at reaching a decision as soon as possible. Whenever the quality-control method determines that the decision is to be reached at the completion of the process, the two possible outcomes are (a) determine the production process

to be in statistical control and accept the batch or (b) stop the production flow because of a signal that the process is out of control. On the other hand, in a curtailed inspection, based on a statistical stopping rule, the results from the first few items tested may suffice to stop the process prior to the batch completion.

Let us consider a batch of size n , with the tested items $\mathbf{Y}_1, \dots, \mathbf{Y}_n$. The curtailed inspection tests the items sequentially. Assume that the targets are specified, either externally assigned, or from a reference sample or batch. With respect to those targets, let $V_i = 1$ if the T^2 of the ordered i th observation exceeds the critical value κ and $V_i = 0$, otherwise. For the i th observation, the process is considered to be in control if for a prespecified P , say $P = 0.95$, $\Pr(V_i = 0) \geq P$. Obviously, the inspection will be curtailed only at an observation i for which $V_i = 1$ (not necessarily the first).

Let $N(g) = \sum_{i=1}^g V_i$ be the number of rejection up to the g th tested item. For each number of individual rejections U (out of n), $R(U)$ denotes the minimal number of observations allowed up to the U th rejection, without rejecting the overall null hypothesis. Thus, for each U , $R(U)$ is the minimal integer value such that under the null hypothesis, $\Pr\left(\sum_{i=1}^{R(U)} V_i \leq U\right) \geq \alpha$. For fixed U , the random variable $\sum_{i=1}^U V_i$ has a negative binomial distribution, and we can compute $R(N(g))$ from the inverse of the negative binomial distribution. For example when $n = 13$, $P = 0.95$, and $\alpha = 0.01$, the null hypothesis is rejected if the second rejection occurred at or before the third observation, or if the third rejection occurred at or before the ninth observation, and so on (see [7]).

Variable Decomposition and Monitoring Indices

The batched data are naturally grouped, but even if quality control is performed on individual items, grouping the data into rational consequent subgroups may yield relevant information within subgroups variability in addition to the deviations from targets.

In the j th subgroup (or batch) of size n_j the individual T_{ij}^2 values, $i = 1, \dots, n_j$ are given by $T_{ij}^2 = (\mathbf{Y}_{ij} - \theta)' G^{-1} (\mathbf{Y}_{ij} - \theta)$. When the targets are externally assigned then $\theta = \mu$. If the covariance matrix is also externally assigned then $G = \Sigma$, otherwise G is

the covariance matrix estimated either from the tested or from the reference sample. In the case of targets derived from an external reference sample $\theta = \mathbf{m}$ and $G = S$, where $\mathbf{m}$ and S are the mean and the covariance matrix of a reference sample of size n . Within the j th subgroup, let us denote the mean of the subgroup observations by $\bar{Y}_j$ and the mean of the target values in the j th subgroup by θ_j .

The sum of the individual T_{ij}^2 values, $T_{0j}^2 = \sum_{i=1}^n T_{ij}^2$ can be decomposed into two measurements of variability, one representing the deviation of the subgroup mean from the multivariate target denoted by T_{Mj}^2 , and the other measuring the internal variability within the subgroup, denoted by T_{Dj}^2 . The deviation of the subgroup mean from the multivariate target is estimated by $T_{Mj}^2 = (\bar{Y}_j - \theta_j)' G^{-1} (\bar{Y}_j - \theta_j)$, while the internal variability within the subgroup is estimated by $T_{Dj}^2 = (\mathbf{Y}_{ij} - \bar{Y}_j)' G^{-1} (\mathbf{Y}_{ij} - \bar{Y}_j)$, with $T_{0j}^2 = T_{Mj}^2 + T_{Dj}^2$.

Since asymptotically, T_{Mj}^2 and T_{Dj}^2 have a χ^2 distribution with p and $(n-1)p$ degrees of freedom, respectively, one can further compute two indices, I_1 and I_2 , to determine whether the overall variability is mainly due to the distances between the means of the tested subgroup from targets or to the within subgroup variability. The indices are relative ratios of the normalized versions of the two components of T_{0j}^2 , i.e., $I_1 = I_1^*/(I_1^* + I_2^*)$ and $I_2 = I_2^*/(I_1^* + I_2^*)$, where $I_1^* = T_{Mj}^2/p$ and $I_2^* = T_{Dj}^2/[(n-1)p]$. We can express the indices in terms of the original T^2 statistics as, $I_1 = (n-1)T_{Mj}^2/[(n-1)T_{Mj}^2 + T_{Dj}^2]$ and $I_2 = T_{Dj}^2/[(n-1)T_{Mj}^2 + T_{Dj}^2]$. Tracking these indices provides powerful monitoring capabilities. For more details and examples see [7].

Multivariate Tolerance Regions

An alternative approach to design a multivariate SPC scheme is to use tolerance regions as a basis for decision making. Tolerance regions provide inferential capabilities on **percentiles** of distributions. A specific value of a variable is a tolerance limit for a proportion p , at a specific confidence level, for example 95%. If we set control limits of a univariate control chart using tolerance limits for the 0.013 and 0.987 proportion, at the 95% confidence level, we obtain UCLs and LCLs that are different from the 3σ Shewhart control limits. The advantage of fusing tolerance region is that they are easily extendable

to several dimensions with the concept of proportion and confidence level remaining unchanged.

Let $\mathbf{X}_1, \dots, \mathbf{X}_n$ be a sample of n independent and identically distributed observations from a multivariate **normal distribution** of dimension p with mean μ and covariance matrix Σ . In the quality-control area we use a sample $(\mathbf{X}_1, \dots, \mathbf{X}_n)$ as a reference sample for setting the standards to which each new observation is compared.

When using a tolerance region as a standard, it is necessary to construct from $(\mathbf{X}_1, \dots, \mathbf{X}_n)$ a region about which we can assert with probability δ that the proportion of the individuals in the population contained in that region is greater than or equal to Γ . Each new observation is tested to determine whether it belongs to that region.

Let X_{ih} ($i = 1, 2, \dots, p$) be the value of the i th attribute for the h th observation ($h = 1, 2, \dots, n$) and S the empirical covariance matrix with elements s_{ij} :

$$s_{ij} = \sum_{h=1}^n \frac{(X_{ih} - \bar{X}_i)(X_{jh} - \bar{X}_j)}{(n-1)} \quad (9)$$

The means $\bar{X}_i$ and $\bar{X}_j$ are the i th and j th elements, respectively, of the sample mean vector $\bar{\mathbf{X}} = (\bar{X}_1, \dots, \bar{X}_p)$. It has been shown [16] that if we ignore terms of order two in $(1/n)$ the tolerance region can be defined by the $\mathbf{X}$ vectors satisfying the inequality

$$(\mathbf{X} - \bar{\mathbf{X}})' \mathbf{S}^{-1} (\mathbf{X} - \bar{\mathbf{X}}) \leq k \quad (10)$$

where k is a function of the prespecified (Γ, δ) and of the dimensionality of the problem (p). (The left-hand side of equation (10) is the Mahalanobis distance.)

The critical value k is defined as follows: Let $g(\phi, f, m)$ be the ϕ th percentile of the noncentral χ^2 distribution with m degrees of freedom and noncentrality parameter f . As a special case, $g(\phi, 0, m)$ is the ϕ th percentile of the central χ^2 distribution.

$$Pr(\chi_m^2 \leq g(\phi, 0, m)) = \phi \quad (11)$$

The critical value k (see John [16]) is given by

$$k = \frac{g(\Gamma, p/2n, p)(n-1)p}{g(1-\delta, 0, (n-1)p)} \quad (12)$$

with $p/2n$ being the noncentrality parameter of the χ^2 distribution.

Other approximations have been provided by Krishnamoorthy and Mathew [17]. The tolerance limits approach to process control, as an alternative to the classical control limits extending the basic Shewhart chart, is slowly gaining popularity. Parametric and nonparametric multivariate tolerance regions have been investigated by Di Bucchianico *et al.* [18] and Baillo and Cuevas [19].

Several aspects of multivariate SPC have been covered in this entry. In particular we carefully distinguished between alternative scenarios for establishing multivariate SPC including internal and external reference samples and data in batches or considered individually.

Other Multivariate Charts

In the multivariate equivalent of the Shewhart chart, the determination of the state of statistical control is based only on the current observation vector. As their univariate counterparts, the multivariate charts that plot statistics based on weighted sequences of previous observations are in general more sensitive to small deviations from control.

Alwan [20] and Crosier [21] suggested MCUSUMs charts of T^2 and T , respectively, as multivariate extensions of univariate cumulative sum (CUSUM) (*see Cumulative Sum (CUSUM) Chart*). Woodall and Ncube [22] considered a procedure based on individual CUSUM of the variables in X with a signal triggered if any of the p CUSUMs exceeds its control limit. The approach has the advantage of better interpretability than procedures based on T . However, when the variables are highly correlated, procedures that capitalize on the correlation between variables will provide improved performance.

The MEWMA chart [23], which is an extension of the univariate version of exponentially weighted moving average (EWMA), plots the statistics $T_i^2 = Z_i' \sum_{j=1}^i Z_j$, where $Z_i = \lambda X_i + (1 - \lambda)Z_{i-1}$, $\sum_{j=1}^i Z_j$ is the covariance matrix, and λ is the EWMA weighting matrix. By using various weights for each variable, the investigator can perform a sensitivity analysis. Prabhu and Runger [24] show that the choice of λ does not dramatically affect in-control run length. When all the variables are assigned equal weights, the chart is directionally invariant as are the univariate versions of the CUSUM and EWMA. On the other

hand, by using direction analysis on MCUSUM and MEWMA, one can determine shifts of particular variables along their respective axes, and manipulate specific variables to “aim” them in a particular direction [25]. The results of Runger and Prabhu [26] concur with Lowry *et al.* [23] that in general the MEWMA is more effective in detecting mean shifts than **Hotelling’s T^2** chart.

Hawkins [27] proposed to track a vector Z , derived from scaled **residuals** from the regression of each variable on all others. He shows how T^2 can be determined as a sum of p weighted Z_i where the weights are a scaled distance of X_i from the mean. He then considers an individual-chart approach in which each Z_i is monitored separately and a group-chart approach in which a single composite measure computed from the p charts is used to signal an out-of-control condition. These proposals have been shown to perform better than the proposals mentioned in the previous paragraph in shifts of single and pairs of components of X .

Cheng and Thaga [28] propose a multivariate control chart referred to as the *Max-MCUSUM chart* that is capable of simultaneously detecting shifts in both mean vector and covariance matrix of a multivariate process. The chart is based on standardizing the sample mean vectors and covariance matrices and finding the sampling distribution of the standardized values.

Li and Johnson [29] recently introduced a new MCUSUM scheme whose components are the difference between the univariate CUSUM statistic for detecting an increase and that for detecting a decrease. The procedure also applies when monitoring principal components. They also suggest another procedure where, at each new observation, the results for the two most active components are graphed with the two components being selected to have the largest value of the appropriate bivariate CUSUM statistic.

When the assumption of independence between successive observations does not hold, more complex techniques are required. A popular approach is to model the **autocorrelation** in the data and apply standard procedures for independent observations to the residual data [30–32]. Such modeling can involve multivariate dynamic linear model (DLM) or classic **time series** models, such as autoregressive integrated moving average (ARIMA) models. The classic time series models are static in nature implicitly assuming that unknown model parameters remain constant as

time passes while the DLM techniques allow the parameters to be updated.

Yeh *et al.* [33] provide a comprehensive review of **multivariate control charts** for monitoring a covariance matrix and conclude with an appeal for future investigations to compare existing charts in a systematic, organized, and thorough manner. An additional family of models for multivariate process control includes Bayesian tracking procedures such as extensions of the at most one change (AMOC) models developed in [34].

References

- [1] Hotelling, H. (1947). Multivariate quality control, illustrated by the air testing of sample bombsights, in *Techniques of Statistical Analysis*, C. Eisenhart, M.W. Hastay & W.A. Wallis, eds, McGraw-Hill, New York.
- [2] Jackson, J.E. (1985). Multivariate quality control, *Communications in Statistics: Theory and Methods* **14**, 2657–2688.
- [3] Fuchs, C. & Kenett, R.S. (1987). Multivariate tolerance regions and F-tests, *Journal of Quality Technology* **19**, 122–131.
- [4] Fuchs, C. & Kenett, R.S. (1988). Appraisal of ceramic substrates by multivariate tolerance regions, *The Statistician* **37**, 401–411.
- [5] Mason, R.L. & Young, J.C. (2002). *Multivariate Statistical Process Control with Industrial Applications*, ASA-SIAM Series on Statistics and Applied Probability, ASA-SIAM, Philadelphia.
- [6] Wierda, S. (1994). *Multivariate Statistical Process Control*, Groningen Theses in Economics, Management and Organization, Wolters-Noordhoff, Groningen.
- [7] Fuchs, C. & Kenett, R. (1998). *Multivariate Quality Control: Theory and Application*, *Quality and Reliability Series*, Marcel Dekker, New York, Vol. 54.
- [8] Apley, D.W. & Lee, H.Y. (2003). Identifying spatial variation patterns in multivariate manufacturing processes: a blind separation approach, *Technometrics* **45**, 220–234.
- [9] Runger, G., Barton, R., Del Castillo, E. & Woodall, W. (2007). Optimal monitoring of multivariate data for fault patterns, *Journal of Quality Technology* **39**, 159–172.
- [10] Kenett, R. and Zacks, S. (1998). *Modern Industrial Statistics: Design and Control of Quality and Reliability*, Duxbury Press, San Francisco, 2nd Edition 2003, Chinese Edition 2004.
- [11] Alt, F.B. (1984). Multivariate quality control, *The Encyclopedia of Statistical Sciences*, Wiley, pp. 110–122.
- [12] Sullivan, J.H. & Woodall, W.H. (1996). A comparison of multivariate control charts for individual observations, *Journal of Quality Technology* **28**(4), 398–408.
- [13] Montgomery, D.C. (2001). *Introduction to Statistical Quality Control*, 4th Edition, John Wiley & Sons, New York.
- [14] Tracy, N.D., Young, J.C. & Mason, R.L. (1992). Multivariate Control Charts for Individual Observations, *Journal of Quality Technology* **24**, 88–95.
- [15] Gnanadesikan, R. (1977). *Methods for Statistical Data Analysis of Multivariate Observations*, John Wiley & Sons, New York, **20A** 32–52.
- [16] John, S. (1963). A tolerance region for multivariate normal distributions, *Sankhya, Series A* **25**, 363–368.
- [17] Krishnamoorthy, K. & Mathew, T. (1999). Comparison of approximation methods for computing tolerance factors for a multivariate normal population, *Technometrics* **41**(3), 234–249.
- [18] Di Buccianico, A., Einmah, J. & Mushkudiani, N. (2001). Smallest nonparametric tolerance regions, *Annals of Statistics* **29**, 1320–1343.
- [19] Baillo, A. & Cuevas, A. (2006). Parametric versus nonparametric tolerance regions in detection problems, *Computational Statistics* **21**, 523–536.
- [20] Alwan, L.C. (1986). CUSUM quality control – multivariate approach, *Communications in Statistics: Theory and Methods* **15**, 3531–3543.
- [21] Crosier, R.B. (1988). Multivariate generalizations of cumulative sum quality control schemes, *Technometrics* **30**, 291–303.
- [22] Woodall, W.H. & Ncube, M.M. (1985). Multivariate CUSUM quality-control procedures, *Technometrics* **27**, 285–292.
- [23] Lowry, C.A., Woodall, W.H., Champ, C.W. & Rigdon, S.E. (1992). A multivariate exponentially weighted moving average control chart, *Technometrics* **34**, 46–53.
- [24] Prabhu, S.S. & Runger, G.C. (1997). Designing a multivariate EWMA control chart, *Journal of Quality Technology* **29**(1), 8–15.
- [25] Lowry, C.A. & Montgomery, D.C. (1995). A review of multivariate control charts, *IIE Transactions* **27**, 800–810.
- [26] Runger, G.C. & Prabhu, S.S. (1996). A Markov Chain Model for the Multivariate Exponentially Weighted Moving Averages Control Chart, *Journal of the American Statistical Association* **91**(436), 1701–1706.
- [27] Hawkins, D.M. (1991). Multivariate quality control based on regression-adjusted variables, *Technometrics* **33**, 61–75.
- [28] Cheng, S. & Thaga, K. (2005). Multivariate max-CUSUM chart, *Quality Technology and Quantitative Management* **2**, 221–235.
- [29] Li, R. & Johnson, R.A. (2006). A multivariate CUSUM procedure and a most active bivariate chart, *Quality Technology and Quantitative Management* **3**, 401–414.
- [30] Montgomery, D. & Mastrangelo, C. (1991). Some statistical process control methods for autocorrelated data, *Journal of Quality Technology* **23**, 179–193.
- [31] VanBrackle, L.N. & Reynolds, M. (1997). EWMA and CUSUM control charts in the presence of correlation, *Communications in Statistics* **26**, 979–1008.
- [32] Jiang, W. (2004). Multivariate control charts for monitoring autocorrelated processes, *Journal of Quality Technology* **36**(4), 367–379.

- [33] Yeh, A.B., Lin, D. & McGrath, R. (2006). Multivariate control charts for monitoring covariance matrix: a review, *Quality Technology and Quantitative Management* **3**, 415–436.
- [34] Zacks, S. & Kenett, R.S. (1994). Process tracking of time series with change points, in *Recent Advances in Statistics and Probability*, J. Vilaplana & M. Puri, eds, Brill Academic Publishers, pp. 155–171.
- Mason, R.L., Tracy, N.D. & Young, J.C. (1995). Decomposition of T^2 for multivariate control chart interpretation, *Journal of Quality Technology* **27**, 99–108.
- Murphy, B.J. (1987). Selecting out of control variables with the T^2 multivariate quality control procedures, *The Statistician* **36**, 571–583.
- Nedumaran, G. & Pignatiello, J.J. (1998). Diagnosing signals from T^2 and χ^2 multivariate control charts, *Quality Engineering* **10**, 657–667.
- Wishard, J. (1928). Controversies and contradictions in statistical process control, *Journal of Quality Technology* **32**, 341–350.

Further Reading

- Chang, S.I. & Zhang, K. (2007). Statistical process control for variance shift detections of multivariate autocorrelated processes, *Quality Technology and Quantitative Management* **4**, 413–435.
- Doganaksoy, N., Faltin, F.W. & Tucker, W.T. (1991). Identification of out-of-control quality characteristics in a multivariate manufacturing environment, *Communication in Statistics: Theory and Methods* **20**, 2775–2790.
- Hawkins, D.M. (1981). A CUSUM for a scale parameter, *Journal of Quality Technology* **13**, 228–231.
- Jackson, J.E. (1959). Quality control methods for several related variables, *Technometrics* **1**, 359–377.
- Mason, R.L., Chou, Y., Sullivan, J.H., Stoumbos, Z.G. & Young, J.C. (2003). Systematic patterns in T^2 charts, *Journal of Quality Technology* **35**, 47–58.

Related Article

Bayesian Control Charts; Change-Point Models; Degradation Models; Multivariate Charts for Variability; Multivariate Control Charts, Interpretation of; Multivariate Cumulative Sum (CUSUM) Chart; Multivariate Exponentially Weighted Moving Average (MEWMA) Control Chart; Multivariate Stochastic Orders and Aging; Reliability Growth Modeling; Statistical Process Control, Bayesian.

CAMIL FUCHS AND RON S. KENETT

Process Capability Indices, Bayesian Estimation of

Process Capability Indices

The process capability index, C_p , relates process variation to customer requirements in the form of a ratio

$$C_p = \frac{USL - LSL}{6\sigma} \quad (1)$$

where the difference between the upper specification limit (USL) and the lower specification limit (LSL) provides a measure of allowable process spread and 6σ , σ^2 being the process variance, a measure of actual process spread.

The failure of C_p to consider a target value resulted in the development of several indices. Those indices that assume the target (T) to be the midpoint of the upper and lower specification limits include

$$C_{pu} = \frac{USL - \mu}{3\sigma} \quad (2)$$

$$C_{pl} = \frac{\mu - LSL}{3\sigma} \quad (3)$$

$$C_{pk} = \min(C_{pl}, C_{pu}) \quad (4)$$

$$C_{pm} = \frac{USL - LSL}{6\sqrt{\sigma^2 + (\mu - T)^2}} \quad (5)$$

and

$$C_{pk} = (1 - k)C_p \quad (6)$$

where μ denotes the process mean, $k = 2|T - \mu| / (USL - LSL)$, $0 \leq k \leq 1$ and $LSL < \mu < USL$. The two definitions of C_{pk} are equivalent when $0 \leq k \leq 1$.

Generalized versions of these indices, that do not assume T to be the midpoint of the specification limits, include

$$C_{pu}^* = \frac{USL - T}{3\sigma} \left(1 - \frac{|T - \mu|}{USL - T} \right) \quad (7)$$

$$C_{pl}^* = \frac{T - LSL}{3\sigma} \left(1 - \frac{|T - \mu|}{T - LSL} \right) \quad (8)$$

$$C_{pk}^* = \min(C_{pl}^*, C_{pu}^*) \quad (9)$$

and

$$C_{pm}^* = \frac{\min[USL - T, T - LSL]}{3\sqrt{\sigma^2 + (\mu - T)^2}} \quad (10)$$

A hybrid of C_{pk}^* and C_{pm}^* is defined to be

$$C_{pmk} = \frac{\min[USL - \mu, \mu - LSL]}{3\sqrt{\sigma^2 + (\mu - T)^2}} \quad (11)$$

As all of these indices use a function of σ as a measure of actual process spread in their determination of process capability and, in practice, are translated into parts per million nonconforming, none should be used when the characteristic under investigation is not normally distributed. To illustrate, suppose that precisely 99.73% of the process measurements fall within the specification limits (i.e., a **process yield** of 2700 ppm) and the process is centered at the target. The values of C_p (as well as C_{pl} , C_{pu} , C_{pk} , C_{pm} , and C_{pmk}) are 0.5766, 0.7954, 1.0000, 1.2210, and 1.4030 when the measurements arise from a uniform, triangular, normal, logistic, and double exponential distribution respectively, yet the process yields are identical. Under these circumstances, the assumption that the underlying process is normally distributed is critical to both sampling results and parameter values.

A Bayesian Approach

The general Bayesian approach assumes the parameter(s) to be stochastic and uses sampling results to adjust an assumed prior distribution to develop a posterior distribution for the associated parameter(s). In the case of process capability indices, and following the assumption that the observations must arise from a **normal distribution**, the likelihood function for a sample of size n is

$$L(\mu, \sigma^2) = (2\pi\sigma^2)^{-n/2} e^{-\sum_{i=1}^n (x_i - \mu)^2 / (2\sigma^2)} \quad (12)$$

The noninformative prior density function

$$h(\mu, \sigma^2) = 1/\sigma \quad -\infty < \mu < \infty, \quad 0 < \sigma < \infty \quad (13)$$

results in the posterior distribution

$$f(\sigma | x_1, x_2, x_3, \dots, x_n)$$

2 Process Capability Indices, Bayesian Estimation of

$$= \frac{2}{\Gamma[(n-1)/2]} \left(\frac{(n-1)s^2}{2} \right)^{\frac{n-1}{2}} \sigma^{-n} e^{-\frac{(n-1)s^2}{2\sigma^2}} \quad (14)$$

where $s^2 = \sum_{i=1}^n (x_i - \bar{x})^2 / (n-1)$ and $\bar{x} = \sum_{i=1}^n x_i / n$. As each of the process capability indices is a function of the process parameters μ and σ , and judged based upon a value (e.g., $C_p > 1$) this allows equation (14) to be written as

$$\Pr(C_p > c | x_1, x_2, x_3, \dots, x_n) \quad (15)$$

In the case of C_p , Cheng and Spiring [1] rewrote equation (15)

$$\begin{aligned} \Pr(C_p > c | x_1, x_2, x_3, \dots, x_n) &= \Pr\left(\sigma < \left(\frac{USL - LSL}{6c}\right) | x_1, x_2, x_3, \dots, x_n\right) \\ &= \int_0^{(USL-LSL)/(6c)} f(\sigma | x_1, x_2, \dots, x_n) d\sigma \\ &= \int_0^{(USL-LSL)/(6c)} \frac{2}{\Gamma[(n-1)/2]} \\ &\quad \times \left(\frac{(n-1)s^2}{2}\right)^{\frac{n-1}{2}} \sigma^{-n} e^{-\frac{(n-1)s^2}{2\sigma^2}} d\sigma \quad (16) \end{aligned}$$

resulting in a credible interval that depends upon n, s^2, USL, LSL and c .

Program Cp: Mathematica v5.2 Code for determining $\Pr(C_p > c | x_1, x_2, x_3, \dots, x_n)$

```
In[1]: n:= ;
s2:= ;
specdiff:= ;
c:= ;
N[Integrate[(((2/Evaluate[Gamma
[(n-1)/2]]((n-1)s2/2)^(n-1/2))(x^(-n))
Exp[-((n-1)s2/2)/(x^2)]), {x,0, (specdiff/(6c))}]]]
```

The user must enter: **n** the sample size, **s2** the sample variance, **specdiff** = $(USL - LSL)$, and **c** the desired level of C_p .

Pearn and Wu [2] extended the credible interval approach for C_p to the case where $\hat{C}_p^*$ is essentially a pooled estimate of C_p based upon a series (m) of repeated samples (n). Minimum values for $\hat{C}_p^*$

have been developed for $m = 2(2)10, n = 10(5)30$, and $p = 0.95, 0.975, 0.99$ providing a Bayesian alternative to the capability chart developed in [3].

In the case of C_{pm} and the noninformative prior distribution defined in equation (13), the conditional distribution of equation (14) becomes

$$\begin{aligned} f(\sigma' | x_1, x_2, x_3, \dots, x_n) &= \frac{2\sigma'}{\Gamma[n/2]} \left(\frac{(n-1)s^2 + n(\bar{x} - \mu)^2}{2} \right)^{\frac{n}{2}} \\ &\quad \times (\sigma'^2 - (\mu - T)^2)^{-\left(\frac{n}{2}+1\right)} \\ &\quad \times e^{\frac{(n-1)s^2 + n(\bar{x} - \mu)^2}{-2(\sigma'^2 - (\mu - T)^2)}} \quad (17) \end{aligned}$$

where $\sigma'^2 = \sigma^2 + (\mu - T)^2$ and $(\mu - T)^2 < \sigma'^2 < \infty$. The resulting credible interval depends upon $n, x, s^2, USL, LSL, T, \mu$ and c and is of the form [4]

$$\begin{aligned} \Pr(C_{pm} > c | x_1, x_2, x_3, \dots, x_n) &= \int_0^{(USL-LSL)/(6c)} f(\sigma' | x_1, x_2, \dots, x_n) d\sigma' \\ &= \int_0^{(USL-LSL)/(6c)} \frac{2\sigma'}{\Gamma[n/2]} \\ &\quad \times \left(\frac{(n-1)s^2 + n(\bar{x} - \mu)^2}{2} \right)^{\frac{n}{2}} \\ &\quad \times (\sigma'^2 - (\mu - T)^2)^{-\left(\frac{n}{2}+1\right)} e^{\frac{(n-1)s^2 + n(\bar{x} - \mu)^2}{-2(\sigma'^2 - (\mu - T)^2)}} d\sigma' \quad (18) \end{aligned}$$

Program Cpm: Mathematica v5.2 Code for determining $\Pr(C_{pm} > c | x_1, x_2, x_3, \dots, x_n)$

```
In[1]: n:= ;
s2:= ;
mu:= ;
T:= ;
xbar:= ;
specdiff:= ;
c:= ;
Integrate[(2x/Evaluate[Gamma[n/2]]]
(((n-1)s2+n(xbar-mu)^2)/2)^(n/2)(x^2-
(mu-T)^2)^(-(n/2+1)) Exp[(((n-1)s2+n
(xbar-mu)^2)/(-2((x^2)-(mu-T)^2))],
{x,0,(specdiff/(6c))}]]]
```

The user must enter: **n** the sample size, **s2** the sample variance, **mu** the population average, **T** the target value, **xbar** = the sample average, **specdiff** = ($USL - LSL$), and **c** the desired level of Cpm .

In the case of Cpk , Shiau *et al.* [5] used the same approach to develop a credible interval restricted to the case where $m = (USL + LSL)/2$ by considering

$$\begin{aligned} \Pr(Cpk > c | x_1, x_2, x_3, \dots, x_n) \\ = \Pr\left(\sigma < \left(\frac{USL - LSL}{6c}\right) \middle| x_1, x_2, x_3, \dots, x_n\right) \end{aligned} \quad (19)$$

which for the posterior distribution defined in equation (14) becomes

$$\begin{aligned} \Pr(Cpk > c | x_1, x_2, x_3, \dots, x_n) \\ = \int_0^{(USL-LSL)/(6c)} f(\sigma | x_1, x_2, \dots, x_n) d\sigma \\ = \int_0^{(USL-LSL)/(6c)} \frac{2}{\Gamma[n/2]} \left(\frac{ns'^2}{2}\right)^{\frac{n}{2}} \\ \times \sigma^{-(n+1)} e^{-\frac{ns'^2}{2\sigma^2}} d\sigma \end{aligned} \quad (20)$$

where $s'^2 = \sum_{i=1}^n ((x_i - m)^2/n)$, resulting in a credible interval that depends upon n, s'^2, USL, LSL and c . The authors further simplify equation (20) to get

$$\begin{aligned} \Pr(Cpk > c | x_1, x_2, x_3, \dots, x_n) \\ = \int_b^\infty \frac{1}{\Gamma[n/2]} y^{(n/2)-1} e^{-y} dy \end{aligned} \quad (21)$$

for $b = n/2 (c/\hat{C}pk)^2$.

Program Cpk: Mathematica v5.2 Code for determining $\Pr(Cpk > c | x_1, x_2, x_3, \dots, x_n)$

```
In[1]: n:= ;
      cpkh:= ;
      c:= ;
      1-Integrate[(1/Evaluate[Gamma[n/2]])
      (y^((n/2)-1))
      Exp[-y],{y,0,(n/2)(c/cpkh)^2}]
```

The user must enter: **n** the sample size, **cpkh** = $(USL - LSL)/(6s')$, and **c** the desired level of Cpk .

In the case of Cpm^* , Lin *et al.* [6] developed a credible interval, $\Pr(Cpm^* > c | x_1, x_2, x_3, \dots, x_n) = p$, of the form

$$p = \begin{cases} \left\{ \int_0^{t^*} \frac{\Phi(b_1(y) + b_U^*(y)) - \Phi(b_1(y) - b_L^*(y))}{\gamma^\alpha y^{\alpha+1} \Gamma(\alpha)} \exp\left(-\frac{1}{\gamma y}\right) dy \right\} & \text{for } \bar{x} < T \\ \left\{ \int_0^{t^*} \frac{\Phi(b_1(y) + b_U^*(y)) - \Phi(b_1(y) - b_L^*(y))}{\gamma^\alpha y^{\alpha+1} \Gamma(\alpha)} \exp\left(-\frac{1}{\gamma y}\right) dy \right\} & \text{for } \bar{x} > T \end{cases} \quad (22)$$

where $t^* = 2[\min[USL - T, T - LSL]]/(2c)$
 $\times [(n-1)s^2 + n(\bar{x} - T)^2]^{-1}$,

Φ is the cumulative distribution function of the $N(0, 1)$,

$$b_1(y) = \left[(2/y)(1 - (1/(1 + ((n(\bar{x} - T)^2)/((n-1)s^2)))) \right]^{0.5},$$

$$\begin{aligned} b_U^*(y) &= [2(USL - T)/(USL - LSL)] \\ &\times [n((2(\min[USL - T, T - LSL])/(3c))^2 \\ &\times ((n-1)s^2 + n(\bar{x} - T)^2)^{-1})/y - 1]^{0.5}, \\ \gamma^2 &= (\bar{x} - T)^2/s^2, \end{aligned}$$

$$b_2(y) = [n((2\hat{C}pm^2/nyc^2) - 1)]^{0.5}, \text{ and}$$

$$\begin{aligned} b_L^*(y) &= [2(T - LSL)/(USL - LSL)] \\ &\times [n((2(\min[USL - T, T - LSL])/(3c))^2 \\ &\times ((n-1)s^2 + n(\bar{x} - T)^2)^{-1})/y - 1]^{0.5}, \end{aligned}$$

resulting in an interval that depends upon n, s^2, γ, USL, LSL and c .

The authors indicate that although equation (22) reduces to

$$p = \left\{ \int_0^{t^*} \frac{\Phi(b_1(y) + b_2(y)) - \Phi(b_1(y) - b_2(y))}{\gamma^\alpha y^{\alpha+1} \Gamma(\alpha)} \times \exp\left(-\frac{1}{\gamma y}\right) dy \right\} \quad (23)$$

4 Process Capability Indices, Bayesian Estimation of

Table 1 Process specifications and summaries

Sampling summaries			Specifications			Process capability estimates			
n	$\bar{x}$	s	T	USL	LSL	$\hat{C}_p$	$\hat{C}_{pm}$	$\hat{C}_{pk}$	$\hat{C}_{pmk}$
50	2.08	0.022	2.08	2.15	2.01	1.06	1.06	1.06	1.06

Table 2 Associated credible intervals

C_p		C_{pm}		C_{pk}	
$\hat{C}_p$	$\Pr(C_p > 1 \hat{C}_p)^{(a)}$	$\hat{C}_{pm}$	$\Pr(C_{pm} > 1 \hat{C}_{pm})^{(b)}$	$\hat{C}_{pk}$	$\Pr(C_{pk} > 1 \hat{C}_{pk})^{(c)}$
1.06	0.6926	1.06	0.7279	1.06	0.6929

^(a)From **Program Cp**

^(b)From **Program Cpm** for $\mu = T$

^(c)From **Program Cpk** for $m = (USL + LSL)/2$

where $t = 2\hat{C}_{pm}^2/(nc^2)$ when T is the midpoint of the specification limits, and the equivalent of equation (18) when $m = (USL + LSL)/2$, equation (22), is “complicated and computationally inefficient”.

Extension of the credible interval approach to C_{pmk} proves to be a difficult problem mathematically. Pearn and Lin [7] have examined some properties of C_{pmk} from a Bayesian estimation perspective.

Example

Braverman [8] discussed an example from a **quality control** department that wanted to assess the capability of a shaft manufacturing process. The shaft specifications were $USL = 2.15$ and $LSL = 2.01$ with $T = 2.08$. A random sample of 50 shafts was selected and measured with the following results (see Table 1)

Using the above **Mathematica** programs for C_p , C_{pm} , and C_{pk} , the associated Bayesian credible interval probabilities are listed in Table 2 suggesting that the probability that the process is deemed capable given the sampling results is approximately 0.7 for each of C_p , C_{pm} , and C_{pk} .

Comments

The amount of material associated with the Bayesian approach for assessing process capability represents a small portion of the work in the area of drawing

inferences regarding a process capability. A list of additional related manuscripts that deal with specific situations has been included. Much like the general development, the evolution of Bayesian inference in the area of process capability assessment appears to be dominated by academic interests and has had little practitioner input.

References

- [1] Cheng, S.W. & Spiring, F.A. (1989). Assessing process capability: a Bayesian approach, *IIE Transactions* **21**, 97–98.
- [2] Pearn, W.L. & Wu, C.-W. (2005). A Bayesian approach for assessing process precision based on multiple samples, *European Journal of Operational Research* **165**(3), 685–695.
- [3] Spiring, F.A. (1995). Process capability: a total quality management tool, *Total Quality Management* **6**, 21–33.
- [4] Fan, S.-K. & Kao, C.-K. (2004). Lower Bayesian confidence limits on the process capability index C_{pm} : a comparative study, *Quality and Reliability Engineering International* **5**, 444–452.
- [5] Shiau, J.H., Chiang, C. & Hung, H. (1999). A Bayesian procedure for process capability assessment, *Quality and Reliability Engineering International* **15**, 369–378.
- [6] Lin, G.H., Pearn, W.L. & Yang, Y.S. (2005). A Bayesian approach to obtain a lower bound for the C_{pm} capability index, *Quality and Reliability Engineering International* **21**, 655–668.
- [7] Pearn, W.L. & Lin, G.H. (2003). A Bayesian-like estimator of the process capability index C_{pmk} , *Metrika* **57**, 303–312.
- [8] Braverman, J.D. (1981). *Fundamentals of Statistical Quality Control*, Prentice-Hall, Reston.

Further Reading

- Bernardo, J.M. & Irony, T.Z. (1996). A general multivariate Bayesian process capability index, *The Statistician: Journal of the Institute of Statisticians* **45**, 487–502.
- Lin, G.H., Pearn, W.L. & Chen, S. (2002). A note on the distributions of the estimated C_{pk} with asymmetric tolerances, *Far East Journal of Theoretical Statistics* **6**, 25–37.
- van der Merwe, A.J. & Chikobvu, D. (2004). Bayesian estimation of the process capability index C_{pk} , *South African Statistical Journal* **38**, 139–158.
- Niverthi, M. & Dey, D.K. (2000). Multivariate process capability: a Bayesian perspective, *Communications in Statistics: Simulation and Computation* **29**, 667–687.
- Pearn, W.L. & Chen, K.S. (1996). A Bayesian-like estimator of C_{pk} , *Communications in Statistics: Simulation and Computation* **25**, 321–329.
- Pearn, W.L. & Liao, M.-Y. (2005). One-sided process capability assessment in the presence of measurement errors, *Quality and Reliability Engineering International* **22**(7), 771–785.
- Shiau, J.H., Hung, H. & Chiang, C. (1999). A note on Bayesian estimation of process capability indices, *Statistics and Probability Letters* **45**, 215–224.
- Singh, H.P. & Saxena, S. (2005). Bayesian shrinkage estimation of process capability index C_p , *Communications in Statistics: Theory and Methods* **34**, 205–228.
- de Souza Borges, W. & Ho, L.L. (2001). A fraction defective based capability index, *Quality and Reliability Engineering International* **17**, 447–458.

FRED SPIRING AND SMILEY CHENG

Process Capability Indices, Comparison of

Measuring Process Capability

Historically process capability was synonymous with process variation and expressed as either 6σ or the population range. Neither of these measures consider customer requirements nor permit general comparisons among processes as both measures are unit dependent. Juran [1] suggests that Japanese companies have initiated the use of process capability indices by relating process variation to customer requirements in the form of the ratio

$$C_p = \frac{USL - LSL}{6\sigma} \quad (1)$$

where the difference between the upper specification limit (USL) and the lower specification limit (LSL) provides a measure of allowable process spread (i.e., customer requirements) and 6σ , σ^2 being the process variance, a measure of actual process spread (i.e., process performance). Bissell [2] suggests that the British Standards Institution proposed an analogous measure referred to as the relative precision index, which was withdrawn and replaced by C_p in 1942.

Incorporating customer specification limits into the assessment of process capability results in a more meaningful measure that fosters comparisons across all types of processes. However C_p uses only the customer's USL and LSL in its assessment and fails to consider a target (or nominal) value. The three processes depicted in Figure 1 have identical values of σ^2 and therefore identical values of C_p . However, processes 2 and 3 deviate from the target (T), and additional work and/or costs due to these departures may be incurred. As a result processes 2 and 3 are considered less capable of meeting customer requirements than process 1.

Processes with small variability, but poor proximity to the target, have sparked the derivation of several indices that incorporate targets into their assessment. The most common of these measures assume the target (T) to be the midpoint of the specification limits and include

$$C_{pu} = \frac{USL - \mu}{3\sigma} \quad (2)$$

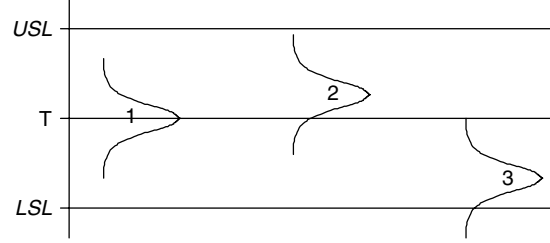


Figure 1 Three processes with identical values of C_p

$$C_{pl} = \frac{\mu - LSL}{3\sigma} \quad (3)$$

$$C_{pk} = \min(C_{pl}, C_{pu}) \quad (4)$$

$$C_{pm} = \frac{USL - LSL}{6\sqrt{\sigma^2 + (\mu - T)^2}} \quad (5)$$

and

$$C_{pk} = (1 - k)C_p \quad (6)$$

where μ is the process mean, $k = \frac{2|T - \mu|}{USL - LSL}$, $0 \leq k \leq 1$ and $LSL < \mu < USL$. The two definitions of C_{pk} are presented interchangeably and are equivalent when $0 \leq k \leq 1$.

The generalized analogs of these measures do not assume T to be the midpoint of the specifications and are

$$C_{pu}^* = \frac{USL - T}{3\sigma} \left(1 - \frac{|T - \mu|}{USL - T} \right) \quad (7)$$

$$C_{pl}^* = \frac{T - LSL}{3\sigma} \left(1 - \frac{|T - \mu|}{T - LSL} \right) \quad (8)$$

$$C_{pk}^* = \min(C_{pl}^*, C_{pu}^*) \quad (9)$$

and

$$C_{pm}^* = \frac{\min[USL - T, T - LSL]}{3\sqrt{\sigma^2 + (\mu - T)^2}} \quad (10)$$

A hybrid of these measures, C_{pmk} , is defined to be

$$C_{pmk} = \frac{\min[USL - \mu, \mu - LSL]}{3\sqrt{\sigma^2 + (\mu - T)^2}} \quad (11)$$

Individually C_{pu} , C_{pu}^* , C_{pl} , and C_{pl}^* consider only unilateral tolerances (i.e., USL or LSL respectively) in assessing process capability. Each uses 3σ as a measure of actual process spread, while the distance from where the process is centered (μ) to the USL (for C_{pu}) or LSL (for C_{pl}) is used as a measure of

2 Process Capability Indices, Comparison of

allowable process spread. Both C_{pu} and C_{pl} compare the length of one tail of the **normal distribution** (3σ) with the distance between the process mean and the respective specification limit. In the case of bilateral tolerances, C_{pu} and C_{pl} have an inverse relationship and individually do not provide a complete assessment of process capability. However, conservatively taking the minimum of C_{pu} and C_{pl} results in the bilateral tolerance measure defined as C_{pk} . Similarly C_{pu}^* and C_{pl}^* use adjusted distances between their respective specification limits and μ that incorporate the target in assessing allowable spread.

C_{pm} incorporates a measure of proximity to the target by replacing the process variance in the definition of C_p , with the process mean square error around the target. C_p and C_{pm} are identical when the process is centered at the target (i.e., $\mu = T$). But since C_p is not a valid measure of process capability when the process is not centered at the target, C_{pm} dominates C_p as a measure of process capability.

C_{pl}^* , C_{pu}^* , C_{pk}^* , C_{pm}^* , and C_{pmk} form the basis for a large number of process capability indices, sometimes referred to as third generation. Effectively these indices extend the class of allowable/customer process requirements versus actual process spread/performance by allowing the target to be other than the midpoint of the specification limits.

By definition, each of C_p , C_{pl} , C_{pu} , C_{pk} , C_{pm} , their generalized analogues, and C_{pmk} are unitless, thereby fostering comparisons among and within processes, regardless of the underlying mechanics of the product or service being monitored. In all cases, as process performance improves, either through reductions in variation and/or moving closer to the target, these indices increase in magnitude. As process performance improves relative to customer requirements, customer satisfaction increases as the process has a greater ability to be near the target. In all cases larger index values indicate a more capable process. Using the indices to assess the ability of the process to meet customer expectations, with larger values of the indices indicating a higher customer satisfaction and lower values, a poorer customer satisfaction is consistent with the concept of “fitness for use”.

A unified index, C_{pw} [3] of the form

$$C_{pw} = \frac{USL - LSL}{6\sqrt{\sigma^2 + w(\mu - T)^2}} \quad (12)$$

can, by varying w , be used to represent a wide spectrum of capability indices.

Allowing w to take on various values permits C_{pw} to assume equivalent computational algorithms for a variety of indices, including C_p , C_{pm} , C_{pk} , and C_{pmk} . For example, setting $w = 0$, results in

$$C_{pw} = \frac{USL - LSL}{6\sqrt{\sigma^2}} = C_p \quad (13)$$

while for $w = 1$,

$$C_{pw} = \frac{USL - LSL}{6\sqrt{\sigma^2 + (\mu - T)^2}} = C_{pm} \quad (14)$$

Defining $d = \frac{USL - LSL}{2}$, $a = \mu - \frac{USL + LSL}{2}$ and $p = \frac{|\mu - T|}{\sigma}$ then

$$w = \begin{cases} \left(\frac{d^2}{(d - |a|)^2} - 1 \right) \frac{1}{p^2} & 0 < p \\ 0 & \text{elsewhere} \end{cases} \quad (15)$$

allows $C_{pw} = C_{pk}$.

Similarly,

$$w = \begin{cases} \frac{k(2 - k)}{(1 - k)^2 p^2} & 0 < k < 1 \\ 0 & \text{elsewhere} \end{cases} \quad (16)$$

where $k = \frac{2|T - \mu|}{USL - LSL}$, allows $C_{pw} = C_{pk}^*$, while

$$w = \begin{cases} \left(\frac{d}{d - |a|} \right)^2 \left(\frac{1}{p^2} + 1 \right) - \frac{1}{p^2}, & 0 < p \\ 0 & \text{elsewhere} \end{cases} \quad (17)$$

elsewhere results in $C_{pw} = C_{pmk}$. Table 1 includes several weights including the appropriate weights associated with $C_p(u, v)$ [4].

It is important to note that the universal assumptions associated with the unified index C_{pw} include that the underlying process is stable and the process measurements are normally distributed.

Interpreting Process Capability Measures

Process capability measures have traditionally been used to provide insights into the number (or proportion) of a nonconforming product (i.e., yield). Practitioners cite a C_p value of 1 as representing 2700 ppm

Table 1 Other C_{pw} weights and the associated index

w	Resulting index
$\begin{cases} \frac{6C_p - p}{(3C_p - p)^2 p}, & 0 < \frac{p}{3} < C_p \\ 0, & \text{elsewhere} \end{cases}$	$C_{pk}^* = (1 - k) \frac{USL - LSL}{6\sigma}$
$\begin{cases} \left(\frac{d}{d - a } \right)^2 \left(\frac{1}{p^2} + 1 \right) - \frac{1}{p^2}, & 0 < p \\ 0, & \text{elsewhere} \end{cases}$	$C_{pmk} = \frac{\min(USL - \mu, \mu - LSL)}{3\sqrt{\sigma^2 + (\mu - T)^2}}$ $= \frac{d - \mu - M }{3\sqrt{\sigma^2 + (\mu - T)^2}}$
$\begin{cases} \left(\frac{d}{d - u a } \right)^2 \left(\frac{1}{p^2} + v \right) - \frac{1}{p^2}, & 0 < p \\ 0, & \text{elsewhere} \end{cases}$	$C_p(u, v) = \frac{d - u \mu - M }{3\sqrt{\sigma^2 + v(\mu - T)^2}}$

nonconforming, while 1.33 represents 63 ppm; 1.66 corresponds to 0.6 ppm; and 2 indicates <0.1 ppm. C_{pk} has similar connotations with a C_{pk} of 1.33 representing a maximum of 63 ppm nonconforming. Practitioners use the value of the process capability index and its associated number, nonconforming, to identify capable processes. A process with a C_p greater than or equal to 1 has traditionally been deemed capable, while a C_p of less than 1 indicates that the process is producing more than 2700 ppm nonconforming and used as an indication that the process is not capable of meeting customer requirements. In the case of C_{pk} , the auto industry frequently uses 1.33 as a benchmark in assessing the capability of a process.

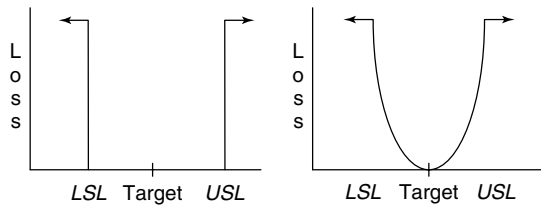
Inherent in any discussion of yield as a measure of process capability is the assumption that product produced just inside the specification limit is of equal quality to that produced at the target. This is equivalent to assuming a square-well loss function [5] for the quality variable (see Figure 2).

In practice the magnitudes of C_p , C_{pl} , C_{pu} , C_{pk} , their generalized analogues, and C_{pmk} are interpreted as a measure of nonconforming and therefore are represented by the square-well loss function. Any change in the magnitude of these indices (holding the customer requirements constant) is due to

changes in the distance between the specification limits and the process mean. As Boyles [6] emphatically points out “ C_{pk} does not in itself say anything about the distance between μ and T ” and “is essentially a measure of **process yield** only.” By design C_p , C_{pl} , C_{pu} , C_{pk} , C_{pl}^* , C_{pu}^* , C_{pk}^* , and C_{pmk} are used to identify changes in the amount of product beyond the specification limits (not proximity to the target) and therefore consistent with the square-well loss function.

Taguchi used the quadratic loss function to motivate the idea that a product imparts no loss, only if that product is produced at its target. He maintains that even small deviations from the target result in a loss of quality and that as the product increasingly deviates from its target, there are larger and larger losses in quality. This approach to quality and quality assessment is different from the traditional approach, where no loss in quality is assumed until the product deviates beyond its USL and LSL (i.e., square-well loss function). Taguchi’s philosophy highlights the need to have small variability around the target. Clearly in this context, the most capable process will be one that produces its entire product at the target, with the next best being the process with the smallest variability around the target.

The motivation for C_{pm} and C_{pm}^* does not arise from examining yield, but from looking at the ability of the process to be in the neighborhood of the target. This motivation has little to do with the number of nonconforming although upper bounds on the number of nonconforming can be determined for numerical values of C_{pm} [7]. Several authors [3, 6, 8] have discussed the relationship between C_{pm} and the quadratic loss function and its affinity with


Figure 2 Square-well and modified quadratic loss functions

4 Process Capability Indices, Comparison of

the philosophies that support a loss in quality for any departure from the target. The loss associated with the quality variable under investigation is then depicted using a modified quadratic loss function (see Figure 2).

C_{pk} , C_{pm} , and C_{pmk} have different functional forms, are represented by different loss functions, and have different relationships with C_p as the process drifts from the target. C_{pk} and C_{pm} are often called second-generation measures of process capability whose motivations arise directly from the inability of C_p to consider the target value. The differences in their associated loss functions demarcate the measures, while the magnitudinal relationship between C_p and C_{pk} , C_{pm} , C_{pmk} are also different. C_{pk} , C_{pm} , and C_{pmk} are functions of C_p that attempt to penalize the process for not being centered on the target (T). Expressing C_{pm} , C_{pk} , and C_{pmk} in the following manner

$$C_{pm} = \frac{1}{\sqrt{1 + \left(\frac{\mu - T}{\sigma}\right)^2}} C_p \quad (18)$$

$$C_{pk} = \left(1 - \frac{2|\mu - T|}{USL - LSL}\right) C_p \quad (19)$$

and

$$C_{pmk} = \left(1 - \frac{2|\mu - T|}{USL - LSL}\right) \left(\frac{1}{\sqrt{1 + p^2}}\right) C_p \quad (20)$$

illustrates the “penalizing” relationship between C_p and C_{pm} , C_{pk} , C_{pmk} respectively. As the process mean drifts from the target (measured by $p = \frac{|\mu - T|}{\sigma}$), C_{pm} , C_{pk} , and C_{pmk} decline as a percentage of C_p (see Figure 3). In the case of C_{pm} , this relationship is independent of the magnitude of C_p . On the other hand, the decline in C_{pk} and C_{pmk} as a percentage of C_p depends on the magnitude of C_p . For example, in Figure 3, $C_{pk}(1)$, $C_{pk}(2)$, and $C_{pk}(3)$ represent the relationships between C_{pk} and C_p for $C_p = 1$, $C_p = 2$ and $C_p = 3$, respectively. Similarly $C_{pmk}(1)$, $C_{pmk}(2)$, and $C_{pmk}(3)$ represent the relationships between C_{pmk} and C_p for $C_p = 1$, $C_p = 2$, and $C_p = 3$.

C_{pk} , C_{pmk} , and C_{pm} have different functional forms, are represented by different loss functions, and have different relationships with C_p as the process drifts from the target. Hence, although C_{pm} and C_{pk} are lumped together as second-generation measures and C_{pmk} as a hybrid of the two, all are different in their development and assessment of process capability. A discussion on the relationship of capability indices (C_p , C_{pk} , and C_{pm}) among themselves, linked to specifications at the supplier level and specifications to the assembly

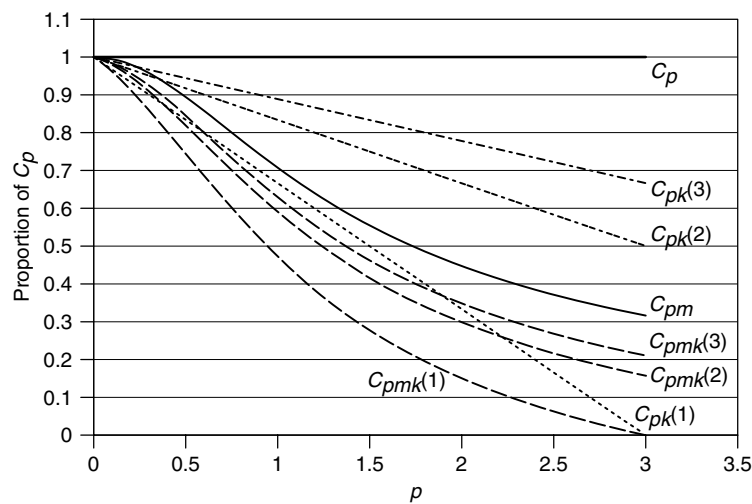


Figure 3 Comparison of C_p , C_{pk} , C_{pm} , and C_{pmk}

level is cited in [9]. Similarly C_{pk}^* and C_{pmk}^* differ from C_{pm}^* in their assessment of process capability.

Effects of Nonnormality

If the process measurements do not arise from a normal distribution, none of the indices discussed provides valid measures of yield. Each index uses a function of σ as a measure of actual process spread in its determination of process capability. But as many authors [10–13] have pointed out, although the standard deviation has become synonymous with the term dispersion, its physical meaning need not be the same for different families of distributions, or for that matter, within a family of distributions. Therefore the actual process spread (a function of 6σ) does not provide a consistent meaning over various distributions. To illustrate, suppose that precisely 99.73% of the process measurements fall within the specification limits (i.e., a yield of 2700 ppm). The values of C_p , C_{pl} , C_{pu} , C_{pk} , C_{pm} , and C_{pmk} are 0.5766, 0.7954, 1.0000, 1.2210, and 1.4030 respectively when the measurements arise from a uniform, triangular, normal, logistic, and double exponential distribution. As long as 6σ carries a yield interpretation when assessing process capability (i.e., is translated into ppm nonconforming), none of the indices should be used if the underlying process distribution is not normal.

If we assume process capability assessments to be studies of the ability of the process to produce product around the target, then C_{pm}^* and C_{pmk}^* provide practitioners with an assessment of capability regardless of the distribution associated with the measurements. Clustering around the target, rather than a measure of nonconforming releases the physical meaning attached to 6σ . The denominators of C_{pm}^* and C_{pmk}^* then provide a measure of the clustering around the target and compare this with customer tolerance.

Eliminating the physical meaning of ppm and investigating Euclidean distances around the target allows C_{pm} and C_{pmk} to compare the capability of various processes, or processes over time, regardless of the underlying distribution. The underlying distribution will impact the inferences that can be made from samples gathered from the population; however, the population parameter is no longer distributionally sensitive.

References

- [1] Juran, J.M. (1979). *Quality Control Handbook*, McGraw-Hill, New York.
- [2] Bissell, A.F. (1990). How reliable is your capability index? *Applied Statistics* **39**, 331–340.
- [3] Spiring, F.A. (1997). A unifying approach to process capability indices, *Journal of Quality Technology* **29**, 49–58.
- [4] Vannman, K. (1995). A unified approach to capability indices, *Statistica Sinica* **5**, 805–820.
- [5] Tribus, M. & Szonyi, G. (1989). An alternate view of the Taguchi approach, *Quality Progress* **22**(5), 46–52.
- [6] Boyles, R.A. (1991). The Taguchi capability index & corrigenda, *Journal of Quality Technology* **23**, 17–26.
- [7] Spiring, F.A. (1991). A new measure of process capability, *Quality Progress* **24**(2), 57–61.
- [8] Johnson, T. (1992). The relationship of cpm to squared error loss, *Journal of Quality Technology* **24**, 211–215.
- [9] Parlar, M. & Wesolowsky, G.O. (1999). Specification limits, capability indices, and process centering in assembly manufacturing, *Journal of Quality Technology* **31**, 317–325.
- [10] Hoaglin, D.C., Mosteller, F. & Tukey, J.W. (1983). *Understanding Robust and Exploratory Data Analysis*, John Wiley & Sons, New York.
- [11] Mosteller, F. & Tukey, J.W. (1977). *Data Analysis and Regression*, Addison-Wesley, Reading.
- [12] Tukey, J.W. (1970). A survey of sampling from contaminated distributions, *Contributions to Probability and Statistics*, Stanford University Press, Stanford, pp. 448–485.
- [13] Huber, P.J. (1977). *Robust Statistical Procedures*, Society for Industrial and Applied Mathematics, Philadelphia.

FRED SPIRING

Process Capability Plots

Introduction

When measuring the capability of a manufacturing process, some form of process capability index is often used (*see Process Capability Indices, Comparison of*). Vännman [1] unified the most common indices by a family of capability indices, defined as

$$C_p(u, v) = \frac{d - u|\mu - T|}{3\sqrt{\sigma^2 + v(\mu - T)^2}} \quad (1)$$

where $u \geq 0, v \geq 0, \mu$ is the process mean and σ is the process standard deviation of the in-control process. In equation (1) the target value T equals the midpoint of the specification interval $[LSL, USL]$. Furthermore, $d = (USL - LSL)/2$, i.e. half the length of the specification interval. Setting $u = 0$ or 1 and $v = 0$ or 1 in equation (1) we obtain $C_p(0, 0) = C_p$, $C_p(1, 0) = C_{pk}$, $C_p(0, 1) = C_{pm}$, and $C_p(1, 1) = C_{pmk}$. When using process capability indices, a process is defined to be capable if the process capability index exceeds a certain threshold value $k > 0$. Some commonly used values are $k = 1, 4/3, 5/3$.

Assume that we use an index defined by the family $C_p(u, v)$ in equation (1) and define the process to be capable if $C_p(u, v) > k$, for given values of u, v , and k . Different choices of u, v , and k impose different restrictions on the process parameters (μ, σ) . This is easily seen in a process capability plot, which is a contour plot of $C_p(u, v) = k$ as a function of μ and σ , or as a function of μ_t and σ_t , where

$$\mu_t = \frac{\mu - T}{d} \quad \text{and} \quad \sigma_t = \frac{\sigma}{d} \quad (2)$$

The contour curve is obtained by rewriting the index in equation (1) as a function of μ_t and σ_t , solving the equation $C_p(u, v) = k$ with respect to σ_t and then plotting σ_t as a function of μ_t . We easily find that $C_p(u, v) = k$ is equivalent to

$$\sigma_t = \sqrt{\frac{(1 - u|\mu_t|)^2}{9k^2} - v\mu_t^2},$$

$$|\mu_t| \leq \frac{1}{u + 3k\sqrt{v}}, \quad (u, v) \neq (0, 0) \quad (3)$$

When $u = v = 0$, i.e. when considering the index $C_p = k$, we have

$$\sigma_t = \frac{1}{3k}, \quad |\mu_t| \leq 1 \quad (4)$$

The reason for making the contour plot a function of (μ_t, σ_t) instead of a function of (μ, σ) is to obtain a plot where the scale is invariable, irrespective of the values of the specification limits. This is useful when considering several different characteristics (see the section titled “Several Characteristics in the Same Plot”).

To obtain the expression for σ_t in the case of C_{pk} , we let $u = 1$ and $v = 0$ in equation (3) and find that the contour curve is composed of two straight lines. To obtain the expression for σ_t in the case of C_{pm} , we let $u = 0$ and $v = 1$ in equation (3) and find that the contour curve is a semicircle (see Figure 1).

Values of the process parameters μ and σ which give (μ_t, σ_t) values inside the region bounded by the contour curve $C_p(u, v) = k$ and the μ_t axis will give rise to a $C_p(u, v)$ value larger than k , i.e. a capable process. We call this region the capability region. Furthermore, values of μ and σ which give (μ_t, σ_t) values outside this region will give a $C_p(u, v)$ value smaller than k , i.e. a non-capable process.

Note that the process capability plot is based on the process parameters, which are usually unknown in practice. Hence, this plot is a theoretical process capability plot useful to understand the restrictions the index puts on the process parameters. But it cannot immediately be used to deem if a process is capable, since it does not take the randomness of the estimators of the process parameter into account. The following sections discuss how to deal with the case in which the process parameters μ and σ are unknown.

Estimated Process Capability Plots

The process capability plot can now be generalized to an estimated process capability plot, which can be used to judge if a process is capable, when the process parameters μ and σ are unknown and need to be estimated. We assume that the studied characteristic of the process is normally distributed. Let $X_1, X_2, \dots, X_n$ be a random sample from a normal distribution with mean μ and variance σ^2 as the parametric process quality indicators.

2 Process Capability Plots

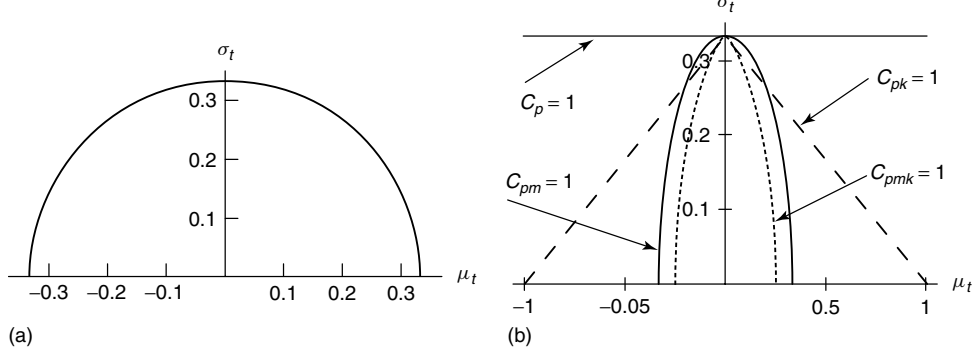


Figure 1 (a) The process capability plot for $C_{pm} = 1$. (b) The contour curves for the process capability indices C_p , C_{pk} , C_{pm} , and C_{pmk} when $k = 1$. The region bounded by the contour curve and the μ_t axis is the corresponding capability region

To obtain an appropriate decision rule we consider a hypothesis test with the null hypothesis $H_0: C_p(u, v) \leq k_0$ and the alternative hypothesis $H_1: C_p(u, v) > k_0$. As a test statistic we use the estimator $\hat{C}_p(u, v)$, which is obtained by estimating μ and σ^2 with their maximum-likelihood estimators, i.e.

$$\hat{\mu} = \bar{X} = \frac{1}{n} \sum_{i=1}^n X_i \quad \text{and} \quad \hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 \quad (5)$$

The null hypothesis is rejected, and the process is deemed capable, whenever $\hat{C}_p(u, v) > c_\alpha$, where the constant $c_\alpha > k_0$ is determined so that the **significance level** of the test is α . Details of how to obtain the critical value c_α can be found, e.g. in Vännman [2, 3]. Hubele and Vännman [4] showed that, when using the capability index C_{pm} , the critical value at significance level α is obtained as

$$c_\alpha = k_0 \sqrt{n / \chi_{\alpha, n}^2} \quad (6)$$

where $\chi_{\alpha, n}^2$ is the α th quantile from a χ^2 distribution with n **degrees of freedom**. When using the capability index C_{pk} it can be shown that the critical value at significance level α is obtained from the equation

$$\int_{-\infty}^{3k_0\sqrt{n}} F_\xi \left(\frac{(3k_0\sqrt{n} - z)^2}{9c_\alpha^2} \right) f_\eta(z) dz = \alpha \quad (7)$$

where F_ξ denotes the cumulative distribution function of a central χ^2 distribution with $n - 1$ degrees of

freedom and f_η denotes the probability density function of the standardized normal distribution $N(0, 1)$. For arbitrary values of $u \geq 0$ or $v \geq 0$ the critical value c_α is determined so that the probability $P(\hat{C}_p(u, v) > c_\alpha)$ is at most α for all values of (μ_t, σ_t) such that $C_p(u, v) = k_0$, where

$$\begin{aligned} P(\hat{C}_p(u, v) > c_\alpha) &= \int_{\frac{-\sqrt{n}/\sigma_t}{u+3c_\alpha\sqrt{v}}}^{\frac{\sqrt{n}/\sigma_t}{u+3c_\alpha\sqrt{v}}} F_\xi \left(\frac{(\sqrt{n}/\sigma_t - u|t|)^2}{9c_\alpha^2} - vt^2 \right) \\ &\quad \times f_\eta \left(t - \frac{\mu_t\sqrt{n}}{\sigma_t} \right) dt \end{aligned} \quad (8)$$

MATLAB codes for solving equation (7), as well as for determining the critical value c_α for arbitrary values of u and v using equation (8), can be obtained from the author.

The estimated process capability plot is then obtained by replacing μ , σ , and k in equations (2–4) with $\hat{\mu}$, $\hat{\sigma}$, and c_α , respectively, and plotting the contour curve $\hat{C}_p(u, v) = c_\alpha$ as a function of $\hat{\mu}_t$ and $\hat{\sigma}_t$. We can now use the estimated process capability plot to define an estimated capability region as the region bounded by the contour curve $\hat{C}_p(u, v) = c_\alpha$ and the $\hat{\mu}_t$ axis. To determine if a process can be deemed capable we first calculate

$$\hat{\mu}_t = \frac{\hat{\mu} - T}{d} \quad \text{and} \quad \hat{\sigma}_t = \frac{\hat{\sigma}}{d} \quad (9)$$

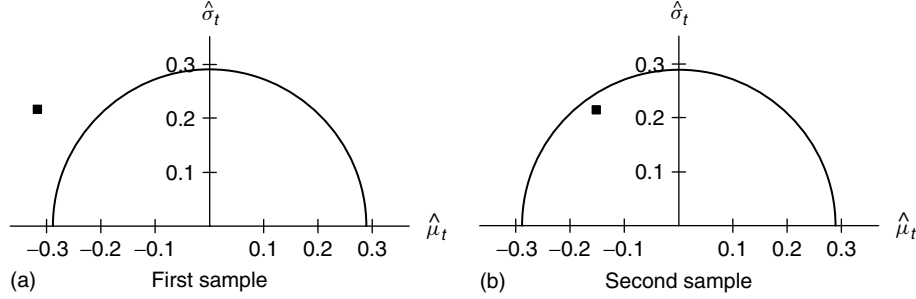


Figure 2 The estimated process capability plot defined by $\hat{C}_{pm} = c_{0.05} = 1.151$. In (a) the first sample with the observed value of $(\hat{\mu}_t, \hat{\sigma}_t)$ outside the estimated capability region, illustrates a non-capable process. In (b) the second sample, with observed value of $(\hat{\mu}_t, \hat{\sigma}_t)$ inside the estimated process capability region, illustrates a capable process

Then the point $(\hat{\mu}_t, \hat{\sigma}_t)$ is plotted in the estimated process capability plot. If this point is inside the estimated capability region, the process will be considered capable at significance level α .

To illustrate the estimated process capability plot we use an example from Vännman [2], where the process is defined capable if $C_{pm} > 1$ and the significance level is $\alpha = 5\%$. The example consists of two datasets, each of size 80, from a process where $T = 8.70$ and $d = 0.24$. When $n = 80$ and $k_0 = 1$ we get the critical value $c_{0.05} = 1.151$ from equation (6). First a random sample was taken from a process, giving rise to estimates of μ and σ being equal to $\hat{\mu} = 8.6234$ and $\hat{\sigma} = 0.0519$, respectively. Later, a second random sample was taken from the same process, after it had been adjusted. Then, from the second sample, the obtained estimates of μ and σ were equal to $\hat{\mu} = 8.6628$ and $\hat{\sigma} = 0.0516$, respectively.

On the basis of the sample data, we find the estimated values of μ_t and σ_t to be $\hat{\mu}_t = (\hat{\mu} - T)/d = -0.319$ and $\hat{\sigma}_t = \hat{\sigma}/d = 0.216$ for the first sample, and $\hat{\mu}_t = -0.155$ and $\hat{\sigma}_t = 0.215$ for the second sample. We then plot the contour curve $\hat{C}_p(u, v) = c_\alpha = 1.151$ together with the points with coordinates $(-0.319, 0.216)$ and $(-0.155, 0.215)$, for the first and second sample, respectively (see Figure 2).

We can see that for the first sample, the point $(-0.319, 0.216)$ is outside the estimated capability region and hence the process cannot be considered capable. For the second sample, the point $(-0.155, 0.215)$ is now inside the estimated capability region and hence the process can be deemed capable at the 5% significance level.

What is obvious from the graphical method used here is that the non-capability in Figure 2(a) is, to a

large extent, due to the deviation of the process mean from the target value. By looking at the estimated process capability plot we get information instantly about both the deviation from the target value and the process spread as well as information about the capability. We can also see how much the deviation from target must be reduced in order to achieve a capable process.

Safety Regions

As an alternative to the estimated process capability plot, we can use the theoretical process capability plot together with a safety region. A safety region is a region plotted in the theoretical process capability plot for drawing conclusions about process capability at a given significance level. It is built around the estimated point $(\hat{\mu}_t, \hat{\sigma}_t)$ to take the randomness of the estimator $(\hat{\mu}_t, \hat{\sigma}_t)$ into account. The size of the region is determined in such a way that it corresponds to significance level α .

Vännman [3] suggests putting a circular safety region around the estimated point $(\hat{\mu}_t, \hat{\sigma}_t)$. The circular safety region is defined as the region bounded by the circle

$$(\delta - \hat{\mu}_t)^2 + (\gamma - \hat{\sigma}_t)^2 = R^2 \hat{\sigma}_t^2 \quad (10)$$

The constant R depends on the sample size n and the significance level α . The process is deemed capable if the whole circular safety region in equation (11) is inside the capability region in the theoretical process capability plot defined by $C_p(u, v) = k_0$. The constant R is chosen so that the

4 Process Capability Plots

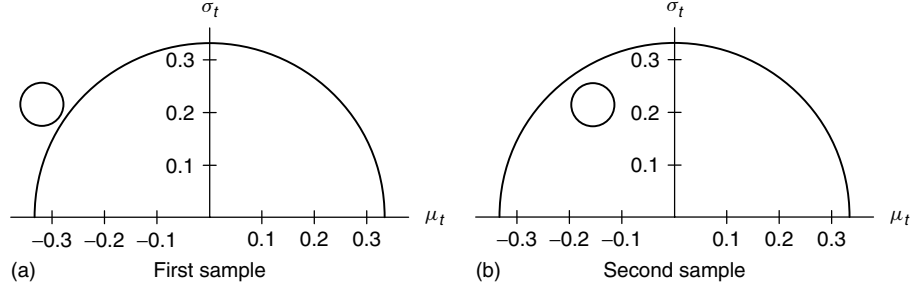


Figure 3 The process capability plot defined by $C_{pm} = 1$ together with the circular region. In (a) the process cannot be considered capable, but in (b) the whole circle is inside the capability region and hence the process is considered capable at 5% significance level

probability that the circular safety region is inside the capability region, given that $C_p(u, v) = k_0$, is at most equal to α . When using the index C_{pm} to define a capable process, the probability $\Pi(\mu_t, k_0)$ that the circular region is inside the capability region defined by $C_{pm} = k_0$, given that $C_{pm} = k_0$, can be expressed as follows.

$$\Pi(\mu_t, k_0) = \int_0^{Q(\mu_t, k_0)} F_\tau((S(\mu_t, k_0) - R\sqrt{y})^2 - y) f_\xi(y) dy \quad (11)$$

where

$$S(\mu_t, k_0) = \sqrt{\frac{n}{1 - 9k_0^2\mu_t^2}} \quad \text{and} \quad Q(\mu_t, k_0) = \frac{S(\mu_t, k_0)^2}{(R + 1)^2} \quad (12)$$

F_τ denotes the cumulative distribution function of a noncentral χ^2 distribution with 1 degree of freedom and noncentrality parameter λ , where λ is

$$\lambda = \frac{(\mu - T)^2 n}{\sigma^2} = \frac{\mu_t^2 n}{\sigma_t^2} \quad (13)$$

f_ξ denotes the probability density function of a central χ^2 distribution with $n - 1$ degrees of freedom. The probability in equation (11) is defined for $|\mu_t| < 1/(3k)$. The constant R is determined so that the probability $\Pi(\mu_t, k_0)$ in equation (11) is at most α for all values of μ_t such that $C_{pm} = k_0$. MATLAB codes for determining the R -value using equation (11) can be obtained from the author.

If we apply this circular safety region to the previous example, where $n = 80$ and $\alpha = 0.05$ and the process is defined as capable when $C_{pm} > 1$, we find $R = 0.1903$. In Figure 3 we see the circular regions plotted. The conclusions about the process capability are the same as before.

Deleryd and Vännman [5] suggest a rectangular safety region. This region is, however, less powerful than the circular safety region.

Several Characteristics in the Same Plot

Several characteristics of a process can be monitored in the same plot and at the same time the information on the location and spread of the process is retained. As an example consider the following situation. Three different quality characteristics A , B , and C are of interest. They have the following different

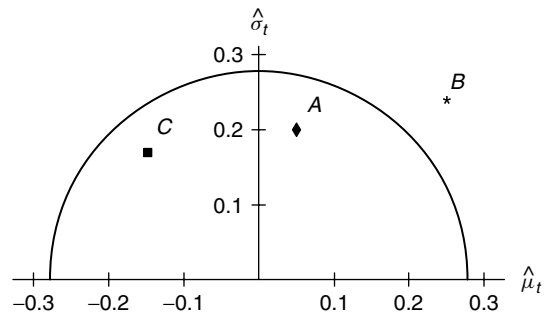


Figure 4 The estimated process capability plot defined by $\hat{C}_{pm} = c_{0.01} = 1.199$. A , B , and C indicate the observed value of $(\hat{\mu}_t, \hat{\sigma}_t)$ for the three characteristics A , B , and C , respectively

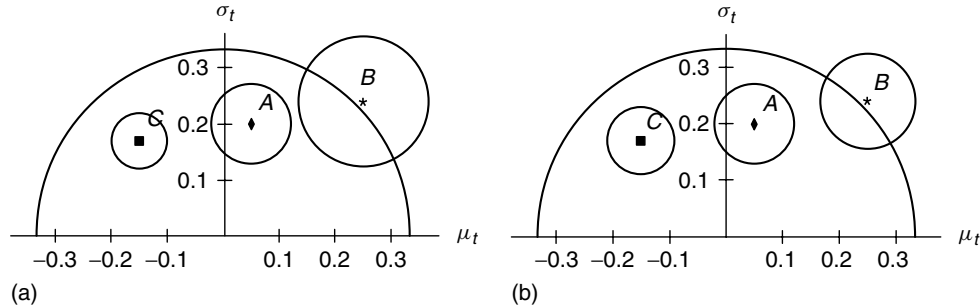


Figure 5 Process capability plots with circular regions for each of characteristics A , B , and C , each at 1% significance level. In (a) the sample sizes differ. In (b) all have sample size 50. The observed values of $(\hat{\mu}_t, \hat{\sigma}_t)$ are the same in (a) and (b)

specification intervals: A : [200, 210]; B : [12.0, 16.0]; C : [1.40, 2.20], with target values equal to the midpoints of each interval. They are assumed to be normally distributed and in statistical control. The process will be considered capable when all three characteristics are considered capable according to the definition that $C_{pm} > 1$.

A sample of size 50 is taken from each of the three characteristics and the point $(\hat{\mu}_t, \hat{\sigma}_t)$ is calculated for each of them. The decision rule is that the process will be considered capable if all three points $(\hat{\mu}_t, \hat{\sigma}_t)$ fall inside the estimated capability region, where the significance level for the estimated capability curve is 1%. Then, according to Bonferroni's inequality, the significance level for the decision rule will be at most 3%. The critical value, when $n = 50$ and $\alpha = 0.01$, is $c_{0.01} = 1.199$ according to equation (6). In Figure 4 we can see the results having the following observed values of $(\hat{\mu}_t, \hat{\sigma}_t)$: A : (0.05, 0.20); B : (0.25, 0.24); C : (-0.15, 0.17).

In Figure 4 we cannot consider the process to be capable, since characteristic B falls outside the semicircle. Figure 4 shows that from one single plot we get, for each of the three characteristics A , B , and C , information about the deviation from target and the standard deviation, as well as information on how to adjust the process to make it capable, if it is not capable.

When monitoring several characteristics in the estimated process capability plot, the same sample size is needed for each characteristic. With different sample sizes, the safety regions can be used instead. Figure 5 shows an example of this situation. In Figure 5(a) the sample size for characteristic A is 50, for B is 30, and for C is 70. Figure 5(b) illustrates the same example as in Figure 4.

Conclusions

The graphical methods discussed here are efficient in the respect that in one single plot we get visual information simultaneously about the location and spread of the studied characteristic, as well as information about the capability of the process at a given significance level. We can at a glance relate the deviation from target and the spread to each other and to the capability index in such a way that we are able to see whether the non-capability is caused by the fact that the process output is off target, or that the process spread is too large, or if the result is a combination of these two factors. Furthermore, we can easily see how large a change is needed to obtain a capable process. For more details regarding the different plots, see Vännman [2, 3] and Deleryd and Vännman [5].

References

- [1] Vännman, K. (1995). A unified approach to capability indices, *Statistica Sinica* **5**, 805–820.
- [2] Vännman, K. (2001). *A Graphical Method to Control Process Capability*, *Frontiers in Statistical Quality Control*, No 6, H.-J. Lenz & P.-T.H. Wilrich, eds, Physica-Verlag, Heidelberg, pp. 290–311.
- [3] Vännman, K. (2006). Safety regions in process capability plots, *Quality Technology and Quantitative Management* **3**, 227–246.
- [4] Hubele, N. & Vännman, K. (2004). The effect of pooled and un-pooled variance estimators on C_{pm} when using subsamples, *Journal of Quality Technology* **36**, 207–222.
- [5] Deleryd, M. & Vännman, K. (1999). Process capability plots—a quality improvement tool, *Quality and Reliability Engineering International* **15**, 213–227.

KERSTIN VÄNNMAN

Yield Surface Modeling

Introduction

Yield surface modeling™ (YSM) is a multiple-response optimization method developed and patented by Maass and Feldbaumer [1]. YSM is particularly well-suited for use with simulation. If the dependence of a set of critical parameters on a shared set of manufacturing and design factors (some of which experience part-to-part manufacturing variability) can be simulated, YSM can determine settings where the critical parameters are robust against manufacturing variability. The output can be a **response surface** or contour plot, but one in which the response plotted as a function of the factors is either the composite yield or the composite C_{pk} (*see Multi-variate Process Capability Indices, Comparison of; Process Capability Indices, Comparison of; Control Charts and Process Capability*) for multiple responses compared with the specification limits.

YSM™ expands on response surface modeling (RSM), which can explore effects of factors on each of several critical parameters or responses. YSM uses the set of equations provided by RSM to predict the composite yield for the set of critical parameters and enables determination of settings of the factors that will provide acceptable composite yield and composite C_{pk} .

The rationale for using composite yield and composite process capability, C_{pk} , follows this reasoning: each individual critical parameter to be optimized can be in different units (such as volts, seconds, kilograms, or millimeters) and has different specification limits in those units. The degree of “goodness” of projected performance for each critical parameter cannot be directly combined – rather, each critical parameter needs to be converted to an index that is unit-less prior to combination. C_{pk} is unit-less and compares the expected value of the critical parameter to its nearest specification limit, then divides the difference by three times the projected standard deviation for the critical parameter. C_{pk} is a “higher is better” index, and a C_{pk} value of 1.5 is considered consistent with **Six Sigma** capability. Assuming normality, and assuming that the part of the distribution that would fall beyond the nearest specification limit would be rejected, the C_{pk} value is directly related to projected yield for that requirement. Projected yield

for each requirement is another index that is unit-less. The process capabilities for several critical parameters can be combined into a composite C_{pk} , assuming independence. Similarly, projected yields for several critical parameters can be combined into a composite yield index, assuming independence. These resulting composite indices not only allow multiple responses in different units to be combined, but can be predictive indices for the average yield of the product if the product’s performance is tested for this set of critical parameters, with each critical parameter being tested relative to its associated specification limits.

Success Stories

YSM has been used in the development of 60 integrated circuits (ICs) and in a clinical sensor; all products were first-pass successes with higher yield than comparable products. Figure 1 illustrates the results for a family of logic devices (ICs). Prior to using YSM, the first few devices had low yields and the logic family was up for cancellation. After optimization using YSM with Simulation Program with Integrated Circuit Emphasis (SPICE) simulation, the yield for subsequent devices increased from 56% initially to 90–100%, resulting in 57 first-pass successful new products introduced in 28 weeks. This set a record for the company and led to hundreds of millions in sales and over 100 million in profits.

YSM was used to identify and assist in the resolution of yield sensitivity in the design of a new clock driver prior to tapeout (Figure 2). This resulted in a successful clock driver IC family, and over 50 million in sales and 25 million in profits.

Assumptions and Formulas

YSM™ uses these assumptions:

- The polynomial equations obtained through RSM estimate the mean or expected value for each response, $E(Y_i)$, for the combinations of set of factor levels corresponding to that point on the response surface (equation 1).
- The variance for each response at the same combinations of factor levels can be estimated using the Generation of System **Moments** method based on a Taylor Series Expansion, as given in equation (2) [2].

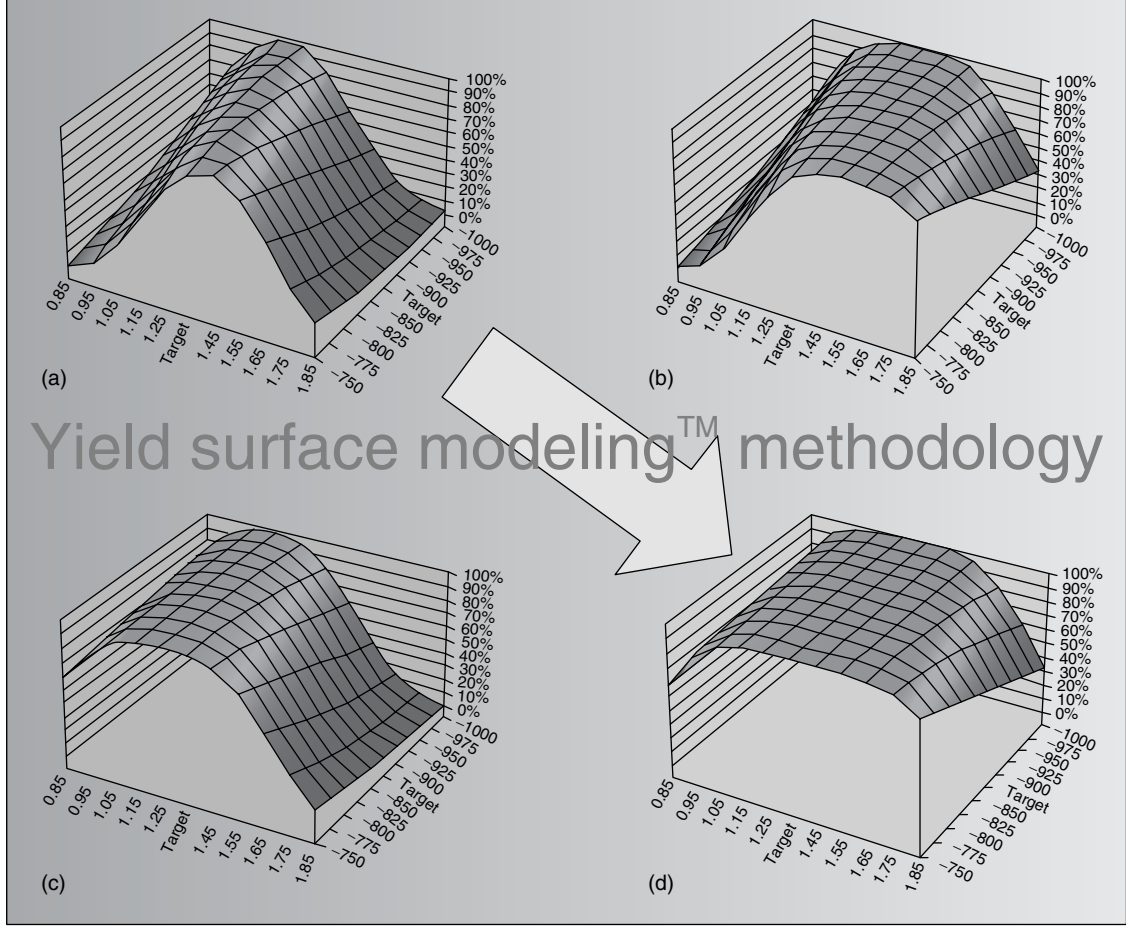


Figure 1 Improvement in yield surface from the original design (a) through a process change (b) and a change in the buffer design (c) that combined to provide a smooth yield surface corresponding to a robust, high yield design (d)

- The factors are uncorrelated. Exceptions to this assumption can be handled with an additional term in the Taylor Series Expansion [3].
- Each response is normally distributed.
- The composite yield for multiple responses can be estimated as the product of the expected percent lying within the specification limits for each response, as given in equations (4) and (6).

The sequence for developing a yield surface use equations (1–6), as shown in Figure 3.

$$E(Y_j) = f(x_1, x_2, \dots, x_n) \quad (1)$$

$$\text{Var}(Y_j) = \sum_{i=1}^N \left(\frac{\partial Y_j}{\partial X_i} \right)^2 \text{Var}(X_i) \quad (2)$$

$$C_{pk}(Y_j) = \min \left\{ \left(\frac{USL - E(Y_j)}{3 \times S(Y_j)} \right), \left(\frac{E(Y_j) - LSL}{3 \times S(Y_j)} \right) \right\} \quad (3)$$

$$\text{Yield}(Y_j) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{3 \times C_{pk}} (Y_j) \exp \left(-\frac{u^2}{2} \right) du \quad (4)$$

$$\text{Composite } C_{pk} = \min(C_{pk}(Y_1), C_{pk}(Y_2), \dots,$$

$$C_{pk}(Y_p)) \quad (5)$$

$$\text{Composite yield} = \prod_{j=1}^P \text{Yield}(Y_j) \quad (6)$$

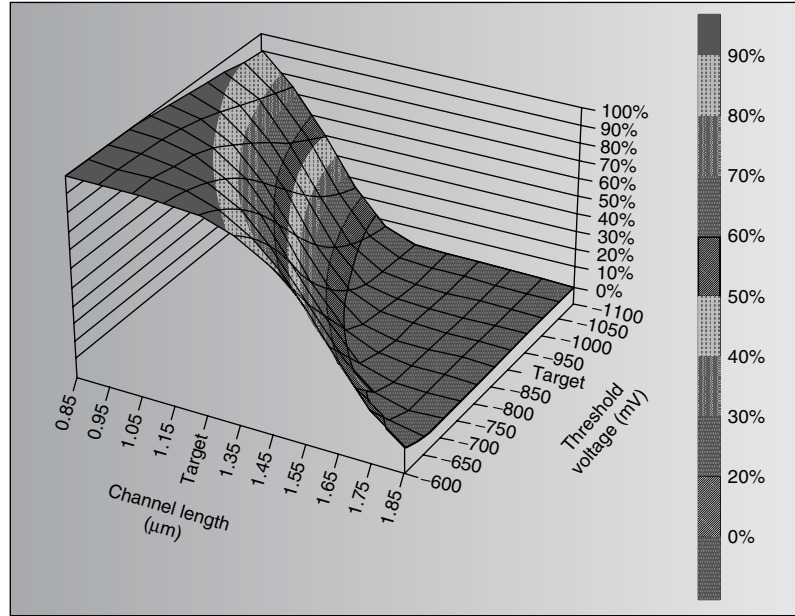


Figure 2 Yield surface model for a clock driver prior to tapeout for first-pass prototyping. The yield surface showed that the first design was highly sensitive to the channel length and would have low yield with the current process

Distributions	Response surface methodology	Mean and variance (propagation of errors)	C_{pk} calculation for each response	Yield for each response	Composite C_{pk}	Composite yield
Input A Input B Input C Input D	Response 1 (or $R1$) = $f1(\text{Inputs})$	Mean1 = $f1(\text{Inputs})$ $(S1)^2 = \text{Sum}[(dR1/dx_j \times Sx_j)^2]$	$C_{pk1} = \frac{(\text{Mean1} - \text{NSL1})}{3 \times S1}$	$Y_1 = \text{cdf}(C_{pk1})$	Composite C_{pk} $= \min(C_{pk1}, \dots, C_{pkn})$	Composite yield $= \text{product}(Y_1, \dots, Y_n)$
Input A Input B Input C Input D	Response 2 (or $R2$) = $f2(\text{Inputs})$	Mean2 = $f2(\text{Inputs})$ $(S2)^2 = \text{Sum}[(dR2/dx_j \times Sx_j)^2]$	$C_{pk2} = \frac{(\text{Mean2} - \text{NSL2})}{3 \times S2}$	$Y_2 = \text{cdf}(C_{pk2})$		
Input A Input B Input C Input D	Response 3 (or $R3$) = $f3(\text{Inputs})$	Mean3 = $f3(\text{Inputs})$ $(S3)^2 = \text{Sum}[(dR3/dx_j \times Sx_j)^2]$	$C_{pk3} = \frac{(\text{Mean3} - \text{NSL3})}{3 \times S3}$	$Y_3 = \text{cdf}(C_{pk3})$		

Figure 3 Method to develop a composite yield surface from RSM for multiple responses

Case Study: Integrated Alternator Regulator (IAR) IC for Automotive

The IAR was a bipolar IC designed by Stephen Dow of Motorola for use in the ignition system for automobiles. YSM™ was performed in conjunction with SPICE, a circuit simulator. A central composite design (CCD) (*see Central Composite Designs*) was set up to vary the levels of four factors that were monitored in the manufacturing area: current gain of bipolar transistors, capacitance of an on-chip capacitor built using a silicon-nitride dielectric, and resistances for two types of resistors.

The results for four critical parameters required for acceptable performance were obtained through SPICE simulation at each combination of levels for these factors. Figure 4 shows the response surfaces for the mean and variance for one critical parameter, soft start time, and the mean and variance combined to develop a response surface for the predicted Z or C_{pk} (referred to as a C_{pk} Surface) followed by the conversion of the C_{pk} surface into a yield surface. This resulting yield surface represents the expected percentage of good parts obtained with that combination of factor levels, subject to the given set of assumptions. Figure 5 illustrates the combination of composite C_{pk} and yield surfaces for the four responses.

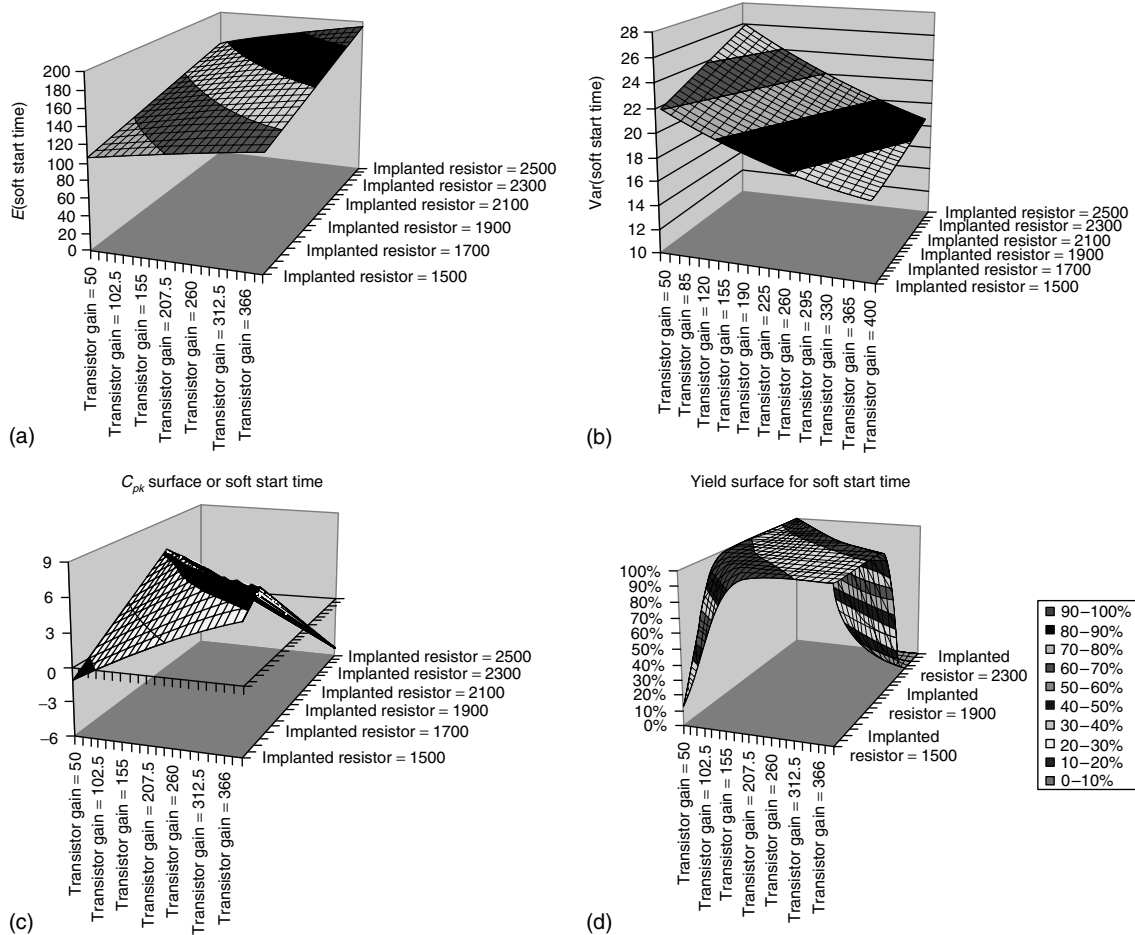


Figure 4 Response surface for the expected value (a) and the variance (b) for the critical parameter soft start time is translated into a response surface for C_{pk} (c) and yield (d)

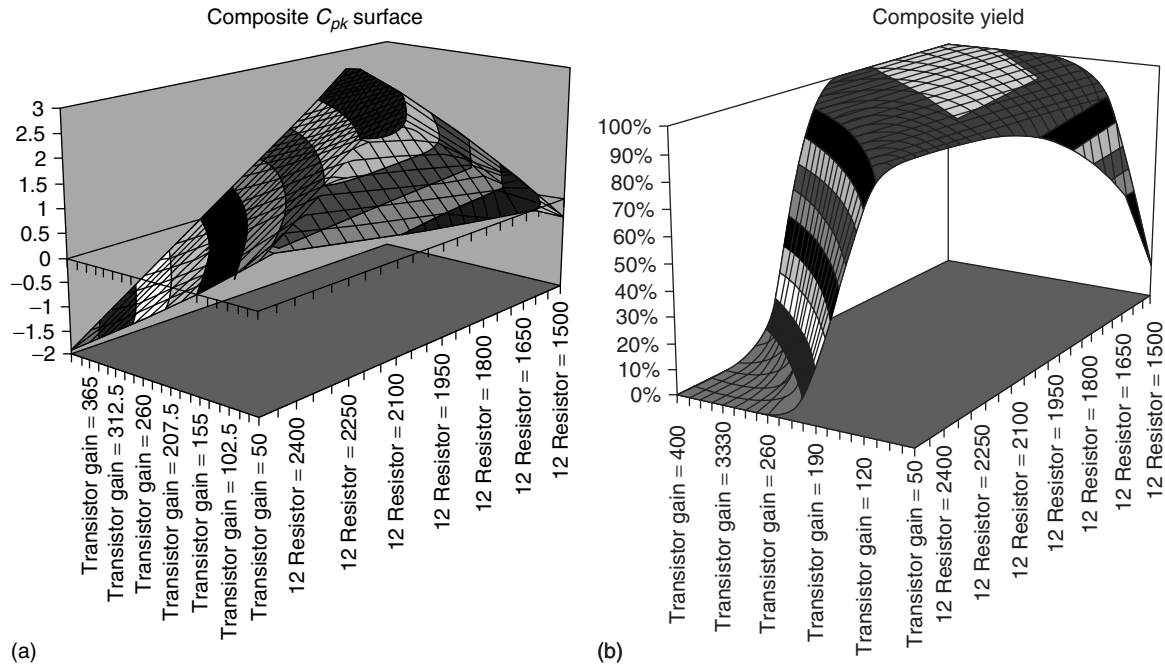


Figure 5 C_{pk} s and yields for all four critical parameters are combined into surfaces for composite C_{pk} (a) and composite yield (b)

YSM Strengths and Comparison with Alternatives

YSM can be used to optimize multiple responses, and predicts the composite yield and composite C_{pk} for those multiple responses. Alternative approaches include overlaid contour plots, **Monte Carlo simulation**, and the use of desirability functions.

An overlay of contour plots can be used to find a **feasible region** where all responses can fall within the specification limits – this only assures that the C_{pk} will be greater than zero and the yield for each response will be greater than 50%. With multiple responses, the composite yield can be substantially below 50% within the feasible region. By contrast, YSM can be used to find a region with very high yield, as illustrated in Figures 1, 2, 4 and 5.

Monte Carlo simulation predicts yield for one response at a time, if the input variables were centered as selected. Stochastic optimization combined with Monte Carlo simulation (as with Crystal Ball™ with OptQuest™) can be used to find settings that optimize the predicted C_{pk} of one response while restricting

the predicted C_{pk} 's for other responses to be more than some minimal value. This approach can provide a stochastic optimum, but YSM provides a means to visualize the sensitivity in the region that Monte Carlo simulation can miss – the latter does not show when the optimum settings are close to a “cliff” where a slight shift can result in disastrously lower yields.

Both Monte Carlo simulation and YSM can provide useful sensitivity analysis. With YSM, the sensitivity analysis comes directly from equation (2), where the magnitude of each contributor to the total variance of each response is proportional to the square of the product of the partial derivative of the response to each x_i times the variance of x_i . This is illustrated with Figure 6, which comes from the automotive IAR case study described earlier. This sensitivity analysis helps to provide guidance as to which factors should be focused on in order to reduce the variance of responses with low predicted C_{pk} and yield.

YSM can also handle correlated input variables by adding more terms from the Taylor Series Expansion (expanded versions of equations 1 and 2).

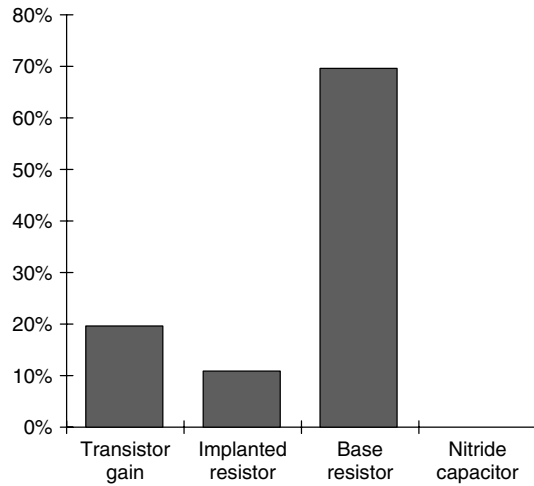


Figure 6 Percent of variance of the response, soft start time, to the variation of each factor for the automotive IAR case study

Summary

YSM has been used in the successful development of several new products. Although most of

the successful uses have been in IC development, this approach has also been used on other products including the optimization of the mechanical design for a two-way radio and the development of a complex medical sensor for DNA sequences, indicating a larger range of potential applications. To date, every new product that has been developed using YSM has been a first-pass success with high yield. This was considered so remarkable by the former president of Motorola Semiconductor Product Sector that, at his request, the approach was treated as proprietary information for several years.

References

- [1] Feldbaumer D. & Maass, E.C. (1995). Yield surface modeling methodology, U.S. Patent 5 438 527, Aug. 4, 1995.
- [2] Hahn, G.J. & Shapiro, S.S. (1967). *Statistical Models in Engineering*, Appendix 7A, John Wiley & Sons.
- [3] Hahn, G.J. & Shapiro, S.S. (1967). *Statistical Models in Engineering*, Appendix 7B, John Wiley & Sons.

ERIC CHARLES MAASS

Process Yield

Process Yield

Process yield has, for a long time, been the most common and standard criterion used in various manufacturing industries to characterize the process performance. Process yield is currently defined as the percentage of the processed product units passing inspection. That is, the product characteristic, X , must fall within the preset specification interval.

For processes involving two-sided manufacturing specifications, the process yield can be calculated as: $Yield = \Pr(LSL < X < USL)$, where USL and LSL are the upper and the lower specification limits, respectively. For the-smaller-the-better processes whose preset specifications require only the upper specification limit USL , the process yield is: $Yield = \Pr(X < USL)$. For the-larger-the-better processes whose preset specifications require only the lower specification limit LSL , the process yield is: $Yield = \Pr(X > LSL)$.

Suppose the production process is operating in a stable manner such that the fraction of the nonconformities (NC), the probability that any unit will not conform to specifications, is $P(\text{NC})$, and that successive units produced are independent. Note that $P(\text{NC}) = 1 - Yield$. If a random sample of n units of a product is selected and D is the number of units of the product that are nonconforming, then D follows $b(n, p)$ with $p = P(\text{NC})$, a binomial distribution with parameters n and p . Generally, the actual fraction of nonconformities, p , is unknown. Traditionally, the parameter p is estimated by the sample fraction nonconforming, the estimated p , which is defined as: $\hat{p} = D/n$. The actual $Yield$ of a process is also unknown, which can be estimated by $1 - \hat{p} = (n - D)/n$.

The above approach is applicable only for processes with large $P(\text{NC})$, say, $p \geq 0.01$. For processes with small p it is inappropriate, since $\hat{p}$ will be most likely zero in most cases.

Process Capability Indices

Process capability indices (PCIs) are intended to provide single-number assessments of the ability to

meet specification limits on quality characteristic(s) of interest. Existing PCIs, C_p , C_a , C_{PU} , C_{PL} , C_{pk} , C_{pm} , and C_{pmk} have been proposed to the manufacturing industry as capability measures based on various criteria including process variation, departure, yield, and loss. The process yield is one of the important factors, though not the only one, considered when accessing or improving a process. In this article, PCIs dedicated to the measurement of the yield of normally distributed processes with two-sided specification limits as well as one-sided specification limit are discussed. These indices are explicitly defined as follows [1–4, **Process Capability Indices, Comparison of**]:

$$\begin{aligned} C_p &= \frac{USL - LSL}{6\sigma}, C_a = 1 - \frac{|\mu - m|}{d}, \\ C_{PU} &= \frac{USL - \mu}{3\sigma}, C_{PL} = \frac{\mu - LSL}{3\sigma}, \\ C_{pk} &= \min \left\{ \frac{USL - \mu}{3\sigma}, \frac{\mu - LSL}{3\sigma} \right\}, \\ C_{pm} &= \frac{USL - LSL}{6\sqrt{\sigma^2 + (\mu - T)^2}}, \\ C_{pmk} &= \min \left\{ \frac{USL - \mu}{3\sqrt{\sigma^2 + (\mu - T)^2}}, \right. \\ &\quad \left. \frac{\mu - LSL}{3\sqrt{\sigma^2 + (\mu - T)^2}} \right\} \end{aligned} \quad (1)$$

where μ is the process mean, σ is the process standard deviation, $m = (USL + LSL)/2$ is the midpoint of the specification interval, $d = (USL - LSL)/2$ is the half-length of the specification tolerance, and T is the target value preset by the product designer or manufacturing engineer. This paper will focus on the situation in which the target value T is set to the specification center, i.e. $T = m$, the most common case in practical applications.

The index C_p , a function of the process standard deviation σ and the specification limits, has been referred to as the *precision index* which is defined to measure the consistency of the process quality characteristic relative to the manufacturing tolerance. The index C_a , a function of the process mean and the specification limits, has been referred to as the *accuracy index*, which is defined to measure the degree of process centering relative to the manufacturing tolerance. Note that $0 \leq C_a \leq 1$ for $LSL \leq \mu \leq USL$, and

2 Process Yield

$C_a < 0$ for $\mu < LSL$ or $\mu > USL$. Especially, $C_a = 1$ for $\mu = m$.

PCIs for Two-Sided Processes

There are many capability indices including C_p , C_a , C_{pk} , C_{pm} , and C_{pmk} designed for assessing processes with two-sided specification limits (which require USL and LSL). The three indices C_{pk} , C_{pm} , and C_{pmk} may be rewritten respectively as follows

$$C_{pk} = C_p C_a, C_{pm} = \frac{C_p}{\sqrt{1 + [3C_p(1 - C_a)]^2}}, \text{ and}$$

$$C_{pmk} = \frac{C_p C_a}{\sqrt{1 + [3C_p(1 - C_a)]^2}} \quad (2)$$

which are all functions of the two basic indices C_p and C_a .

For a normally distributed process, the probability fraction of nonconforming $P(\text{NC})$ is

$$P(\text{NC}) = \Phi[-3C_p(2 - C_a)] + \Phi[-3C_p C_a] \quad (3)$$

which is also a function of the two basic indices C_p and C_a . Especially, $P(\text{NC}) = 2700$ ppm when $C_p = 1$ and $C_a = 1$, i.e. when $USL - LSL = 6\sigma$ and $\mu = m$.

Pearn and Lin [5] obtained a yield measure formula based on C_{pk} as:

$$P(\text{NC}) = \Phi[-3C_{pk}] + \Phi\left[-\frac{3C_{pk}(2 - C_a)}{C_a}\right] \quad (4)$$

for $C_{pk} > 0$ and $C_a > 0$. Based on C_{pk} , Boyles [6] noted a bound for the yield as:

$$2\Phi[3C_{pk}] - 1 \leq \text{Yield} \leq \Phi[3C_{pk}] \quad (5)$$

or

$$\Phi[-3C_{pk}] \leq P(\text{NC}) \leq 2\Phi[-3C_{pk}] \quad (6)$$

for $C_{pk} \geq 0$.

Pearn and Lin [5] also obtained a yield measure formula based on C_{pm} as

$$P(\text{NC}) = \Phi\left[-\frac{2 - C_a}{\sqrt{\frac{1}{(3C_{pm})^2} - (1 - C_a)^2}}\right]$$

$$+ \Phi\left[-\frac{C_a}{\sqrt{\frac{1}{(3C_{pm})^2} - (1 - C_a)^2}}\right] \quad (7)$$

for $(3C_{pm} - 1)/(3C_{pm}) \leq C_a \leq 1$. Based on C_{pm} , Rucinski [7] obtained a lower bound for the yield as: $\text{Yield} \geq 2\Phi(3C_{pm}) - 1$, or $P(\text{NC}) \leq 2\Phi(-3C_{pm})$ for $C_{pm} > \sqrt{3}/3$.

Furthermore, the yield measure formula based on C_{pmk} can be displayed as:

$$P(\text{NC}) = \Phi\left[-\frac{3C_{pmk}(2 - C_a)}{\sqrt{C_a^2 - [3C_{pmk}(1 - C_a)]^2}}\right]$$

$$+ \Phi\left[-\frac{3C_{pmk}C_a}{\sqrt{C_a^2 - [3C_{pmk}(1 - C_a)]^2}}\right] \quad (8)$$

for $3C_{pmk}/(1 + 3C_{pmk}) \leq C_a \leq 1$. The corresponding lower bound for the yield is: $\text{Yield} \geq 2\Phi(3C_{pmk}) - 1$, or $P(\text{NC}) \leq 2\Phi(-3C_{pmk})$ for $C_{pmk} > \sqrt{2}/3$.

Hence, $P(\text{NC}) \leq 2\Phi(-3C)$ and $0 \leq C_a \leq 1$, i.e. $LSL \leq \mu \leq USL$, for $C_{pk} = C$. $P(\text{NC}) \leq 2\Phi(-3C)$ and $(3C - 1)/(3C) \leq C_a \leq 1$, i.e. $m - d/(3C) \leq \mu \leq m + d/(3C)$, for $C_{pm} = C$. And, $P(\text{NC}) \leq 2\Phi(-3C)$ and $3C/(1 + 3C) \leq C_a \leq 1$, i.e. $m - d/(1 + 3C) \leq \mu \leq m + d/(1 + 3C)$, for $C_{pmk} = C$. The three indices C_{pk} , C_{pm} , and C_{pmk} provide the same upper bound $P(\text{NC}) \leq 2\Phi(-3C)$ when $C_{pk} = C$, $C_{pm} = C$, or $C_{pmk} = C$, but different information about the process centering measure C_a . For example, if $C_{pk} = 1.0$ we only have the information of process yield through the upper bound $P(\text{NC}) \leq 2700$ ppm. But, if $C_{pm} = 1.0$ we have the information of process yield through the upper bound $P(\text{NC}) \leq 2700$ ppm and process centering measure $0.667 \leq C_a \leq 1$. And, if $C_{pmk} = 1.0$ we have the information of process yield through the upper bound $P(\text{NC}) \leq 2700$ ppm and process centering measure $0.75 \leq C_a \leq 1$. Table 1 displays the corresponding upper bound of $P(\text{NC})$ in ppm (nonconforming parts per million - NCPPM) for various $C = 0.90(0.01)2.00$ when $C_{pk} = C$, or $C_{pm} = C$, or $C_{pmk} = C$, where $NCCPPM = 10^6 \times P(\text{NC})$. For example, for $C = 1.24$, the number of nonconformities is not greater than 200 ppm.

In general, the process mean, μ , and the process standard deviation, σ , are unknown. But, in practice μ and σ can be estimated by the

Table 1 The corresponding upper bound of $P(\text{NC})$ in ppm (NCPMP) for various $C = 0.90(0.01)2.00$ when $C_{pk} = C$, or $C_{pm} = C$, or $C_{pmk} = C$

C	NCPMP	C	NCPMP	C	NCPMP
0.90	6933.948	1.27	138.967	1.64	0.865
0.91	6333.433	1.28	123.034	1.65	0.742
0.92	5780.136	1.29	108.835	1.66	0.636
0.93	5270.804	1.30	96.193	1.67	0.544
0.94	4802.365	1.31	84.946	1.68	0.466
0.95	4371.923	1.32	74.950	1.69	0.398
0.96	3976.752	1.33	66.073	1.70	0.340
0.97	3614.288	1.34	58.198	1.71	0.290
0.98	3282.122	1.35	51.218	1.72	0.247
0.99	2977.997	1.36	45.036	1.73	0.210
1.00	2699.796	1.37	39.566	1.74	0.179
1.01	2445.537	1.38	34.731	1.75	0.152
1.02	2213.370	1.39	30.460	1.76	0.129
1.03	2001.565	1.40	26.691	1.77	0.110
1.04	1808.510	1.41	23.369	1.78	0.093
1.05	1632.705	1.42	20.443	1.79	0.079
1.06	1472.751	1.43	17.867	1.80	0.067
1.07	1327.350	1.44	15.603	1.81	0.056
1.08	1195.297	1.45	13.614	1.82	0.048
1.09	1075.475	1.46	11.868	1.83	0.040
1.10	966.848	1.47	10.337	1.84	0.034
1.11	868.460	1.48	8.996	1.85	0.029
1.12	779.425	1.49	7.822	1.86	0.024
1.13	698.926	1.50	6.795	1.87	0.020
1.14	626.211	1.51	5.898	1.88	0.017
1.15	560.587	1.52	5.115	1.89	0.014
1.16	501.414	1.53	4.432	1.90	0.012
1.17	448.107	1.54	3.837	1.91	0.010
1.18	400.127	1.55	3.319	1.92	0.008
1.19	356.981	1.56	2.869	1.93	0.007
1.20	318.217	1.57	2.477	1.94	0.006
1.21	283.421	1.58	2.137	1.95	0.005
1.22	252.215	1.59	1.842	1.96	0.004
1.23	224.254	1.60	1.587	1.97	0.003
1.24	199.223	1.61	1.365	1.98	0.003
1.25	176.835	1.62	1.174	1.99	0.002
1.26	156.828	1.63	1.008	2.00	0.002

NCPMP: nonconforming parts per million

sample mean, $\bar{X} = \sum_{i=1}^n X_i/n$, and the sample standard deviation, $S = [\sum_{i=1}^n (X_i - \bar{X})^2/(n-1)]^{1/2}$ or $S_n = [\sum_{i=1}^n (X_i - \bar{X})^2/n]^{1/2}$, to obtain the natural estimators of the three indices $\hat{C}_{pk}$, $\hat{C}_{pm}$, and $\hat{C}_{pmk}$. In order to calculate these estimators, however, sample data must be collected, and a great degree of uncertainty may be introduced into capability assessments owing to sampling errors. The approach of simply looking at the calculated values

of the estimated indices and concluding whether the given process is capable, is unreliable as the sampling errors are ignored. For normally distributed processes, the testing procedures using the natural estimators $\hat{C}_{pk}$, $\hat{C}_{pm}$, or $\hat{C}_{pmk}$ for practitioners have been developed to determine if their processes satisfy the targeted quality condition (see [8–11]).

To obtain an exact measure of the process yield, Boyles [6] considered a yield index, referred to as S_{pk} , for normally distributed processes. The index S_{pk} is defined as:

$$S_{pk} = \frac{1}{3} \Phi^{-1} \left\{ \frac{1}{2} [\Phi(3C_{PU}) + \Phi(3C_{PL})] \right\} \quad (9)$$

There is a one-to-one relationship between S_{pk} and the process yield because for a process with $S_{pk} = c$, we can obtain $Yield = 2\Phi(3c) - 1$ or $P(\text{NC}) = 2\Phi(-3c)$. Table 2 displays the number of nonconformities (in ppm) for various S_{pk} values = 0.5(0.1)2.0.

Unfortunately, the exact distribution of $\hat{S}_{pk}$, the natural estimator of the index S_{pk} , is mathematically intractable. Lee *et al.* [12] derived an approximate distribution of the estimator $\hat{S}_{pk}$ using the Taylor expansion technique. They showed that the estimator $\hat{S}_{pk}$ is approximately normally distributed. Chen and Pearn [13] used the **bootstrap** resampling technique to find the lower confidence bound on S_{pk} , so that practitioners/engineers can use them to perform quality testing and determine whether their processes meet the preset quality requirement.

PCIs for One-Sided Processes

Capability indices C_{PU} and C_{PL} have been designed specifically for assessing processes with one-sided specification limit (which require only *USL* or *LSL*). C_{PU} measures the capability of a smaller-the-better process with an upper specification limit. On the other hand, C_{PL} measures the capability of a larger-the-better process with a lower specification limit.

For a normally distributed process with one-sided specification limit *USL*, corresponding to C_{PU} value the process yield is:

$$\begin{aligned} Yield &= \Pr(X < USL) = \Phi(3C_{PU}), \text{ or} \\ P(\text{NC}) &= \Phi(-3C_{PU}) \end{aligned} \quad (10)$$

On the other hand, for a normally distributed process with one-sided specification limit LSL, corresponding to C_{PL} value the process yield is:

$$\begin{aligned} \text{Yield} &= \Pr(X > LSL) = \Phi(3C_{PL}), \text{ or} \\ P(\text{NC}) &= \Phi(-3C_{PL}) \end{aligned} \quad (11)$$

Given value of C_{PU} or C_{PL} , the corresponding fraction of nonconforming in ppm is

$$\begin{aligned} \text{NCPPM} &= 10^6 \times \Phi(-3C_{PU}) \text{ or} \\ \text{NCPPM} &= 10^6 \times \Phi(-3C_{PL}) \end{aligned} \quad (12)$$

In practice, sample data must be collected in order to calculate these indices. The natural estimators of C_{PU} and C_{PL} are respectively defined as:

$$\hat{C}_{PU} = \frac{USL - \bar{X}}{3S} \text{ and } \hat{C}_{PL} = \frac{\bar{X} - LSL}{3S} \quad (13)$$

For normally distributed processes, Pearn and Chen [14] showed that $\tilde{C}_{PU} = b_f \cdot \hat{C}_{PU}$ and $\tilde{C}_{PL} = b_f \cdot \hat{C}_{PL}$ with $f = n - 1$ and $b_f = \sqrt{\frac{2}{f}} \cdot \frac{\Gamma[f/2]}{\Gamma[(f-1)/2]}$ are the uniformly minimum variance unbiased estimators (UMVUEs) of C_{PU} and C_{PL} , respectively. Based on the UMVUEs, $\tilde{C}_{PU}$, and $\tilde{C}_{PL}$, Lin and Pearn [15] provided testing procedures and efficient SAS computer programs for practitioners to assess processes with one-sided specification limit.

Summary and Related Topics

In this article, PCIs dedicated to the measurement of the yield of normally distributed processes with two-sided specification limits as well as one-sided specification limit are discussed. The yield measure formulas based on C_{pk} , C_{pm} , C_{pmk} , S_{pk} , C_{PU} , and C_{PL} are all given. The corresponding upper bounds of NCPPM for various values of C_{pk} , C_{pm} , and C_{pmk} as well as the corresponding NCPPM for various S_{pk} values are tabulated. For non-normal situation, the technique of non-normal quantile estimation is the most common method in use today for modifying PCIs [16, 17].

References

- [1] Kane, V.E. (1986). Process capability indices, *Journal of Quality Technology* **18**, 41–52.
- [2] Chan, L.K., Cheng, S.W. & Spiring, F.A. (1988). A new measure of process capability: C_{pm} , *Journal of Quality Technology* **20**, 162–175.
- [3] Pearn, W.L., Kotz, S. & Johnson, N.L. (1992). Distributional and inferential properties of process capability indices, *Journal of Quality Technology* **24**, 216–233.
- [4] Pearn, W.L., Lin, G.H. & Chen, K.S. (1998). Distributional and inferential properties of the process accuracy and process precision indices, *Communications in Statistics: Theory and Methods* **27**, 985–1000.
- [5] Pearn, W.L. & Lin, P.C. (2004). Measuring process yield based on capability index C_{pm} , *The International Journal of Advanced Manufacturing Technology* **24**, 503–508.
- [6] Boyles, R.A. (1991). The Taguchi capability index, *Journal of Quality Technology* **23**, 17–26.
- [7] Rucinski, I. (1996). *The relation between C_{pm} and the degree of inclusion*, Ph.D. Thesis, University of Würzburg, Würzburg.
- [8] Cheng, S.W. (1994). Practical implementation of the process capability indices, *Quality Engineering* **7**, 239–259.
- [9] Pearn, W.L. & Lin, P.C. (2002). Computer program for calculating the p-value in testing process capability index C_{pmk} , *Quality and Reliability Engineering International* **18**, 333–342.
- [10] Pearn, W.L. & Lin, P.C. (2004). Testing process performance based on capability index C_{pk} with critical values, *Computers and Industrial Engineering* **47**, 351–369.
- [11] Lin, P.C. & Pearn, W.L. (2005). Testing process performance based on the capability index C_{pm} , *The International Journal of Advanced Manufacturing Technology* **27**, 351–358.
- [12] Lee, J.C., Hung, H.N., Pearn, W.L. & Kueng, T.L. (2002). On the distribution of the estimated process yield index S_{pk} , *Quality and Reliability Engineering International* **18**, 111–116.
- [13] Chen, J.P. & Pearn, W.L. (2002). Testing process performance based on the yield: an application to the liquid-crystal display module, *Microelectronics Reliability* **42**, 1235–1241.
- [14] Pearn, W.L. & Chen, K.S. (2002). One-sided capability indices C_{pu} and C_{pl} : decision making with sample information, *International Journal of Quality and Reliability Management* **19**, 221–245.
- [15] Lin, P.C. & Pearn, W.L. (2002). Testing process capability for one-sided specification limit with application to the voltage level translator, *Microelectronics Reliability* **42**, 1975–1983.
- [16] Pearn, W.L. & Chen, K.S. (1997). Capability indices for non-normal distributions with an application in electrolytic capacitor manufacturing, *Microelectronics Reliability* **37**, 1853–1858.
- [17] Kotz, S. & Lovelace, C. (1998). *Process Capability Indices in Theory and Practice*, Arnold, New York.

Control Charts and Process Capability

Introduction

Statistical process control (SPC) charts play a very important role in the determination of process capability. Process capability index measurements are meaningless when calculated from an unstable, out-of-control process. The very definition of “process capability” hints at long-term process quality, and given that process data taken from an out-of-control process in no way indicates what the process is capable of producing in terms of quality product, no valuable predictive information can be gleaned from this type of data. Therefore, it is absolutely imperative that statistical control be achieved *via* the use of appropriate **control charts** before process capability indices are calculated. Process capability indices and control charts should be considered an inseparable pair of **benchmarking** tools for the quality conscious enterprise.

Although many indices have been introduced in the literature to represent process capability from various perspectives (see Kotz and Lovelace [1]), the original C_p and C_{pk} indices remain the dominant choice among quality professionals, along with the second generation index C_{pm} , which contains an additional component for quantifying deviations from target. These indices are defined as

$$C_p = \frac{USL - LSL}{6\sigma} = \frac{d}{3\sigma} \quad (1)$$

$$C_{pk} = \min\left(\frac{USL - \mu}{3\sigma}, \frac{\mu - LSL}{3\sigma}\right) = \frac{d - |\mu - M|}{3\sigma} \quad (2)$$

and

$$C_{pm} = \frac{USL - LSL}{6\tau} = \frac{USL - LSL}{6\sqrt{\sigma^2 + (\mu - T)^2}} \quad (3)$$

where $USL - LSL$ is the allowable tolerance range of the process, $d = (USL - LSL)/2$ is the half-interval width, $M = (USL + LSL)/2$ is the midpoint of the specification interval, and τ is the standard deviation about the target.

A significant portion of the theoretical research dedicated to process capability indices has focused on

those index estimators that utilize a pooled estimator of process variance of the form

$$s = \frac{\sum_{i=1}^n (X_i - \bar{X})^2}{n - 1} \quad (4)$$

Estimators incorporating this pooled estimator are preferred by statisticians because of their superior and more tractable statistical properties, as compared to those containing an unpooled variance estimator (which utilizes subsample data from a control chart). (More details concerning the sampling distributions, **moments**, **confidence intervals**, etc. for the indices using the **pooled variance** estimator may be found in Kane [2], Kocherlakota [3], Kotz and Johnson [4], Boyles [5], Kotz and Johnson [6], Chou and Owen [7], Nagata and Nagahata [8], Chou *et al.* [9], Bissell [10], Pearn *et al.* [11], Chan *et al.* [12], Kushler and Hurley [13] and Johnson [14], among others.)

The use of index estimators containing s , however, does not promote or encourage the use of control charts to check for process stability, and requires a second sample over and above the data collected on these charts. Therefore, data collected in subsamples from an in-control SPC chart, most likely an $\bar{X} - R$ or $\bar{X} - s$ chart, is preferred in practice for index estimation. Juran and Gryna [15] termed this the *control chart method*, and it is the only technique for process capability estimation that provides a means to check for statistical control. In this case, σ is estimated from the in-control chart for the process by either

$$\hat{\sigma}_R = \frac{\bar{R}}{d_2} \quad (5)$$

or

$$\hat{\sigma}_s = \frac{\bar{s}}{c_4} \quad (6)$$

where $\bar{R}$ is the average subsample range, $\bar{s}$ is the average subsample standard deviation, and d_2 and c_4 are coefficients which are tabulated in numerous statistical **quality control** texts (see, e.g., Montgomery [16]). These estimators relate the process within-subgroup range (or standard deviation) to the sample standard deviation for normal processes.

This article will focus on the significant work published to date on process capability indices estimated using subsamples, as is the procedure when

2 Control Charts and Process Capability

using control chart data for index estimation, and the determination of optimal control chart parameters for quality programs that use capability indices as the primary measure of process status. Only by a thorough understanding of the statistical characteristics of the subsample index estimators can we effectively use control charts and process capability indices together for process benchmarking and improvement.

Control Chart Performance and Capability Indices

The first definitive examination of control chart performance and its effect on capability estimation was performed by Porter and Oakland [17], who studied the effect of process shifts on control chart performance for a given C_{pk} and subsample size (n). They considered the general case of a process mean shift of $k\sigma$, such that the upper 3σ tail of the distribution just touches the upper specification limit (USL) and 99.73% of the sample means are below the upper control limit. In this case, $k = 6/\sqrt{n}$ and C_{pk} can be redefined in terms of k

and n :

$$C_{pk} = \frac{USL - \mu}{3\sigma} = \frac{k\sigma + \sigma}{3\sigma} \\ = 1 + \frac{k}{3} = 1 + \frac{2}{\sqrt{n}} \quad (7)$$

From this definition, the probability of getting a sample mean outside the control limits on the $\bar{X}$ chart from the first sample after the shift, P_1 , can be calculated for a given C_{pk} and **sample size** (n). Table 1 (from Porter and Oakland [17]), gives the probability (P_1) for various values of C_{pk} and n . This table illustrates that a high C_{pk} value has a significant effect upon the probability of detecting a process shift.

Table 1 illustrates the significant threshold of $C_{pk} = 2$; at or above this capability level and for subgroup sizes of 2 or more, the probability of finding a sample mean above the control limit after a process shift is higher than 0.89, with a resulting **average run length** (ARL) of approximately 1. Therefore, it is very likely that a shift in the process mean of size $6\sigma/\sqrt{n}$ will be detected on the first sample taken after the shift. This is also shown graphically, for $n = 2$ to 8, in Figure 1.

Table 1 Probability of detecting a point outside the control limits when the first sample is taken after a process shift to the USL^(a)

C_{pk}	N									
	1	2	3	4	5	6	7	8	9	10
1.0	0.0013	0.0013	0.0013	0.0013	0.0013	0.0013	0.0013	0.0013	0.0013	0.0013
1.1	0.0035	0.0049	0.0066	0.0082	0.0099	0.0119	0.0136	0.0158	0.0179	0.0202
1.2	0.0082	0.0158	0.0250	0.0359	0.0485	0.0630	0.0793	0.0968	0.1151	0.1357
1.3	0.0179	0.0418	0.0749	0.1151	0.1611	0.2148	0.2676	0.3264	0.3821	0.4404
1.4	0.0359	0.0968	0.1788	0.2743	0.3745	0.4761	0.5675	0.6517	0.7267	0.7852
1.5	0.0668	0.1894	0.3446	0.5000	0.6368	0.7486	0.8340	0.8925	0.9332	0.9591
1.6	0.1151	0.3264	0.5478	0.7257	0.8461	0.9207	0.9608	0.9817	0.9918	0.9964
1.7	0.1841	0.4860	0.7389	0.8849	0.9545	0.9838	0.9948	0.9984	*	*
1.8	0.2743	0.6517	0.8770	0.9641	0.9911	0.9980	*	*		
1.9	0.3821	0.7939	0.9525	0.9918	*	*				
2.0	0.5000	0.8925	0.9857	0.9987						
2.1	0.6179	0.9525	0.9967	*						
2.2	0.7257	0.9817	*							
2.3	0.8159	0.9941								
2.4	0.8849	0.9984								
2.5	0.9332	*								
2.6	0.9641									

^(a) © John Wiley & Sons Ltd, 1990

* >0.9987

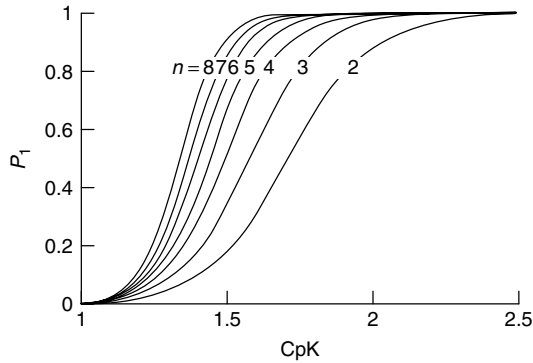


Figure 1 The graph of P_1 against C_{pk} for $n = 2$ to 8 [© John Wiley & Sons Ltd, 1990.]

Table 2 (also from Porter and Oakland [17]) contains the ARL to an action point (ARL(A)) from the P_1 probabilities for various values of C_{pk} and n . Following the pattern seen in Table 1, large C_{pk} values allow for relatively small required subsample sizes in order to achieve desirably low ARLs. The poorer the process capability, the larger the subsample size must be in order to keep the ARL acceptably low.

Of practical interest is the required value of C_{pk} for a given value of n , such that $P_1 \geq 0.9987$. These

values are provided in Table 3. It can be seen that for the smaller subgroup sizes, C_{pk} must be greater than 2 in order to achieve $P_1 \geq 0.9987$. As C_{pk} decreases, larger subgroup sizes are needed in order to increase the probability of process shift detection.

Control charts may also be optimized for the specific purpose of monitoring process quality using capability indices. More and more often, organizations are using capability indices themselves to define current process status. They are also utilized as a bridge between current process performance as monitored by SPC charts and the product design specifications. Therefore, control chart parameters may be optimized to control to a specific minimum index requirement. Wu *et al.* [18] developed an algorithm for designing the $\bar{X} - s$ chart that aligns the in-control and out-of-control parameters of the charts with a required value of C_{pk} . Specifically, the charts are designed so that when C_{pk} declines below a minimum threshold, the $\bar{X} - s$ chart will give a signal within a short period of time.

As originally shown by Kane [2], a single C_{pk} value will correspond to an infinite number of (μ, σ) pairs for the process distribution, therefore corresponding to an infinite number of possible shifts in the mean (δ_μ) or standard deviation (δ_σ) of the process. Also, since C_{pk} compares the allowable design

Table 2 ARL(A) – average run lengths to an action point after a process shift to the USL^(a)

C_{pk}	N									
	1	2	3	4	5	6	7	8	9	10
1.0	740.7	740.4	740.7	740.7	740.7	740.7	740.7	740.7	740.7	740.7
1.1	285.7	204.1	151.5	121.9	101.0	84.0	73.5	63.3	56.2	49.5
1.2	121.9	63.3	40.0	27.9	20.6	15.9	12.6	10.3	5.7	7.4
1.3	55.9	23.9	13.4	8.7	6.2	4.7	3.7	3.1	2.6	2.3
1.4	27.9	10.3	5.6	3.6	2.7	2.1	1.8	1.5	1.4	1.3
1.5	15.0	5.3	2.9	2.0	1.6	1.3	1.2	1.1	1.1	1
1.6	8.7	3.1	1.8	1.4	1.2	1.1	1	1	1	1
1.7	5.4	2.1	1.4	1.1	1	1	1	1	*	*
1.8	3.6	1.5	1.1	1	1	1	*	*		
1.9	2.6	1.3	1	1	*	*				
2.0	2.0	1.1	1	*						
2.1	1.6	1.1	1							
2.2	1.4	1	*							
2.3	1.2	1								
2.4	1.1	1								
2.5	1.1	1								
2.6	1	*								

^(a) © John Wiley & Sons Ltd, 1990

Value of the * = 1

4 Control Charts and Process Capability

specification range to the center and spread of the process, it can serve as a link between the control charts and the process specifications. When the process shifts by δ_μ , as indicated on the $\bar{X}$ chart, or by δ_σ , as indicated on the s chart, C_{pk} will decrease. If C_{pk} decreases to some critical level, say $C_{pk,d}$, the defective rate will increase to an undesired level, and the process should be stopped. An optimally designed $\bar{X} - s$ chart will have a large in-control ARL, ARL_0 , so that the false alarm rate is very small, as well as a small out-of-control ARL, ARL_d , so that out-of-control conditions are detected quickly. The Wu *et al.* [18] modification to these charts is based upon the following optimization model:

Minimize : n

Subject to : $arl_0 \geq ARL_0$ (at $C_{pk} = c_{pk,0}$) (8)

$arl_d \leq ARL_d$ (at $C_{pk} = c_{pk,d}$) (9)

where arl_0 and arl_d are the actual in-control and out-of-control ARLs, respectively. The optimal design of the $\bar{X} - s$ chart is dependent upon the maximum allowable probability of type I error, α , and is composed of two elements, $\alpha_{\bar{X}}$ and α_s (the type I error probabilities for the $\bar{X}$ chart and s chart, respectively). For a specific shift pattern $(\delta_\mu, \delta_\sigma)$, the out of control $C_{pk,d}$ may be calculated as

$$C_{pk,d} = \frac{USL - (\mu_0 + \delta_\mu \sigma_0)}{3\delta_0 \sigma_0} \quad (10)$$

Note that each shift pattern is actually determined by δ_μ , since δ_σ is dependent upon δ_μ for a given

Table 3 Required C_{pk} for a given sample size, n , and process mean shift, $k\sigma$, to satisfy the requirements $P_1 \geq 0.9987^{(a)}$

Subgroup size, n	Process mean shift in standard deviations k	Process capability C_{pk}
1	6.00	3.00
2	4.24	2.41
3	3.46	2.15
4	3.00	2.00
5	2.68	1.89
6	2.45	1.82
7	2.27	1.76
8	2.12	1.71

^(a) © John Wiley & Sons Ltd, 1990

value of $C_{pk,d}$. The mean shift δ_μ can vary from 0 to $\delta_{\mu,\max}$, and $\delta_{\mu,\max}$ can be determined as

$$\delta_{\mu,\max} = \frac{USL - \mu_0 - 3c_{pk,d}\sigma_0}{\sigma_0} \quad (11)$$

by setting δ_σ to its minimum value of 1.

There are an infinite number of control limit patterns $(LCL_{\bar{X}}(\alpha_{\bar{X}}), UCL_{\bar{X}}(\alpha_{\bar{X}}), UCL_s(\alpha_s))$ for a given sample size n , and a particular sample size is considered feasible per the optimization model given above if the corresponding $arl_d(\alpha_{\bar{X}}^*) \leq ARL_d$, where $\alpha_{\bar{X}}^*$ is the optimal value of $\alpha_{\bar{X}}$ associated with the shortest $arl_d(\alpha_{\bar{X}}^*)$ among all $arl_d(\alpha_{\bar{X}})$.

The Wu *et al.* [18] procedure for checking the feasibility of a sample size n is as follows:

1. Calculate the required value of α as

$$\alpha = \frac{1}{ARL_0} \quad (12)$$

2. Calculate $\delta_{\mu,\max}$ by equation (11).
3. Determine the optimal $\alpha_{\bar{X}}^*$ within the range $0 < \alpha_{\bar{X}} < \alpha$. To accomplish this, for each given value of $\alpha_{\bar{X}}$ determine:
 - (a) The corresponding α_s by

$$\alpha_s = \frac{\alpha - \alpha_{\bar{X}}}{1 - \alpha_{\bar{X}}} \quad (13)$$

- (b) The associated control limits using the equations

$$LCL_{\bar{X}} = \mu_0 - z \times \frac{\sigma_0}{\sqrt{n}} \quad (14)$$

$$UCL_{\bar{X}} = \mu_0 + z \times \frac{\sigma_0}{\sqrt{n}} \quad (15)$$

and

$$UCL_s = \sigma \sqrt{\frac{\chi_{\alpha_s, n-1}^2}{n-1}}, LCL_s = 0 \quad (16)$$

- (c) The out-of-control ARL $arl_d(\alpha_{\bar{X}})$ using

$$arl_d(\alpha_{\bar{X}}) = arl_d(LCL_{\bar{X}}(\alpha_{\bar{X}}), UCL_{\bar{X}}(\alpha_{\bar{X}}),$$

$$UCL_s(\alpha_s))$$

$$= \int_0^{\delta_{\mu,\max}} arl'_d(\delta_\mu, \delta_\sigma) f(\delta_\mu) d\delta_\mu$$

$$= \int_0^{\delta_{\mu,\max}} arl'_d(\delta_\mu, \delta_\sigma(\delta_\mu)) f(\delta_\mu) d\delta_\mu \quad (17)$$

where $f(d_\mu)$ is the **probability density function** of δ_μ .

4. If $arl_d(\alpha_{\bar{X}}^*) \leq ARL_d$, the sample size is feasible.

Step 3 of the above procedure sets the optimal allocation of the detection **power** between the $\bar{X}$ and s charts that results in the minimum out of control arl_d .

Caution is urged when considering this process control approach since process control charts should *control* to the inherent variation range of the process and not to process design specifications. The standard rule of thumb in process control is: "Don't put process specifications on a control chart." By controlling to a process capability index value, which contains the process specifications, you are controlling to the specs.

Confidence Intervals of Process Capability Indices for Subgrouped Data

The waters of statistical understanding are still a bit murky when it comes to capability indices and control chart data. For example, there are currently no confidence intervals other than tables developed *via* bootstrapping for estimators containing $\hat{\sigma}_R$. In this case, Li *et al.* [19] developed lower confidence limits for C_p using $\hat{\sigma}_R$. Noting that a $100(1 - \alpha)\%$ lower confidence limit on C_p , c_o , satisfies $\Pr(C_p \geq c_o) = \alpha$, then

$$\Pr(C_p \geq c_o) = \Pr\left\{\frac{C_p}{\hat{C}_p} \geq \frac{c_o}{\hat{C}_p}\right\} = \Pr\left\{\frac{R}{\sigma d_2} \geq \frac{c_o}{\hat{C}_p}\right\} = \alpha \quad (18)$$

or

$$\begin{aligned} \Pr\left\{\frac{\bar{R}}{\sigma d_2} \leq \frac{c_o}{\hat{C}_p}\right\} &= \Pr\left\{R' \leq \frac{d_2 c_o}{\hat{C}_p}\right\} \\ &= \Pr\left\{\frac{\sqrt{v} R'}{c} \leq \frac{\sqrt{v} d_2 c_o}{c \hat{C}_p}\right\} \\ &\cong \Pr\left\{\chi_v \leq \frac{\sqrt{v} d_2 c_o}{c \hat{C}_p}\right\} = 1 - \alpha \end{aligned} \quad (19)$$

where $R' = R/\sigma$. So,

$$\frac{\sqrt{v} d_2 c_o}{c \hat{C}_p} = \chi_{v, 1-\alpha} \quad (20)$$

so that

$$c_o = \frac{c \hat{C}_p}{\sqrt{v} d_2} \sqrt{\chi_{v, 1-\alpha}^2} \quad (21)$$

Table 4 (from Li *et al.* [19]) gives the ratio values of $c_o/\hat{C}_p$ for $m = 1, 5, 10, 20, 30$ and $n = 3$ to 10, where data is collect in m subgroups of size n each.

Hypothesis Tests Based on m Subgroups of Size n Each

Control chart (subsample) data may also be used to test hypotheses concerning process capability. Typically, hypothesis tests of the general form

$$\begin{aligned} H_0: & \text{The process is not capable} \\ H_1: & \text{The process is capable} \end{aligned}$$

are of interest when benchmarking a process's ability to meet product specifications. Using a standard **hypothesis testing** procedure, the degree of difference between the estimated capability index value and the true or hypothesized capability (as defined by a minimum capability requirement) is determined. To accomplish this, an index value is estimated, compared to a lower index bound c_o , then the so-called p value is computed, which refers to the actual risk of incorrectly concluding that the process is capable for a particular test. For the C_p index, the hypotheses are

$$\begin{aligned} H_0: & C_p \leq c_o \\ H_1: & C_p > c_o \end{aligned} \quad (22)$$

where c_o is the standard minimal criteria of C_p for a particular process. The hypotheses for other indices follow the same format.

Table 4 Maximum values of the ratio $c_o/\hat{C}_p$

Subgroup size, n	Number of subgroups (m)				
	1	5	10	20	30
3	0.255	0.631	0.735	0.811	0.845
4	0.369	0.697	0.783	0.845	0.873
5	0.443	0.735	0.811	0.865	0.890
6	0.494	0.760	0.829	0.879	0.901
7	0.533	0.779	0.843	0.888	0.908
8	0.562	0.793	0.853	0.895	0.914
9	0.586	0.804	0.861	0.901	0.919
10	0.605	0.813	0.867	0.906	0.923

C_p

While Kane [2] provided the critical hypothesis testing protocol for C_p when a single sample is used, Kirmani *et al.* [20] developed the critical value for the hypothesis test using subgroups. They used the estimator

$$\hat{C}_p^* = \frac{(USL - LSL)d_p}{6} \quad (23)$$

where

$$d_p = \sqrt{\frac{m(n-1)-1}{m(n-1)}} \times \frac{\varepsilon_{m(n-1)-1}}{s_p} \quad (24)$$

with constant

$$\varepsilon_k = \sqrt{\frac{2}{k}} \times \left[\frac{\Gamma(k+1)/2}{\Gamma(k/2)} \right] \quad (25)$$

C_{pk}

Lin and Sheen [21] developed a hypothesis testing regimen for C_{pk} for the case where process measurements are collected in subgroups from statistical control charts, and assuming the data were collected as a sequence of independent samples. They applied Hamaker's approximation so that the testing procedure can be adequately derived from a **normal distribution**, avoiding the more complicated sampling distributions of $\hat{C}_{pk(R)}$ and $\hat{C}_{pk(s)}$. The procedure provides p values for a given estimate of C_{pk} , which can then be used to evaluate the hypotheses. Subgrouped quality control data is typically collected on an $\bar{X} - R$ or $\bar{X} - s$ control chart pair, where the process standard deviation, σ , is estimated by either $\bar{R}/d_2$ or $\bar{s}/c_4$ as in equations (1) and (2). The coefficients d_2 and c_4 are added to eliminate **bias** within these estimators. $\hat{C}_{pk(R)}$ and $\hat{C}_{pk(s)}$ are the estimators of C_{pk} from the two pairs of charts, respectively.

When the process standard deviation is known, $3\sqrt{mn}\hat{C}_{pk(s)}$ is distributed as $N(3\sqrt{mn}C_{pk}, 1)$. Defining the observed value of $\hat{C}_{pk(R)} = W_R$, the p value can be computed as

$$p \text{ value} = \sup_{C_{pk} \leq C} \Pr\{\hat{C}_{pk(R)} \geq W_R | m, n\} \\ = \Pr\{Z \geq \theta_R\} \quad (26)$$

where

$$\theta_R = \left[\frac{3\sqrt{mn}(W_R - C)}{\sqrt{1 + 9nd_3^2 W_R^2 / d_2^2}} \right] \quad (27)$$

In similar fashion for the $\bar{X} - s$ chart estimate of C_{pk} , define $\hat{C}_{pk(s)} = W_s$. From this, the p value can be computed as

$$p \text{ value} = \sup_{C_{pk} \leq C} \Pr\{\hat{C}_{pk(s)} \geq W_s | m, n\} \\ = \Pr\{Z \geq \theta_s\} \quad (28)$$

where

$$\theta_s = \left[\frac{3\sqrt{mn}(W_s - C)}{\sqrt{1 + 9nc_5^2 W_s^2 / c_4^2}} \right] \quad (29)$$

C_{pm}

Hypothesis testing for C_{pm} is a bit tricky, and Carlsson [22] feels that tests of hypothesis concerning deviation from the target and tests concerning variation about the target cannot be tested separately. However, Chan *et al.* [12] developed a hypothesis testing procedure for their C_{pm} estimator, which is based on a single sample. The power of this procedure has not been tested. To date, no technique has been presented for hypothesis testing on C_{pm} with subsample data.

Conclusion

Process capability indices are used to measure the long-term ability of a process to manufacture product within specification. They must be estimated from stable (i.e., in-control) process data; otherwise, their values have no long-term implications for process performance. This stability must be assured *via* a statistical control chart, thus in practical terms control charts and process capability indices should be used as a pair. It is hoped that research will continue into the statistical characteristics, effects of subsample size, and practical implementation issues of process capability indices estimated from subsample, control chart data.

References

- [1] Kotz, S. & Lovelace, C. (1998). *Process Capability Indices in Theory and Practice*, Arnold, London.
- [2] Kane, V.E. (1986). Process capability indices, *Journal of Quality Technology* **18**(1), 41–52.
- [3] Kocherlakota, S. (1992). Process capability index: recent developments, *Sankhya: The Indian Journal of Statistics* **54B**, (pt 3), 352–369.

-
- [4] Kotz, S. & Johnson, N.L. (1993). *Process Capability Indices*, Chapman & Hall, London.
 - [5] Boyles, R.A. (1991). The Taguchi capability index, *Journal of Quality Technology* **23**(1), 17–26.
 - [6] Kotz, S. & Johnson, N.L. (1995). Comment on: percentages of units within specification limits associated with given values of C_{pk} and C_{pm} (by Chan LK and Tang JH), *Communications in Statistics* **25**.
 - [7] Chou, Y.-M. & Owen, D.B. (1989). On the distributions of the estimated process capability indices, *Communications in Statistics: Theory and Methods* **18**(12), 4549–4560.
 - [8] Nagata, Y. & Nagahata, H. (1994). Approximation formulas for the lower confidence limits of process capability indices, *Okayama Economic Review* **25**(4), 301–314.
 - [9] Chou, Y.-M., Owen, D.B. & Borrego, S.A. (1990). Lower confidence limits on process capability indices, *Journal of Quality Technology* **22**(3), 223–229.
 - [10] Bissell, A.F. (1990). How reliable is your capability index? *Applied Statistics* **39**(3), 331–340.
 - [11] Pearn, W.L., Kotz, S. & Johnson, N.L. (1992). Distributional and inferential properties of process capability indices, *Journal of Quality Technology* **24**, 216–231.
 - [12] Chan, L.K., Cheng, S.W. & Spiring, F.A. (1988). The robustness of the process capability index to departures from normality, in *Statistical Theory and Data Analysis II, Proceedings of Second Pacific Area Statistical Conference*, Tokyo, pp. 223–239.
 - [13] Kushler, R. & Hurley, P. (1992). Confidence bounds for capability indices. *Journal of Quality Technology* **24**, 188–195.
 - [14] Johnson, T. (1992). The relationship of Cpm to squared error loss. *Journal of Quality Technology* **24**(4), 211–215.
 - [15] Juran, J.M. & Gryna, F.M. (1988). *Quality Control Handbook*, 4th Edition, McGraw-Hill, New York.
 - [16] Montgomery, D.C. (1996). *Introduction to Statistical Quality Control*, 3rd Edition, John Wiley & Sons, New York.
 - [17] Porter, L.J. & Oakland, J.S. (1990). Measuring process capability using indices – some new considerations, *Quality and Reliability Engineering International* **6**, 19–27.
 - [18] Wu, Z., Xie, M. & Tian, Y. (2002). Optimization design of the charts for monitoring process capability, *Journal of Manufacturing Systems* **21**(2), 83–92.
 - [19] Li, H., Owen, D.B. & Borrego, S.A. (1990). Lower confidence limits on process capability indices based on the range, *Communications in Statistics: Simulation* **19**(1), 1–24.
 - [20] Kirmani, S.N.U.A., Kocherlakota, K. & Kocherlakota, S. (1991). Estimation of σ and the process capability index based on subsamples, *Communications in Statistics: Theory Method* **20**, 275–291.
 - [21] Lin, H.-Ch. & Sheen, G.-J. (2005). Practical implementation of the capability index based on the control chart data, *Quality Engineering* **17**, 371–390.
 - [22] Carlsson, O. (1995). *On the use of the Cpm index for detecting instability of a production process*, University of Orebro, Sweden.

Further Reading

- Duncan, A.J. (1994). *Quality Control and Industrial Statistics*, Richard Irwin, Homewood.
- Grant, E.L. & Leavenworth, R.S. (1988). *Statistical Quality Control*, McGraw-Hill, New York.
- Ott, E.R. & Schilling, E.G. (1990). *Process Quality Control: Troubleshooting and Interpretation of Data*, McGraw-Hill, New York.
- Pyzdek, T. (1989). *What Every Engineer Should Know About Quality*, Marcel Dekker, New York.

CYNTHIA R. LOVELACE

Process Capability Indices for Multiple Stream Processes

Introduction

Separate process streams are created in manufacturing when the flow of parts is divided into several presumably identical paths. This split in production occurs when an operation has different molds, spindles, lines, stations, lines, fixtures, or pallets. Presumably, the quality of parts produced by each stream is identical but, in practice, there are often substantial differences between streams. The capability index introduced in this article may be applied when a measure of capability is desired for the combined output of all streams.

The C_{pk} Index

As in all capability studies, the first prerequisite is to verify that each process stream is in a reasonably good state of statistical control (*see Control Charts and Process Capability*). The output from each stream may have a different average and/or standard deviation, but every stream must have a stable, predictable output.

For example, an injection-molding operation utilizes a four-cavity die to produce a plastic part. Through the use of **control charts**, each cavity demonstrates good stability for the inner diameter (i.d.) of the part being molded, but the output of hole size varies considerably from cavity to cavity, as is illustrated in Figure 1.

After the parts are molded, they fall down a chute into a common container, which is then shipped to a customer. What should the customer who receives this shipment – which contains a mixture of parts from all four cavities – expect in terms of parts meeting the specification requirements for hole size?

One popular metric for assessing capability is called the C_{pk} index [1], which is defined in equations (1–3). Here, μ is the process mean, σ is the process standard deviation, LSL is the lower specification limit for the feature being studied, and USL is its upper specification limit (*see Control Charts for*

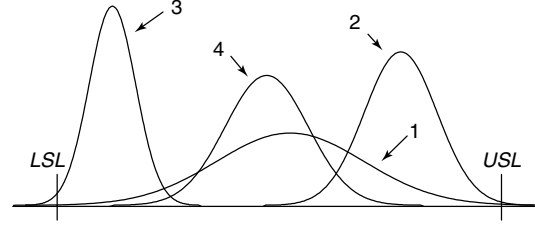


Figure 1 Distributions of hole size for each of four cavities (labeled 1–4)

the Mean).

$$C_{pk} = \text{minimum}(C_{pl}, C_{pu}) \quad (1)$$

$$C_{pl} = \frac{\mu - LSL}{3\sigma} \quad (2)$$

$$C_{pu} = \frac{USL - \mu}{3\sigma} \quad (3)$$

If only an LSL is given for the feature being studied, then C_{pk} equals C_{pl} . Conversely, when just a USL is provided, C_{pk} equals C_{pu} .

When the process output is close to a **normal distribution**, the percentage of parts below the LSL, p_{LSL} , is determined with equation (4), where $\Phi(\cdot)$ is the cumulative distribution function of the standard normal distribution (with $\mu = 0$ and $\sigma = 1$) [2]. For example, $\Phi(0)$ equals 0.50000 (50%), $\Phi(2)$ equals 0.97725, and $\Phi(3)$ equals 0.99865. In the spreadsheet program Excel, $\Phi(2)$ corresponds to the statistical function NORMSDIST(2), which returns the value of 0.97725.

$$p_{LSL} = 1 - \Phi(3C_{pl}) \quad (4)$$

Likewise, the percentage of parts above the USL, p_{USL} , is found with equation (5).

$$p_{USL} = 1 - \Phi(3C_{pu}) \quad (5)$$

The total percentage of nonconforming parts, p_T , is then computed as demonstrated in equation (6).

$$p_T = p_{LSL} + p_{USL} \quad (6)$$

Typically, because the output for each process stream will have a different mean and standard deviation, each stream will have different C_{pk} , p_{LSL} , and p_{USL} values.

2 Process Capability Indices for Multiple Stream Processes

Past Attempts

When faced with differences between process streams that cannot be easily corrected, some practitioners compute the C_{pk} index for all streams and then report just the highest value to customers. This practice is deceptive as it would cause customers to believe they are purchasing a higher level of quality than what they will actually receive.

To avoid misleading customers, other practitioners go to the opposite extreme and report just the lowest C_{pk} index of all the individual streams. Doing so means the customer will receive at least this quality level because the percentages of conforming parts are greater for the outputs from all the other streams. However, if there are large differences between the streams, reporting results for only the worst stream will make the manufacturer's overall quality level look much poorer than it really is.

A third approach has been to report an overall capability index by averaging the C_{pk} indices for the individual streams. To demonstrate why this technique is invalid, suppose an operation has just two process streams; the output of the first stream has a C_{pk} of 2.00 while the output of the second has a C_{pk} of 0.00. The average of these two C_{pk} indices is 1.00 $[(2.00 + 0.00)/2 = 1.00]$.

From equation (1), a C_{pk} index of 1.00 means that both C_{pl} and C_{pu} must also be at least 1.00. Then, from equations (4) and (5), neither p_{LSL} nor p_{USL} can be greater than 0.135%. Thus, a C_{pk} index of 1.00 implies that p_T for this shipment is at most 0.27%.

However, the process stream with a C_{pk} index of 2.00 is producing essentially 0.0% nonconforming parts (because C_{pu} and C_{pl} are both equal to, or greater than, 2.00, both p_{LSL} and p_{USL} are very close to 0). On the other hand, the second stream is generating at least 50% nonconforming parts (with either C_{pu} or C_{pl} equal to 0.00, either p_{LSL} or p_{USL} equals 0.5000 while the other is equal to, or less than, 0.5000).

Assuming both streams produce an equal number of parts in a given shipment, there will be an overall average of at least 25% nonconforming parts $(0 + 50\%)/2$, which is much larger than the 0.27% maximum expected from a shipment with a C_{pk} rating of 1.00. Thus, reporting a C_{pk} index of 1.00 for a shipment containing a mixture of the outputs from these two process streams would falsely imply a much higher quality level than will be received by the customer.

Table 1 Percentage nonconforming for each cavity

Cavity, i	$p_{i,LSL}$	$p_{i,USL}$
1	0.001	0.004
2	0.000	0.008
3	0.015	0.000
4	0.000	0.000
Sum for all four cavities	0.016	0.012
Average percentage nonconforming	0.004	0.003

Computing the Average C_{pk} Index

To correctly calculate the "average" C_{pk} of the combined output from several process streams, the percentage of out-of-specification parts produced by each stream should first be computed. For the four process streams of the plastic injection-molding operation mentioned earlier, these calculations generate the results displayed in Table 1. Here, i represents the cavity number, $p_{i,LSL}$ is the expected percentage of parts with hole diameters below the LSL for the i th cavity, while $p_{i,USL}$ is the expected percentage of parts with hole diameters above the USL for the i th cavity.

The individual percentages of parts with hole sizes below the LSL for each of these four cavities sum to a total of 0.016, or 1.6%, as is seen in the second-last row of Table 1. This total is then divided by four – in general, q , the number of process streams – to come up with 0.004 $(0.016/4)$. (This calculation assumes each of the four cavities contributes 1/4 of the total production of parts.) This last value is labeled $\bar{p}_{LSL}$, which represents the expected average percentage of parts with hole diameters below the LSL for the combined output of the four cavities. Equation (7) demonstrates this calculation for the injection-molding example.

$$\bar{p}_{LSL} = \frac{\sum_{i=1}^q p_{i,LSL}}{q} = \frac{\sum_{i=1}^4 p_{i,LSL}}{4} = \frac{0.016}{4} = 0.004 \quad (7)$$

The above average is an estimate of the percentage of parts below the LSL that the customer will receive, and is thus an indication of the process's ability to mold parts with hole sizes above the lower specification limit.

In a similar manner, an estimate of the expected average percentage of parts produced by this process

with hole diameters above the USL is 0.3%, as is shown by equation (8).

$$\bar{p}_{USL} = \frac{\sum_{i=1}^4 p_{i,USL}}{4} = \frac{0.012}{4} = 0.003 \quad (8)$$

The C_{pl} value, labeled $\bar{C}_{PL}$, corresponding to a given $\bar{p}_{LSL}$ is found by rearranging equation (4) to derive equation (9). The $\Phi^{-1}(\cdot)$ operator is the inverse of the $\Phi(\cdot)$ operator previously defined for equation (4). In Excel, $\Phi^{-1}(0.97725)$ corresponds to the statistical function NORMSINV(0.97725), which returns the value of 2.0000.

$$\bar{C}_{PL} = \frac{\Phi^{-1}(1 - \bar{p}_{LSL})}{3} \quad (9)$$

In the four-cavity die example, the $\bar{p}_{LSL}$ of 0.004 is associated with a $\bar{C}_{PL}$ value of 0.88, as is shown in equation (10).

$$\bar{C}_{PL} = \frac{\Phi^{-1}(1 - 0.004)}{3} = \frac{2.652}{3} = 0.88 \quad (10)$$

In a similar manner, equation (11) demonstrates how the $\bar{C}_{PU}$ value corresponding to a $\bar{p}_{USL}$ of 0.003 is calculated to be 0.92.

$$\begin{aligned} \bar{C}_{PU} &= \frac{\Phi^{-1}(1 - \bar{p}_{USL})}{3} = \frac{\Phi^{-1}(1 - 0.003)}{3} \\ &= \frac{2.748}{3} = 0.92 \end{aligned} \quad (11)$$

The average C_{pk} index is then defined as shown in equation (12).

$$\begin{aligned} \text{Average } C_{pk} &= \text{minimum}(\bar{C}_{PL}, \bar{C}_{PU}) \\ &= \text{minimum} \left[\frac{\Phi^{-1}(1 - \bar{p}_{LSL})}{3}, \frac{\Phi^{-1}(1 - \bar{p}_{USL})}{3} \right] \end{aligned} \quad (12)$$

Using the numbers from the previous example, an estimate of the average C_{pk} index for the hole diameters of parts produced by this four-cavity die molding process is computed in equation (13) to be 0.88.

$$\text{Average } C_{pk} = \text{Minimum}(0.88, 0.92) = 0.88 \quad (13)$$

Because this index is quite a bit less than 1.00, the combined output of the four cavities is *not* capable of consistently producing at least 99.73% conforming parts as far as hole size is concerned. In fact, it is producing 0.4% below the LSL and 0.3% above the USL, for a total of 0.7% nonconforming parts, which equates to 99.3% conforming parts. This is the quality level a customer would expect from a process having a traditional C_{pk} rating of 0.88.

Although the average C_{pk} index would be reported to customers, the individual cavity results are used to target process improvement efforts. In the four-cavity die example, cavity three has the largest percentage of parts (1.5%) with hole diameters below the LSL. Because it has no parts with holes above the USL, the initial capability improvement effort for this process should focus on increasing the average output of cavity three (refer back to Figure 1).

Average C_{pk} for Attribute Data

The average C_{pk} index can also be used when assessing the combined capability of several process streams that are generating attribute-type data (*see Measurement Systems Analysis, Attribute*), as this next example demonstrates.

Dough for frozen pizzas is arranged in six rows on a conveyor. The conveyor then moves under several dispensing stations that deposit various vegetables, meats, sauces, and cheeses on the dough. At the end of this line, the pizzas coming off each row are checked to see if the ingredients are evenly distributed over the face of the pizza. Those with an unacceptable ingredient distribution are rejected. After achieving stability for each of the six rows, their rejection rates are estimated and summarized in Table 2.

The average C_{pk} index for the combined output of these six process streams is computed by first determining the average percentage of pizzas that do not satisfy the appearance requirement, as seen in equation (14).

$$\bar{p} = \frac{\sum_{i=1}^6 p_i}{6} = \frac{0.000072}{6} = 0.000012 \quad (14)$$

4 Process Capability Indices for Multiple Stream Processes

Table 2 Rejection rate for each row

Row, i	p_i
1	0.000016
2	0.000003
3	0.000022
4	0.000012
5	0.000005
6	0.000014
Sum	0.000072
Average	0.000012

The average C_{pk} index is then computed as displayed in equation (15).

$$\text{Average } C_{pk} = \frac{\Phi^{-1}(1 - \bar{p})}{3} = \frac{4.224}{3} = 1.41 \quad (15)$$

Given a C_{pk} goal of 1.33, the combined output of these six rows demonstrates very good capability for meeting the appearance requirement.

Summary

The average C_{pk} index provides a method for assessing the outgoing quality of a mixture of outputs from similar process streams. This metric does not require

the outputs from all streams to be centered at the same value nor have the same process spread, which is important because, in most cases, they will not. In addition to part features (with either unilateral or bilateral specifications) whose measurements generate variable-type data, this index applies equally well to features characterized by attribute-type data.

References

- [1] Kane, V. (1986). Process capability indices, *Journal of Quality Technology* **18**, 41–52.
- [2] Montgomery, D. (1991). *Introduction to Statistical Quality Control*, 2nd Edition, John Wiley & Sons, New York, p. 43.

Further Reading

Bothe, D. (1977). *Measuring Process Capability*, Landmark Publishing, Cedarburg, pp. 721–728.

DAVIS ROSS BOTHE

Measurement Systems Analysis, Attribute

Introduction

Methods for assessing the capability of a quantitative or continuous **measurement system** are well documented and include the tabular method and the analysis of variance approach. When the measurement system involves attribute data, the standard quantitative methods are not appropriate. An attribute gauge measurement system is one in which the parts or objects under study are placed into one of two or more possible categories. The classification of the part is the measurement of interest in an attribute measurement system. Several commonly used statistics for “appraiser agreement” and attribute gauge repeatability and reproducibility (GRR) are presented.

Kappa Statistics

Two Appraisers

The κ statistic suggested by [1] measures the agreement between two appraisers (raters, judges) when evaluating or rating the same item. Cohen’s kappa statistic is given by

$$\kappa = \frac{p_o - p_e}{1 - p_e} \quad (1)$$

where p_o represents the observed proportion of agreement between appraisers and p_e represents the expected proportion of agreement due to chance. Table 1 illustrates a typical 2×2 table with two appraisers (X, Y) and two categories (pass, fail). The values a, b, c, d represent the number of times the appraisers agreed (a, d) and disagreed (b, c).

Table 1 Standard summary for two appraisers and two categories

		Ratings by appraiser X		Totals
		Pass	Fail	
Ratings by appraiser Y	Pass	a	b	n_1
	Fail	c	d	n_2
	Totals	m_1	m_2	N

Using the notation from Table 1, the observed proportion of agreement for this binary system is $p_o = (a + d)/N$ and the expected proportion of agreement due to chance is $p_e = (m_1 n_1 + m_2 n_2)/N^2$. If there is perfect agreement between appraisers, then $\kappa = 1$. If all observed agreement is due to chance, then $\kappa = 0$. κ will be negative if the agreement is less than what would be expected purely by chance. Cohen’s κ can be calculated for any two appraisers and is a single value summarizing the quality of measurements.

More than Two Appraisers

The κ statistic suggested by [2] is a variant to Cohen’s kappa statistic which takes into account more than two appraisers measuring the same number of items. Fleiss [2] suggested that the proportion of agreeing pairs be used to determine the degree of agreement among appraisers. Specifically, p_o and p_e are formulated to account for more than two appraisers by averaging the proportion of agreement for each pair of appraisers. Following the notation of [2], suppose there are N parts, n ratings per part, and k possible ratings (categories). Furthermore, suppose there are n_{ij} appraisers who will assign the i th part to the j th category. The proportion of agreeing pairs out of all possible pairs of category assignments is

$$p_i = \frac{1}{n(n-1)} \sum_{j=1}^k n_{ij}(n_{ij} - 1) \quad (2)$$

The average of all p_i ’s ($i = 1, \dots, N$) can be used as the measure for overall agreement,

$$p_o = \frac{1}{N} \sum_{i=1}^N p_i = \frac{1}{Nn(n-1)} \left(\sum_{i=1}^N \sum_{j=1}^k n_{ij}^2 - Nn \right) \quad (3)$$

If the extent of agreement is purely by chance, then we can expect the average proportion of agreement to be

$$p_e = \sum_{j=1}^k p_j^2 \quad (4)$$

The quantities given in equations (3) and (4) are used to calculate κ given in equation (1). The resulting κ statistic is often referred to as *Fleiss’s kappa*

statistic. It is of interest to note that Fleiss's kappa is useful for determining which category level(s) the appraisers (raters) are in relative agreement on, as well as the one(s) in which they do not have a common understanding (disagreement). In practice, this makes it easier to target remedial training or rewrite operational definitions of the categories in order to improve the measurement system. For more information on κ with multiple appraisers see [2–3] and [4].

Limitations of κ Statistics

κ cannot measure the degree of agreement or the size of disagreement between the two appraisers. Also, the value of κ can be misleading. For example, it would seem intuitive that a large value of p_o (indicating high agreement between appraisers), would result in a large value of κ . But, the resulting κ could be misleadingly small if p_e happens to be a large value. To illustrate, consider the data displayed in Table 2.

In this hypothetical situation, we have $p_o = (80 + 5)/100 = 0.85$, a value indicating a high level of agreement between the appraisers. Next, $p_e = [(95)(80) + (5)(20)]/100^2 = 0.77$. As a result, $\kappa = 0.34$. Therefore, a large value of p_e can result in a small value of κ , even if the proportion of observed agreement between appraisers is quite large. This phenomenon is one of two paradoxes associated with κ (see [5] for details and illustrations of both paradoxes).

Summary of κ Statistics

Overall, the κ statistics are easy to calculate and are often reported in standard statistical software. The κ statistics presented in this section can be used for binary data or ratings with more than two categories. The statistics provide some measure of reproducibility for attribute measurement systems. A

measure of repeatability is possible if each appraiser rates each part more than once. Often, repeatability is of less importance than reproducibility in these situations since appraisers tend to give similar (if not the same) ratings to the same part; this results in a high level of repeatability.

There are several drawbacks associated with the use of κ statistics, one of which has been illustrated here. Other problems with the use of Cohen's κ are that it assumes that the two appraisers use the same rating scale. Although this is a desirable condition, there may be situations where one appraiser would use one rating scale (say 1–8) while the other uses a different rating scale (say 1–10). This situation may be appropriate when there is interest in quantifying the consistency of ratings for an appraiser who uses different rating scales. There have been several extensions of Cohen's κ including a weighted κ statistic, statistics for paired data and for situations where **covariate** information is available. See [2–7] for more details concerning drawbacks and alternative statistics for interrater agreement.

Intraclass Correlation

The intraclass correlation coefficient (ICC) measures the correlation among multiple measurements of the same item or part. ICC was developed as an alternative to Cohen's κ [8] and has been used to evaluate reproducibility due to different appraisers. (Note: ICC is sometimes referred to as the *intraclass* κ and could have been discussed in the section titled “Kappa Statistics” (see [9]).)

The intraclass correlation is the correlation between different measurements Y_{ij} and Y_{ik} taken on the i th part or object and can be calculated as

$$ICC = \frac{\text{Cov}(Y_{ij}, Y_{ik})}{\sqrt{\text{Var}(Y_{ij}) \cdot \text{Var}(Y_{ik})}} = \frac{\sigma_p^2}{\sigma_p^2 + \sigma_e^2} \quad (5)$$

σ_p^2 represents the variability between parts σ_e^2 represents the variability within a part or object. An ICC of 0 would indicate lack of consistency among appraisers. An ICC = 1 would indicate perfect consistency or perfect reliability. It should be noted that for continuous measurement systems, ICC can also be written in terms of GRR for a one-way analysis of variance (see **Variance Components; Gauge Repeatability and Reproducibility (R&R)**,

Table 2 Sample data for two appraisers and two categories

		Ratings by appraiser X		
		Pass	Fail	Totals
Ratings by appraiser Y	Pass	80	5	80
	Fail	15	5	20
	Totals	95	5	100

Variance Components in; Gauge Repeatability and Reproducibility (R&R) Studies, Misclassification Rates). One summary of GRR for a one-factor random model is

$$GRR = \frac{\sigma_e}{\sqrt{\sigma_p^2 + \sigma_e^2}} \quad (6)$$

We can then write $ICC = 1 - (GRR)^2$. Although the ICC can be used to provide a measure of reproducibility, complete distributional assumptions are necessary since it is estimated using random-effects models.

Binary Measurement Systems

Suppose the quality characteristic under study can assume only two possible values (pass/fail, go/no go, and 0/1). Initially, also assume there are only two appraisers. For this binary measurement system, the ICC could be written as (this is identical to the ϕ coefficient defined in [4])

$$\begin{aligned} ICC_b &= \frac{\text{Cov}(Y_{ij}, Y_{ik})}{\sqrt{\text{Var}(Y_{ij}) \cdot \text{Var}(Y_{ik})}} \\ &= \frac{P(Y_{i1} = 1, Y_{i2} = 1) - p_1 p_2}{\sqrt{p_1(1 - p_1) \cdot p_2(1 - p_2)}} \end{aligned} \quad (7)$$

where $Y_{ij} = 1$ represents a successful rating on the i th part by the j th appraiser. The estimates for p_1 , p_2 , and $P(Y_{i1} = 1, Y_{i2} = 1)$ are

$$\begin{aligned} \hat{p}_1 &= \frac{\sum_{i=1}^n Y_{i1}}{n}, \quad \hat{p}_2 = \frac{\sum_{i=1}^n Y_{i2}}{n}, \\ \hat{P}(Y_{i1} = 1, Y_{i2} = 1) &= \frac{\sum_{i=1}^n Y_{i1} Y_{i2}}{n} \end{aligned} \quad (8)$$

and n is the number of parts or items being measured.

If the binary measurement system has m appraisers ($m > 2$), then the recommendation by [10] and [11] (and summarized in [4]) involves averaging the ICC_b values for all pairs of appraisers. They also assume that $p_j = p$ for all appraisers (see [4] for more discussion).

Analytic Method for Estimating Bias and Repeatability

Gauge studies are carried out to determine the suitability of the measurement system for use in a particular application. The *analytic method* detailed in [12] provides estimates of bias and repeatability for an attribute gauge and involves approximating the naturally discrete distribution with a continuous distribution, namely the normal (or Gaussian) distribution. This method assumes there are multiple parts or objects to be measured and that each part is measured multiple times by a single appraiser or measuring device. Once measured, the part is classified as acceptable (a) or unacceptable (e.g., binary data). Furthermore, it is assumed that a “standard” or reference value is known or can be obtained for the quality characteristic of interest.

The analytic method has the following requirements [12]:

- There must be a known reference value for each part selected.
- At least eight parts are selected at equidistant intervals.
- The minimum and the maximum values must represent the range of the process.
- Each part is measured $m = 20$ times under identical conditions.
- The number of “accepts” (a) is recorded.
 - Smallest part must have no accepts ($a = 0$). (If this is not satisfied, smaller parts are selected until $a = 0$.)
 - The largest part must have all accepts ($a = 20$). (If this is not satisfied, then larger parts are selected until $a = 20$.)
 - The remaining six parts must have between 1 and 19 accepts, inclusive ($1 \leq a \leq 19$). (If this is not satisfied, more parts are selected over the range – the additional parts are selected at the midpoint between the parts that have already been measured.)

In the analytic approach, gauge bias and gauge repeatability are then estimated by calculating probabilities of acceptance using normal probability plotting (see **Probability Plots**) – a procedure first put forth by [13]. Reproducibility is not estimable since there is a single gauge measuring the parts. The accuracy of the bias and repeatability estimates has

been questioned since it depends on the normal approximation of discrete data. The authors in [14] propose an improved approach for estimating bias and repeatability. In their work, they provide an alternative statistical estimation procedure that takes advantage of modern computational capabilities.

Latent-Class Models

A promising approach to assessing an attribute measurement system – repeatability and reproducibility in particular – is the use of latent-class models. The latent variable in this case represents the true value of the part or object, and as the name implies, is unknown. For interrater agreement, the latent-class model approach assumes that a part being rated or judged belongs to one of two or more possible latent classes. For example, in a binary system there would be two latent classes, namely, pass/fail. The probability that a part receives a particular rating is assumed to be dependent upon the latent class to which the part truly belongs. The joint distribution of the ratings is represented by a latent-class model and consists of a mixture of distributions for the classes of the latent variable. For example, in a binary measurement system, the latent-class model would consist of a mixture of binomial distributions.

Binary measurement systems are commonly encountered in industrial practice. Boyles [15] describes an approach for this situation using latent-class models. Using the notation and method presented in [15], suppose there are p parts and n_i ratings (inspections) of the i th part. Let s_i represent the number of “pass” evaluations of the i th part, i.e. $s_i = \sum_{j=1}^n y_{ij}$, where y_{ij} represents the j th rating of the i th part. For a binary measurement system, let $y_{ij} = 1$ if the i th part receives a pass evaluation (rating) on the j th rating, and $y_{ij} = 0$ if the i th part does not receive a pass evaluation (fails) on the j th rating. The distribution of s_i consists of a mixture of binomial distributions, namely

$$\begin{aligned} P(s_i = s) &= E[P(s_i = s | \xi_i)] \\ &= \binom{n_i}{s} [\phi \theta_1^s (1 - \theta_1)^{n_i - s} \\ &\quad + (1 - \phi) \theta_0^s (1 - \theta_0)^{n_i - s}] \end{aligned} \quad (9)$$

where ξ_i represents the latent classes with $\xi_i = 1$ if the i th part is good and $\xi_i = 0$ if the i th part

is bad. It is assumed that parts in each category of ξ are homogeneous. As a result, for a given level of ξ_i , the joint ratings of the appraisers are assumed to be statistically independent. If the parts are randomly selected from a population of parts, $\phi = P(\xi_i = 1)$ represents material variation. Boyles [15] further denotes the variability due to the measurement system by

$$\theta_0 = P(y_{ij} = 1 | \xi_i = 0) \quad (10)$$

$$\theta_1 = P(y_{ij} = 1 | \xi_i = 1) \quad (11)$$

Therefore, θ_0 is the probability that a part passes when the part is bad and θ_1 is the probability that a part passes when the part is good. The misclassification rates are θ_0 and $1 - \theta_1$. θ_0 and θ_1 are considered the gauge parameters. Furthermore, it is assumed that $\theta_0 < \theta_1$, i.e., the probability that a good part is classified as good is greater than the probability that a good part is classified as bad. Under these conditions (which [15] refers to as the basic model), it can be shown that

$$\begin{aligned} \text{Var}(y_{ij}) &= [\phi \theta_1 (1 - \theta_1) + (1 - \phi) \theta_0 (1 - \theta_0)] \\ &\quad + (\theta_1 - \theta_0)^2 \phi (1 - \phi) \end{aligned} \quad (12)$$

It is desired that estimates be obtained for the gauge parameters (θ_0 , θ_1) and the material variation (ϕ). Maximum-likelihood estimates (MLEs) (see **Maximum Likelihood**) can be found using the expectation-maximization (EM) algorithm where the likelihood function from equation (9) is

$$L \propto \prod_{i=1}^p \{ \phi \theta_1^s (1 - \theta_1)^{n_i - s} + (1 - \phi) \theta_0^s (1 - \theta_0)^{n_i - s} \} \quad (13)$$

The resulting MLEs of (θ_0 , θ_1 , ϕ) provide information about the misclassification levels in the system. In addition, **confidence intervals** can be constructed on the misclassification rates (see [15], p. 225; **Gauge Repeatability and Reproducibility (R&R) Studies, Misclassification Rates**). The author in [15] goes on to demonstrate how the reproducibility can be examined using this modeling approach. Detailed information on latent-class models, their development, and applications can be found in [4, 9, 16–20].

Summary and Conclusions

Several commonly used measures of assessment for attribute data have been presented. Attribute GRR studies can be carried out using appraiser agreement statistics (κ statistics, intraclass correlation), the analytic method, and latent-class models. Each approach can provide some measure of reproducibility or repeatability as discussed. The analytic method also provides some measure of system bias.

Commonly used criteria such as the κ statistics provide only a single statistic to describe the capability of the entire system. The latent-class model approach provides a model that can be used to describe the gauge capability for a measurement system.

If summary statistics are to be used, the ICC is a better measure of interrater agreement than Cohen's κ statistic, if the underlying marginal distributions are identical across all appraisers.

The latent-class model approach is extremely useful for the evaluation of attribute measurement systems; in particular, binary (pass/fail) measurement systems. Reference [15] employs latent-class models to estimate misclassification rates for binary measurement systems under a model where a standard is given and a model when no standard is given. The author in [15] describes a method for estimating repeatability and reproducibility effects utilizing misclassification rates. The authors in [4] compare latent-class models with κ statistics and log-linear models. As they have noted, the latent-class model approach provides a model for the outcome of measurement system, whereas the κ statistics and intraclass correlation are statistics – one-number summaries that are supposed to describe an entire measurement system. The authors in [4] recommend the latent-class model approach for dealing with binary measurement systems due to its flexibility and ability to handle various conditions and restrictions.

References

- [1] Cohen, J. (1960). A coefficient of agreement for nominal scales, *Educational and Psychological Measurement* **20**, 37–46.
- [2] Fleiss, J.L. (1971). Measuring nominal scale agreement among many raters, *Psychological Bulletin* **76**, 378–382.
- [3] Conger, A.J. (1980). Integration and generalization of kappas for multiple raters, *Psychological Bulletin* **88**, 322–328.
- [4] van Wieringen, W.N. & van Heuvel, E.R. (2005). A comparison of methods for the evaluation of binary measurement systems, *Quality Engineering* **17**, 495–507.
- [5] Feinstein, A.R. & Cicchetti, D.V. (1990). High agreement but low kappa: I. The problem of two paradoxes, *Journal of Clinical Epidemiology* **43**, 543–549.
- [6] Cicchetti, D.V. & Feinstein, A.R. (1990). High agreement but low kappa: II. Resolving the paradoxes, *Journal of Clinical Epidemiology* **43**, 551–558.
- [7] Uebersax, J.S. (1988). Validity inferences from interobserver agreement, *Psychological Bulletin* **104**, 405–416.
- [8] Bloch, D.A. & Kraemer, H.C. (1989). 2 × 2 kappa coefficients: measures of agreement or association, *Biometrics* **45**, 269–287.
- [9] Banerjee, M., Capozzoli, M., McSweeney, L. & Sinha, D. (1999). Beyond kappa: a review of interrater agreement measures, *The Canadian Journal of Statistics* **27**, 3–23.
- [10] Fleiss, J.L. (1965). Estimating the accuracy of dichotomous judgments, *Psychometrika* **30**, 469–479.
- [11] Bartko, J.J. & Carpenter, W.T. (1976). On the methods and theory of reliability, *Journal of Nervous and Mental Disease* **163**, 307–317.
- [12] Automotive Industry Action Group. (2002). *Measurement Systems Analysis Reference Manual*, 3rd Edition, Automotive Industry Action Group, Southfield.
- [13] McCaslin, J.A. & Gruska, G.F. (1976). Analysis of attribute gage systems, *ASQC Technical Conference Transactions* **30**, 392–399.
- [14] Sweet, A.L., Tjokrodjojo, S. & Wijaya, P. (2005). An investigation of the measurements systems analysis “Analytic method” for attribute gages, *Quality Engineering* **17**, 219–226.
- [15] Boyles, R.A. (2001). Gauge capability for pass-fail inspection, *Technometrics* **43**, 223–229.
- [16] Agresti, A. (1992). Modelling patterns of agreement and disagreement, *Statistical Methods in Medical Research* **1**, 201–218.
- [17] Agresti, A. & Lang, J.B. (1993). Quasi-symmetric latent class models, with application to rater agreement, *Biometrics* **49**, 131–139.
- [18] Agresti, A. (1988). A model for agreement between ratings on an ordinal scale, *Biometrics* **44**, 539–548.
- [19] Aickin, M. (1990). Maximum likelihood estimation of agreement in the constant predictive model, and its relation to Cohen's kappa, *Biometrics* **46**, 293–302.
- [20] Uebersax, J.S. & Grove, W.M. (1990). Latent class analysis of diagnostic agreement, *Statistical Methods* **9**, 559–572.

Related Articles

Gauge Repeatability and Reproducibility (R&R) Studies, Destructive Testing; Gauge Repeatability

and Reproducibility (R&R) Studies, Misclassification Rates; Gauge Repeatability and Reproducibility (R&R), Variance Components in;

Maximum Likelihood; Measurement Systems Analysis, Capability Measures for; Probability Plots; Variance Components.

CONNIE M. BORROR

Measurement Systems Analysis, Overview

Introduction

Measurement systems analysis (MSA) is a collection of statistical techniques designed to quantify the variability in a manufacturing process that is attributable to the metrology of that process. No process variability can be measured directly; that is, no process variability can be quantified without including the variability induced by measuring that process. MSA is therefore a study, separate from the processing of the product, that estimates the variability of the measurement process under conditions similar to those of the manufacturing process.

MSA is also called, variously, *measurement capability analysis* (in the context of determining whether or not a measurement system's variability is small enough compared with total process variability to adequately measure the process), *gauge studies*, or *gauge repeatability and reproducibility (GR&R) studies*.

Historical Perspective

Some of the earliest references to MSA in the statistical literature are found in the 1930s and 1940s. However, measurement precision was of concern even in older times. Shewhart [1] refers readers to a book by Goodwin [2] that was based on lecture notes from a course at MIT in the late 1800s. In 1948, Grubbs [3] laid out the first equations for estimating variability due to measurement repeatability.

The rise of MSA as a fundamental component of a process characterization study seems to have occurred in the 1980s, along with the rise of **statistical process control (SPC)** and other statistical quality methods in American manufacturing. Both the automotive and semiconductor industries mandate the use of gauge studies in quality improvement projects. Gauge studies are an integral part of both the Deming "Plan, Do, Check, Act" cycle, and the Six Sigma "Define, Measure, Analyze, Improve, Control" procedure. Wheeler and Lyday [4] include details on running gauge studies. In the early 1990s, SEMATECH included metrology system characterization as a

step in their equipment qualification plan [5]. A good review of methods for MSA can be found in Burdick *et al.* [6].

The design of the gauge study experiment has changed over the years. Early on, use of the standard 10-part, three-operator experimental design was extremely popular. In this design, multiple repeated measurements are performed by each operator, and the study yields estimates of variability due to repeats as well as variability due to operators. Another type of design used was the so-called 30 – 30 study, in which two studies are performed, the first with one setup and 30 repeated measurements, the other with 30 setups of one measurement each. The variance of the 30 repeated measurements is an estimate of repeatability, and the variance of the 30 setups is an estimate of reproducibility.

Terminology and methods of MSA have been somewhat standardized over the years; however, they have not been truly standardized. Care must be taken when performing or discussing MSA studies. Terms such as *repeatability* and *reproducibility* have been replaced by *variability due to repeats* or *variability due to operators*, for example. Industrial standard bodies have produced guideline documents that attempt standardization. See [7] for the automotive industry's guide to MSA or SEMI E89 [8] for the semiconductor industry's standard.

Experimental Design

A measurement process consists of the measurement tools themselves, including all hardware and software, as well as the procedures for using the tools: operators, setup and handling procedures, any off-line calculations or data entry, and **calibration** or other **preventive maintenance** frequency and technique.

A good measurement study will address the following questions: How big is the measurement error? What are the sources of the measurement error? Is the tool stable over time? Is the tool capable of making measurements for this process? What needs to be done to improve the measurement process?

Designing a MSA study is similar to designing any experiment (*see Gauge Repeatability and Reproducibility (R&R) Studies*). Up-front planning should include listing specific goals, obtaining measurement standards to use in the study, reviewing previous studies, identifying potential sources of variability,

estimating time and costs, assigning responsibilities, and securing management buy-in. Some of the potential sources of variability will not be studied, for example, yearly preventive maintenance activities. Others will be included in the study, for example, operator training and techniques or issues related to loading parts into the measurement tool. The factors that are not directly under study should still be documented and possibly added to a corrective action plan in the future.

Most gauge study information assumes the measurement tool gives continuous data and the true value of the standard remains constant over time (*see Gauge Repeatability and Reproducibility (R&R) Studies, Destructive Testing*).

The first step is to determine if the measurement tool is calibrated in the ranges of interest (*see Calibration*).

The fundamental equation of process variability is $\sigma_{\text{process}}^2 = \sigma_{\text{product}}^2 + \sigma_{\text{measurement}}^2$. That is, the variability seen in a process is made up of both variability due to the manufacturing process as well as variability due to the measurement process. Most gauge studies are designed to yield an estimate of $\sigma_{\text{measurement}}^2$. Given a list of possible contributors to $\sigma_{\text{measurement}}^2$, designing an experiment to estimate the size of the contributors to the total variability is usually straightforward. For example, consider an ellipsometer, used to measure the thickness of films on silicon wafers. Possible significant sources of variability include day-to-day effects, wafer-handling effects, and the actual measurement itself. The latter factor is often called *repeatability*. A good experiment will, therefore, include multiple days, multiple loads per day, and multiple repeated measurements per load. Note the hierarchical nature of the experimental factors. If the tool is not automated, multiple operators, or at least multiple levels of operator training, might be considered. If significant effects due to changes in ambient temperature and humidity over the day are expected, then measurements might be taken over multiple labor shifts.

Analytical Methods

Most of these effects will be treated as random factors in the ensuing model. Operators are often treated as fixed, especially if the operators in the study are the only operators used in the process or if different

operators represent different (fixed) levels of competency or training. Otherwise, Operator is treated as a random effect, and a variance component due to operators is calculated (*see Gauge Repeatability and Reproducibility (R&R), Variance Components in; Gauge Repeatability and Reproducibility (R&R) Studies, Confidence Intervals for*).

Once the estimate of measurement variability has been calculated, it can be used to answer the question of whether or not the measurement tool is capable of measuring the process under study. One popular metric for determining gauge capability is the precision-to-tolerance ratio, P/T. P/T is defined as the ratio of the measurement distribution to the process specification window, that is $6 \times \sigma_{\text{measurement}} / (USL - LSL)$. In automotive applications, the P/T ratio is often calculated as $5.15 \times \sigma_{\text{measurement}} / (USL - LSL)$. A rule of thumb is that a gauge with $P/T < 10\%$ is considered capable, one with $10\% < P/T < 30\%$ might be suitable, and one with $P/T > 30\%$ is not suitable for manufacturing. A large P/T ratio increases the misclassification probabilities (*see Gauge Repeatability and Reproducibility (R&R) Studies, Misclassification Rates*).

When specification limits are poorly defined, a **signal-to-noise** ratio (SNR) can be calculated. The SNR most often used is $SNR = \sigma_{\text{product}} / \sigma_{\text{measurement}}$. An SNR of at least 10 means that the measurement system is adequate for the process. A control system with a small SNR will have reduced sensitivity, meaning a longer time until an out-of-control condition is detected (*see Average Run Lengths and Operating Characteristic Curves*).

In addition to the metrology portion of the gauge study, care must be taken when choosing parts to use in the study. As with any statistical study, the results are applicable only to the population from which the sample was drawn. Therefore, it is crucial that the parts used in the study are representative of the process the measurement tool will be used for. Burdick *et al.* recommend using more parts and fewer measurements, as a rule.

Graphical Methods

In addition to the graphs generated for any other modeling effort, such as a time plot of the response, or a scatterplot of the response *versus* the factors, a key graph used in the analysis of gauge study data

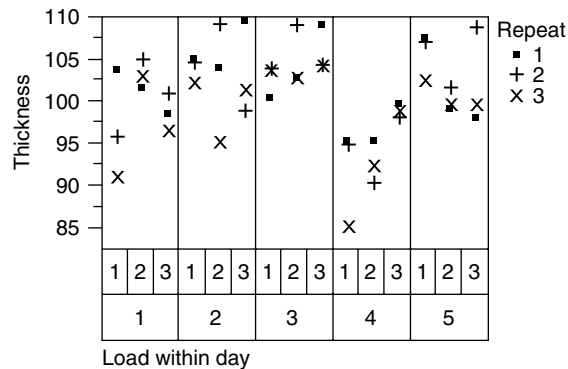


Figure 1 Multi-vari chart of film thickness gauge study data

is the multi-vari plot, also called a *variability graph*. This graph plots the response on the vertical axis, with all hierarchical factors but the lowest on the horizontal axis. Figure 1 illustrates the concept. The data are simulated from a gauge study on a metrology tool measuring wafer film thickness. The study was conducted over 5 days, with three loads (setups) per day and three repeated measurements per load. One part was used in this example.

From the graph, it is easily seen that day-to-day and repeat-to-repeat variability are large, while load-to-load variability is small. The **variance components** and percent contribution of each to the total are found in Table 1.

Stability

It is important for a measurement system to remain stable, as a shift in the measurement distribution can lead to false out-of-control signals or false in-control behavior on the process **control chart**.

Stability of the measurement system is determined in a similar way to stability of a process. Data are

Table 1 Variance components for the film thickness gauge study

Source	Variance component	Percent of total (%)
Day	13.95	43
Load	1.98	6
Repeat	16.26	51
Total	32.19	100

collected over time, oversampling from the potential significant sources of variability, and control charts are calculated (*see Control Charts, Overview*).

Typically, a pool of standards is used for control charting a measurement system. Control limits are calculated using the entire pool, and then one standard is measured at specified intervals for the control system. When that standard breaks or is otherwise retired, the control limits need only to be shifted to match the mean value of the new standard from the pool. The width of the limits remains the same.

As with any SPC system, an out-of-control action plan should be developed. Often the original brainstorming session that determined the factors to be considered in the gauge study will provide information for the out-of-control action plan. For example, if yearly preventive maintenance activities have the potential to cause measurement system variability, but are not included in the gauge study, one line in the action plan flow chart might read: "Has the tool been returned properly from preventive maintenance activities?"

Summary

MSA is an important set of techniques used to determine if the measurement process is suitable for a manufacturing process. MSA gives an estimate of the measurement process variability and goodness metrics can then be calculated. Verification of the suitability of a measurement process should be the first step in any process improvement activity.

References

- [1] Shewhart, W.A. (1986). *Statistical Method from the Viewpoint of Quality Control*, Dover.
- [2] Goodwin, H.M. (1908). *Elements of the Precision of Measurement and Graphical Methods*, G.H. Ellis Co.
- [3] Grubbs, F.E. (1948). On estimating precision of measuring instruments and product variability, *Journal of the American Statistical Association* **43**, 243–264.
- [4] Wheeler, D.J. & Lyday, R.W. (1989). *Evaluating the Measurement Process*, SPC Press.
- [5] SEMATECH Qualification Plan Guidelines for Engineers (1992). Technology Transfer Document 92061182A-GEN, and SEMATECH Qualification Plan Overview, Technology Transfer Document 91050538A-GEN.
- [6] Burdick, R.K., Borror, C.M. & Montgomery, D.C. (2003). A review of methods for measurement systems capability analysis, *Journal of Quality Technology* **35**, 342–354.

4 Measurement Systems Analysis, Overview

- [7] Automotive Industry Action Group (2002). *Measurement Systems Analysis*.
- [8] SEMI E89-1106E (2006). *Guide for Measurement System Analysis (MSA)*, <http://www.semi.org> (accessed 2006).

DIANE K. MICHELSON

Gauge Repeatability and Reproducibility (R&R), Variance Components in

Introduction

Traditionally, a standard gauge repeatability and reproducibility (R&R) study consisted of a two-factor crossed design with operators and parts; with each operator measuring each part say r times (*see Gauge Repeatability and Reproducibility (R&R) Studies; Gauge Repeatability and Reproducibility (R&R) Studies, Misclassification Rates; Gauge Repeatability and Reproducibility (R&R) Studies, Confidence Intervals for; Measurement Systems Analysis, Capability Measures for; and [1]*). There are numerous situations where there are more than two factors, the design factors are not crossed but possibly nested, or the design involves crossed and nested factors. In this article, **variance components** for **gauge R&R** studies are presented for these more complex experimental situations.

Basic Notation and Assumptions

Consider the model

$$Y = X + E \quad (1)$$

where Y is the measured value of a randomly selected part, X is the true measurement of the part, and E is the measurement error. It is assumed that X and E are independent where $X \sim N(\mu_p, \sigma_p^2)$ and $E \sim N(\mu_M, \sigma_E^2)$. As a result, Y is normally distributed with mean $\mu_Y = \mu_p + \mu_M$ and variance $\sigma_Y^2 = \sigma_p^2 + \sigma_E^2$. The mean, μ_M , is considered measurement bias and often assumed to be zero since it is believed that the measurement instrument has been properly calibrated. In this article, let γ_p represent the variance of a process, γ_M represent the variance of the **measurement system**. The variance components for a typical gauge R&R study are described in Table 1.

In the next sections, three more complex designs are investigated for gauge R&R studies. Although variance component estimation can be completed for random, fixed, or mixed models, the focus in this article is on random-effect models. For more details on gauge R&R studies, variance components for

Table 1 Variance components notation

Parameter	Description
γ_p	Variance of the process (often referred to as part-to-part variability)
γ_M	Variance of the measurement system (repeatability and reproducibility)
$\sigma_T^2 = \gamma_p + \gamma_M$	Total variance of the response

mixed or fixed-effects models for various situations, and **confidence interval** construction see references [2–23].

Three-Factor Crossed Design

Suppose that there are three factors of interest in a particular gauge R&R study. Two of these factors may be operators (O) and parts (P) as in the two-factor crossed design, but a third random factor, V (such as time, test instrument, etc.), may also be present. Furthermore, assume there are o operators, p parts, v levels of the third factor, and r replicates. A model that can represent this three-factor random crossed design is

$$\begin{aligned}
 Y_{ijkl} = & \mu_Y + P_i + O_j + V_k + (PO)_{ij} + (PV)_{ik} \\
 & + (OV)_{jk} + (POV)_{ijk} + E_{ijkl} \\
 i = & 1, \dots, p; \quad j = 1, \dots, o; \\
 k = & 1, \dots, v; \quad l = 1, \dots, r
 \end{aligned} \quad (2)$$

where μ_Y is a constant, and $P_i, O_j, V_k, (PO)_{ij}, (PV)_{ik}, (OV)_{jk}, (POV)_{ijk}$, and E_{ijkl} are jointly independent normal random variables with means of zero and variances $\sigma_p^2, \sigma_O^2, \sigma_V^2, \sigma_{PO}^2, \sigma_{PV}^2, \sigma_{OV}^2, \sigma_{POV}^2$, and σ_E^2 , respectively. Table 2 displays the analysis of variance (ANOVA) (*see Analysis of Variance; Degrees of Freedom*) including the expected mean squares (*see Mean Square Error; Sum of Squares*) for this model.

Part variation in this case is the variation of the manufacturing process. Variability due to operators (O), the third factor V (assumed not to be part variation), and all interactions involving operators, parts, and V make up variation due to the measurement system. Using the notation in Table 1, we have $\gamma_p = \sigma_p^2$ and $\gamma_M = \sigma_O^2 + \sigma_V^2 + \sigma_{PO}^2 + \sigma_{PV}^2 + \sigma_{OV}^2 + \sigma_{POV}^2 + \sigma_E^2$. These parameters can be written in terms of their expected mean squares and estimated by equating

2 Gauge R&R, Variance Components in

Table 2 Analysis of variance and expected mean squares for a three-factor crossed model

Source of variation	Degrees of freedom	Mean square	Expected mean square
Parts (P)	$p - 1$	S_1^2	$\theta_1 = \sigma_E^2 + r\sigma_{POV}^2 + vr\sigma_{PO}^2 + or\sigma_{PV}^2 + ovrr\sigma_P^2$
Operators (O)	$O - 1$	S_2^2	$\theta_2 = \sigma_E^2 + r\sigma_{POV}^2 + vr\sigma_{PO}^2 + pr\sigma_{OV}^2 + pvr\sigma_O^2$
V	$V - 1$	S_3^2	$\theta_3 = \sigma_E^2 + r\sigma_{POV}^2 + or\sigma_{PV}^2 + pr\sigma_{OV}^2 + opr\sigma_V^2$
PO	$(p - 1)(o - 1)$	S_4^2	$\theta_4 = \sigma_E^2 + r\sigma_{POV}^2 + vr\sigma_{PO}^2$
PV	$(p - 1)(v - 1)$	S_5^2	$\theta_5 = \sigma_E^2 + r\sigma_{POV}^2 + or\sigma_{PV}^2$
OV	$(o - 1)(v - 1)$	S_6^2	$\theta_6 = \sigma_E^2 + r\sigma_{POV}^2 + pr\sigma_{OV}^2$
POV	$(p - 1)(o - 1)(v - 1)$	S_7^2	$\theta_7 = \sigma_E^2 + r\sigma_{POV}^2$
Replicates	$pov(r - 1)$	S_8^2	$\theta_8 = \sigma_E^2$

the associated mean squares to the expected mean squares. For γ_P , we have

$$\gamma_P = \sigma_P^2 = \frac{\theta_1 + \theta_7 - \theta_4 - \theta_5}{ovr} \quad (3)$$

and

$$\hat{\gamma}_P = \hat{\sigma}_P^2 = \frac{S_1^2 + S_7^2 - S_4^2 - S_5^2}{ovr} \quad (4)$$

For measurement system variation, we have

$$\gamma_M = \frac{\sum_{q=2}^8 c_q \theta_q}{povr} \quad (5)$$

and

$$\hat{\gamma}_M = \frac{\sum_{q=2}^8 c_q S_q^2}{povr} \quad (6)$$

where $c_2 = o$, $c_3 = v$, $c_4 = (p - 1)o$, $c_5 = (p - 1)v$, $c_6 = ov - o - v$, $c_7 = o + v - po - pv - ov + pov$, and $c_8 = pov(r - 1)$. It can be shown that because the design is balanced, $n_q S_q^2 / \theta_q \sim \chi_{n_q}^2$ where n_q represents **degrees of freedom** and $q = 2, \dots, 8$. The mean squares are also jointly independent. Because of both these properties, an ANOVA is appropriate.

The estimates of the variance components for the gauge R&R study can be used to estimate other capability measures such as signal-to-noise ratio (SNR) and C_p (see **Gauge Repeatability and Reproducibility (R&R) Studies, Misclassification Rates; Gauge Repeatability and Reproducibility (R&R) Studies, Confidence Intervals for; Measurement Systems Analysis, Capability Measures for**). The

authors in [3] provide details for constructing confidence intervals on variance components for the gauge R&R study with three-factor crossed designs.

Example

Burdick *et al.* [3] present the results of a gauge R&R study conducted on the fabrication of magnetic tape for computerized data storage (this example data was originally presented in [2]). In this study, there are $o = 3$ operators (automated test stations) evaluating the quality of $p = 9$ parts (tape heads). To measure characteristics of the heads, $v = 3$ tapes are used in each test station. All the heads are measured with each of the test station by tape combinations and $r = 2$ replicates for each treatment condition. The resulting design is a three-factor crossed design with the random-effects model given in equation (2). In the process of writing information to magnetic tape, *voltage* is measured at two different frequencies. *Resolution* is a ratio of the two voltages and is used as the response variable of interest (expressed as a percentage). Ideally, this ratio should be 100%. The specification limits are $LSL = 90\%$ and $USL = 110\%$. An ANOVA is carried out using a three-factor random-effects model with the resulting ANOVA table for this experiment given in Table 3.

Part variation of the manufacturing system is then

$$\begin{aligned} \hat{\gamma}_P &= \hat{\sigma}_P^2 = \frac{S_1^2 + S_7^2 - S_4^2 - S_5^2}{ovr} \\ &= \frac{102.8 + 0.418 - 1.074 - 1.884}{18} = 5.57 \end{aligned}$$

The variability due to the measurement system is more complicated but can also be estimated. For this

Table 3 ANOVA Table for three-factor crossed design

Source of variation	Degrees of freedom	Mean square
Parts (P)	8	102.800
Operators (O)	2	13.080
V	2	29.230
PO	16	1.074
PV	16	1.884
OV	4	2.932
POV	32	0.418
Replicates	81	0.139

problem $o = 3, v = 3, p = 9, r = 2$ resulting in $c_2 = o = 3, c_3 = v = 3, c_4 = (p - 1)o = 24, c_5 = (p - 1)v = 24, c_6 = ov - o - v = 3, c_7 = o + v - po - pv - ov + pov = 24$, and $c_8 = pov(r - 1) = 81$. The estimate is then

$$\begin{aligned}\hat{\gamma}_M &= \frac{\sum_{q=2}^8 c_q S_q^2}{povr} \\ &= \frac{3(13.08) + 3(29.23) + 24(1.074) + 24(1.884) + 3(2.932) + 24(0.418) + 81(0.139)}{9(3)(3)(2)} \\ &= 1.4\end{aligned}$$

Fortunately, the values for the measurement system variability and part variation can be found from standard statistical software such as Minitab. The variance components for each source are given in Table 4. We can further recognize that R&R can be estimated separately since $\gamma_M = \sigma_{\text{repeatability}}^2 + \sigma_{\text{reproducibility}}^2$. Repeatability would be $\hat{\sigma}_{\text{repeatability}}^2 = \hat{\sigma}_E^2 = 0.1394$ and reproducibility would be the sum of variance components due to operators and different

Table 4 Variance component estimates for three-factor crossed designs

Source	Variance component
Part	5.5726
Operator	0.1758
Tape	0.4599
Part $\times$ Operator	0.1094
Part $\times$ Tape	0.2442
Operator $\times$ Tape	0.1397
Part $\times$ Operator $\times$ Tape	0.1395
Error	0.1394

setups involving tape heads:

$$\begin{aligned}\hat{\sigma}_{\text{reproducibility}}^2 &= \hat{\sigma}_O^2 + \hat{\sigma}_V^2 + \hat{\sigma}_{PO}^2 + \hat{\sigma}_{PV}^2 + \hat{\sigma}_{OV}^2 + \hat{\sigma}_{POV}^2 \\ &= 0.1758 + 0.4599 + 0.1094 + 0.2442 \\ &\quad + 0.1397 + 0.1395 = 1.269\end{aligned}\quad (7)$$

Finally, there is variability in the estimates themselves and providing simple point estimates as done here is not enough. It would be more beneficial to provide *confidence intervals* on different parameters of interest. Burdick *et al.* [3] use generalized inference (*see Gauge Repeatability and Reproducibility (R&R) Studies, Misclassification Rates; Gauge Repeatability and Reproducibility (R&R) Studies, Confidence Intervals for*) to construct confidence intervals on γ_P and γ_M . They show that the 95% confidence interval on γ_P in this example is [2.46, 20.9] while the 95% confidence interval on γ_M is [1.00, 33.3]. As the authors discuss, the intervals given here are very wide due to the small number of operators ($o = 3$) and number of tapes ($v = 3$) in the study. They recommend having at least four degrees of freedom for the ANOVA random model to reduce the length of the intervals.

Two-Factor Random Nested Designs

Nested designs occur regularly in practice. Jensen [14] presents an example of a wafer process in semiconductor manufacturing. Measurements are taken at multiple locations on a wafer and wafers are processed in groups (lots). The goal of the measurement system is to account for variability among lots, among wafers, and within each wafer. Lots and wafers are selected at random from the process and are therefore random factors. In this process, it is assumed that wafers (W) are nested within lots (L) and replicates are nested within wafers. The quality characteristic of interest is the response and can be modeled using a two-factor random nested model (see [3] and [14]):

$$y_{ijk} = \mu_y + L_i + W_{j(i)} + E_{ijk} \quad (8)$$

where y_{ijk} is the k th measurement on the j th wafer within the i th lot, μ_y is the overall process mean, L_i represents the effect of the i th lot, and $W_{j(i)}$ represents the j th wafer nested within lot. $L_i, W_{j(i)}$, and E_{ijk} are jointly independent random variables with means zero and variances σ_L^2, σ_W^2 , and σ_e^2 ,

4 Gauge R&R, Variance Components in

Table 5 Analysis of variance and expected mean squares for a two-factor random nested model

Source of variation	Degrees of freedom	Mean square	Expected mean square
Lots	$l - 1$	S_1^2	$\theta_1 = \sigma_E^2 + r\sigma_W^2 + wr\sigma_L^2$
Wafers within lots	$l(w - 1)$	S_2^2	$\theta_2 = \sigma_E^2 + r\sigma_W^2$
Replicates	$lw(r - 1)$	S_3^2	$\theta_3 = \sigma_E^2$

respectively. The ANOVA for the model given in equation (8) is given in Table 5. Again, it can be shown that $n_q S_q^2 / \theta_q \sim \chi_{n_q}^2$ with n_q degrees of freedom for the various components. The manufacturing process variability, γ_P , is due to lots and wafers while measurement system variability, γ_M , is due to replication error only. It can be shown that

$$\gamma_P = \frac{\theta_1 + (w - 1)\theta_2 - w\theta_3}{wr} \quad (9)$$

and

$$\gamma_M = \theta_3 \quad (10)$$

The estimates are then

$$\hat{\gamma}_P = \frac{S_1^2 + (w - 1)S_2^2 - wS_3^2}{wr} \quad (11)$$

and

$$\hat{\gamma}_M = S_3^2 \quad (12)$$

For this problem, there were $w = 2$ wafers selected from $l = 20$ lots, and $r = 9$ replicates. The resulting ANOVA is given in Table 6. The point estimates for γ_P and γ_M are then

$$\begin{aligned} \hat{\gamma}_P &= \frac{S_1^2 + (w - 1)S_2^2 - wS_3^2}{wr} \\ &= \frac{739.50 + 97.98 - 2(19.03)}{2(9)} = 44.41 \end{aligned}$$

Table 6 ANOVA table for two-factor nested design

Source of variation	Degrees of freedom	Mean square
Lots	19	739.50
Wafers within lots	20	97.98
Replicates	320	19.03

and

$$\hat{\gamma}_M = S_3^2 = 19.03$$

Again, point estimates do not give the practitioner the whole picture with respect to the amount of variability in the system. It would be more beneficial to examine confidence intervals on γ_P and γ_M . Burdick *et al.* [3] provide several approaches for constructing confidence intervals for this two-factor nested design. Also see [24] and [25].

Designs with Both Crossed and Nested Factors

There are experimental situations where the model involves both crossed and nested factors. The variance components for this particular model will be illustrated using an example from [6] and further detailed in [3]. The example involves three factors in a semiconductor manufacturing process. Measurements are recorded at four sites (S) on a wafer during the manufacturing process. The data are collected over three shifts (operators (O)) each day (D) for a seven-day period. Since operators may differ by day, the factor operator is treated as nested within day. The model of interest for this problem is

$$\begin{aligned} Y_{ijkl} &= \mu_Y + D_i + O_{j(i)} + S_k + (OS)_{jk(i)} + E_{ijkl} \\ \text{for } i &= 1, \dots, d; \quad j = 1, \dots, o; \quad k = 1, \dots, s; \\ \text{and } l &= 1, \dots, r \end{aligned} \quad (13)$$

μ_Y is a constant, $D_i \sim N(0, \sigma_D^2)$, $O_{j(i)} \sim N(0, \sigma_{O(D)}^2)$, $S_k \sim N(0, \sigma^2)$, $(OS)_{jk(i)} \sim N(0, \sigma_{SO(D)}^2)$, and $E_{ijkl} \sim N(0, \sigma_E^2)$. Furthermore, D_i , $O_{j(i)}$, S_k , $(OS)_{jk(i)}$, and E_{ijkl} are jointly independent. Day is treated as a blocking variable (see **Blocking**) and as a result, there is no day-by-site interaction in the model. The ANOVA and expected mean squares are given in Table 7 where $d = 7$, $s = 4$, $o = 3$, and $r = 4$.

In this problem, the process variation is a function of the variability due to day and site. In particular, $\gamma_P = \sigma_D^2 + \sigma_S^2$. The variation due to the measurement system is $\gamma_P = \sigma_{O(D)}^2 + \sigma_{SO(D)}^2 + \sigma_E^2$. It can be shown that

$$\gamma_P = \frac{d(\theta_1 - \theta_2) + s(\theta_3 - \theta_4)}{dosr} \quad (14)$$

and

$$\gamma_M = \frac{\theta_2 + (s - 1)\theta_4 + s(r - 1)\theta_5}{sr} \quad (15)$$

Table 7 Analysis of variance and expected mean squares for a model with crossed and nested factors

Source of variation	Degrees of freedom	Mean square	Expected mean square
Days	$d - 1$	S_1^2	$\theta_1 = \sigma_E^2 + r\sigma_{SO(D)}^2 + sr\sigma_{O(D)}^2 + osr\sigma_D^2$
Operators within days $O(D)$	$d(o - 1)$	S_2^2	$\theta_2 = \sigma_E^2 + r\sigma_{SO(D)}^2 + sr\sigma_{O(D)}^2$
Sites (S)	$s - 1$	S_3^2	$\theta_3 = \sigma_E^2 + r\sigma_{SO(D)}^2 + dor\sigma_S^2$
$S \times O(D)$	$(s - 1)(od - 1)$	S_4^2	$\theta_4 = \sigma_E^2 + r\sigma_{SO(D)}^2$
Replicates	$dos(r - 1)$	S_5^2	$\theta_5 = \sigma_E^2$

The point estimates are then

$$\hat{\gamma}_P = \frac{d(S_1^2 - S_2^2) + s(S_3^2 - S_4^2)}{dosr} \quad (16)$$

and

$$\hat{\gamma}_M = \frac{S_2^2 + (s - 1)S_4^2 + s(r - 1)S_5^2}{sr} \quad (17)$$

The resulting ANOVA table for the semiconductor problem is given in Table 8.

The point estimates for the process variation and measurement system variation for the semiconductor example are

$$\begin{aligned} \hat{\gamma}_P &= \frac{d(S_1^2 - S_2^2) + s(S_3^2 - S_4^2)}{dosr} \\ &= \frac{7(6.480 - 1.404) + 4(35.220 - 1.447)}{7(3)(4)(4)} \\ &= 0.508 \end{aligned}$$

and

$$\begin{aligned} \hat{\gamma}_M &= \frac{S_2^2 + (s - 1)S_4^2 + s(r - 1)S_5^2}{sr} \\ &= \frac{1.404 + 3(1.447) + 4(3)(0.529)}{4(4)} \\ &= 0.756 \end{aligned}$$

Table 8 ANOVA table for design involving crossed and nested factors

Source of variation	Degrees of freedom	Mean square
Days	6	6.480
Operators within days $O(D)$	14	1.404
Sites (S)	3	35.220
$S \times O(D)$	60	1.447
Replicates	252	0.529

Table 9 Variance component estimates for the semiconductor example

Source of variation	Variance component
Days	0.10570
Operators within days $O(D)$	-0.00272
Sites (S)	0.40200
$S \times O(D)$	0.22950
Replicates	0.52920

As in the previous examples, it is recommended to construct confidence intervals on these variances. The individual variance component estimates are given in Table 9. It should be noted that there may be some slight differences in the estimates found using different software packages or doing the calculations by hand due to rounding. The values given in Table 9 illustrate one of the problems that can result from the estimation procedure used for computation of the variance components. Notice that the variance component estimate for operators within days is negative. When this occurs, one recommendation has been to set this component equal to zero. Although the full ANOVA table is not reported here, when the test was conducted, the nested factor, operators within days, was statistically insignificant at the 0.05 level of significance. The authors in [3] removed this factor from the model, analyzed the results for the reduced model, and then calculated the variance components and resulting process and measurement system variation.

Conclusions

Variance component estimation for three complex experimental designs has been presented where variance estimates, such as process variation and measurement system variation, have been shown to be linear combinations of the appropriate individual

variance components. The variance components were estimated by equating the expected mean squares for the factor (or interaction) of interest to its observed mean square value. It is recommended that the practitioner should not rely simply on point estimates of the variance components, but also consider confidence intervals on the variance components. Confidence intervals provide a measure of precision associated with the estimated variance component and can give the practitioner insight as to whether more operators should be included in a future study or perhaps more parts are needed.

In addition, all the examples presented here assumed that every factor was random and the design balanced. There are many scenarios where some of the factors involved in the study are fixed so that the resulting model is mixed (contains both random and fixed factors). It is also possible that the resulting design will not be balanced. Several studies have been conducted on variance component estimation for this situation. For more detailed information on these designs and construction of confidence intervals see [12, 24–29].

References

- [1] Vardeman, S.B. & VanValkenburg, E.S. (1999). Two-way random-effects analyses and gauge R&R studies, *Technometrics* **41**, 202–211.
- [2] Adamec, E. & Burdick, R. (2003). Confidence intervals for a discrimination ratio in a gauge R&R study with three random factors, *Quality Engineering* **15**, 383–389.
- [3] Burdick, R.K., Borror, C.M. & Montgomery, D.C. (2005). *Design and analysis of Gauge R&R studies: Making decision with confidence intervals in random and mixed ANOVA models*, ASA-SIAM Series, Philadelphia, PA.
- [4] Burdick, R.K. & Graybill, F.A. (1992). *Confidence Intervals on Variance Components*, Marcel Dekker, Inc., New York.
- [5] Burdick, R.K. & Larsen, G.A. (1997). Confidence intervals on measures of variability in R&R studies, *Journal of Quality Technology* **29**, 261–273.
- [6] Borror, C.M., Montgomery, D.C. & Runger, G.C. (1997). Confidence intervals for variance components from gauge capability studies, *Quality and Reliability Engineering International* **13**, 361–369.
- [7] Chiang, A.K.L. (2001). A simple general method for constructing confidence intervals for functions of variance components, *Technometrics* **43**, 356–367 (see also discussion by Iyer and Mathew noted below).
- [8] El-Bassiouni, M.Y. (1994). Short confidence intervals for variance components, *Communications in Statistics-Theory and Methods* **23**, 1915–1933.
- [9] Graybill, F.A. (1976). *Theory and Application of the Linear Model*, Duxbury Press, North Scituate.
- [10] Graybill, F.A. & Wang, C.M. (1980). Confidence intervals on nonnegative linear combinations of variances, *Journal of the American Statistical Association* **75**, 869–873.
- [11] Gui, R., Graybill, F.A., Burdick, R.K. & Ting, N. (1995). Confidence intervals on ratios of linear combinations for non-disjoint sets of expected mean squares, *Journal of Statistical Planning and Inference* **48**, 215–227.
- [12] Hamada, M. & Weerahandi, S. (2000). Measurement system assessment via generalized inference, *Journal of Quality Technology* **32**, 241–253.
- [13] Iyer, H.K. & Mathew, T. (2002). Comments on Chiang (2001), *Technometrics* **44**, 284–285.
- [14] Jensen, C.R. (2002). Variance component calculations: common methods and misapplications in the semiconductor industry, *Quality Engineering* **14**, 645–657.
- [15] Leiva, R.A. & Graybill, F.A. (1986). Confidence intervals for variance components in the balanced two-way model with interaction, *Communications in Statistics-Simulation and Computation* **15**, 301–322.
- [16] Lu, T.-F.C., Graybill, F.A. & Burdick, R.K. (1987). Confidence intervals on the ratio of expected mean squares $(\theta_1 + d\theta_2)/\theta_3$, *Biometrics* **43**, 535–543.
- [17] Lu, T.-F.C., Graybill, F.A. & Burdick, R.K. (1988). Confidence intervals on a difference of expected mean squares, *Journal of Statistical Planning and Inference* **18**, 35–43.
- [18] Lu, T.-F.C., Graybill, F.A. & Burdick, R.K. (1989). Confidence intervals on the ratio of expected mean squares $(\theta_1 - d\theta_2)/\theta_3$, *Journal of Statistical Planning and Inference* **21**, 179–190.
- [19] Paark, K. & Burdick, R.K. (1998). Confidence intervals for the mean in a balanced two-factor random effect model, *Communications in Statistics-Theory and Methods* **27**, 2807–2825.
- [20] Ting, N., Burdick, R.K. & Graybill, F.A. (1991). Confidence intervals on ratios of positive linear combinations of variance components, *Statistics and Probability Letters* **11**, 523–528.
- [21] Ting, N., Burdick, R.K., Graybill, F.A., Jeyaratman, S. & Lu, T.-F.C. (1990). Confidence intervals on linear combinations of variance components that are unrestricted in sign, *Journal of Statistical Computation and Simulation* **35**, 135–143.
- [22] Williams, J.S. (1962). A confidence interval for variance components, *Biometrika* **49**, 278–281.
- [23] El-Bassiouni, M.Y. & Abdelhafez, M.E.M. (2000). Interval estimation of the mean in a two-stage nested model, *Journal of Statistical Computation and Simulation* **67**, 333–350.
- [24] Wang, C.M. & Graybill, F.A. (1981). Confidence intervals on a ratio of variances in the two-factor nested components of variance model, *Communications in Statistics-Theory and Methods* **A10**, 1357–1368.

-
- [25] Dolezal, K.K., Burdick, R.K. & Birch, N.J. (1998). Analysis of a two-factor R&R study with fixed operators, *Journal of Quality Technology* **30**, 163–170.
- [26] Hernandez, R.P. & Burdick, R.K. (1993). Confidence intervals and tests of hypothesis on variance components in an unbalanced two-factor crossed design with interaction, *Journal of Statistical Computation and Simulation* **47**, 67–77.
- [27] Hernandez, R.P., Burdick, R.K. & Birch, N. (1992). Confidence intervals and tests of hypotheses on variance components in an unbalanced two-fold nested design, *Biometrical Journal* **34**, 387–402.
- [28] Park, D.J. & Burdick, R.K. (2003). Performance of confidence intervals in regression models with unbalanced one-fold nested error structures, *Communications in Statistics-Simulation and Computation* **32**, 717–732.
- [29] Tsui, K. & Weerahandi, S. (1989). Generalized p-values in significance testing of hypotheses in the presence of nuisance parameters, *Journal of the American Statistical Association* **84**, 602–607 (Corrections: 86(1991), p. 256).
- CONNIE M. BORROR, RICHARD K. BURDICK
AND DOUGLAS C. MONTGOMERY

Gauge Repeatability and Reproducibility (R&R) Studies, Misclassification Rates

Introduction

In order to properly monitor and improve a manufacturing process, it is necessary to measure attributes of the process output. For any group of measurements collected for this purpose, at least part of the variation is due to the measurement system. This is because repeated measurements of any particular item do not always result in the same value. In order to ensure that measurement system variability is not harmfully large, it is necessary to conduct a gauge repeatability and reproducibility (R&R) study (*see Gauge Repeatability and Reproducibility (R&R) Studies; Measurement Systems Analysis, Overview*). For readers who are unfamiliar with these terms, a gauge is any device used to obtain measurements on the output of the manufacturing process. In a **gauge R&R** study, replicate measurements are collected on units using several different operators, setups, or time periods. Two components of variability are generated in such studies: repeatability and reproducibility. Repeatability represents the gauge variability when it is used to measure the same unit (with the same operator or setup or in the same time period). Reproducibility refers to the variability arising from different operators, setups, or time periods. By comparing the relative magnitudes of these sources of variation with the variation inherent in the manufacturing process, decisions are made concerning the capability of the measurement system. In particular, if variation due to the measurement system is small relative to variation of the process, then the measurement system is deemed capable and used to monitor the process.

Originally, gauge R&R studies were conducted using a tabular method based on ranges and **control charts**. Today, however, gauge R&R studies are analyzed using analysis of variance (ANOVA) models (*see Analysis of Variance; Gauge Repeatability and Reproducibility (R&R), Variance Components in*). This is because the ANOVA models provide greater flexibility in modeling the data and constructing measures of uncertainty.

The Measurement Model

The measurement model used in a gauge R&R study can be written as

$$Y = X + E \quad (1)$$

where Y is the measured value of a randomly selected part from a manufacturing process, X is the true value of the part, and E is the measurement error contributed by the measurement system. The terms X and E are independent normal random variables with means μ_P and μ_M , and variances γ_P and γ_M , respectively. The subscripts P and M denote process and measurement system, respectively. These assumptions imply that Y has a **normal distribution** with mean $\mu_Y = \mu_P + \mu_M$ and variance $\gamma_Y = \gamma_P + \gamma_M$. Additionally, the **covariance** between Y and X is γ_P . The mean μ_M is referred to as the *bias* of the measurement system. This bias can typically be eliminated by proper **calibration** of the system, and so it is often assumed that $\mu_Y = \mu_P$.

Parameters of Interest

Table 1 reports parameters in model (1) that are the focus of most gauge R&R studies.

It is useful to describe these parameters in the context of an example. Houf and Berman [1] describe a process used to manufacture a power module for a line of motor starters. An important attribute is the thermal performance of the power module. Measurements of thermal performance are taken on a sample of power modules using a thermal resistance measuring instrument. The units of measurement are degree celsius per watt. Variability in the data can be attributed to

1. the manufacturing process used to produce the power modules (γ_P) and
2. measurement error contributed by the thermal resistance measuring instrument (γ_M).

The purpose of the gauge R&R study is to determine if the variability due to 2 is small relative that due to 1. There are two criteria commonly used for this purpose. Both are based on the parameters shown in Table 1.

The precision-to-tolerance ratio (PTR) is

$$PTR = \frac{k\sqrt{\gamma_M}}{USL - LSL} \quad (2)$$

2 Gauge R&R Studies, Misclassification Rates

Table 1 Parameters in a gauge R&R study

Symbol	Definition
μ_Y	Mean of population of measurements
γ_P	Variance of the process
γ_M	Variance of the measurement system
$\gamma_R = \frac{\gamma_P}{\gamma_M}$	Ratio of process variance to measurement variance

where USL and LSL are the upper and lower specification limits, respectively, and k is either 5.15 or 6. The Automotive Industry Action Group (AIAG) *Measurement Systems Analysis* manual [2, p. 77] recommends the following rule of thumb for using PTR to determine the capability of a measurement system:

1. $PTR \leq 0.1$: the measurement system is capable.
2. $0.1 < PTR \leq 0.3$: the measurement system may be capable depending on factors such as process capability and costs of misclassification.
3. $PTR > 0.3$: the measurement system is not capable.

The signal-to-noise ratio (SNR) is a function of γ_R defined as

$$SNR = \sqrt{2\gamma_R} \quad (3)$$

The AIAG *Measurement Systems Analysis* manual [2, p. 117] defines SNR as the number of distinct categories that can be reliably distinguished by the measurement system. A value of five or greater is recommended, and a value less than two indicates the measurement system is of no value in monitoring the process.

The PTR and SNR are popular criteria for evaluating the capability of a measurement system. Burdick *et al.* [3] critique these criteria and review other criteria reported in the literature. However, given that a measurement system is designed to discriminate between good and bad products, a more meaningful criterion might be the probability of part misclassification. We now describe how this criterion can be used to determine the capability of a measurement system.

Definition of Misclassification Rates

The joint **probability density function** of Y and X is bivariate normal with mean vector $[\mu_Y \mu_P]^T$ and

covariance matrix

$$\begin{bmatrix} \gamma_P + \gamma_M & \gamma_P \\ \gamma_P & \gamma_P \end{bmatrix} \quad (4)$$

where $[\bullet]^T$ represents the transpose operator. We represent this density function as $f(y, x)$. A manufactured part is in conformance if

$$LSL \leq X \leq USL \quad (5)$$

A measurement system will pass a part if

$$LSL \leq Y \leq USL \quad (6)$$

If equation (5) is true and equation (6) is false, then a conforming part is misclassified as a failure. This is called a *false failure* (FF). Alternatively, if equation (5) is false but equation (6) is true, a failure is misclassified as a good part. This is called a *missed fault* (MF). Figure 1 illustrates the regions of FF and MF on an equal density contour of the bivariate normal distribution.

Probabilities of the misclassification errors FF and MF are computed as either joint or conditional probabilities. To demonstrate, an FF occurs when equation (5) is true and equation (6) is false. Representing this joint probability with the symbol δ ,

$$\begin{aligned} \delta &= \Pr[(LSL \leq X \leq USL) \text{ and} \\ &\quad (Y < LSL \text{ or } Y > USL)] \\ &= \Pr[(LSL \leq X \leq USL \text{ and } Y < LSL) \\ &\quad \text{or } (LSL \leq X \leq USL \text{ and} \\ &\quad Y > USL)] \end{aligned}$$

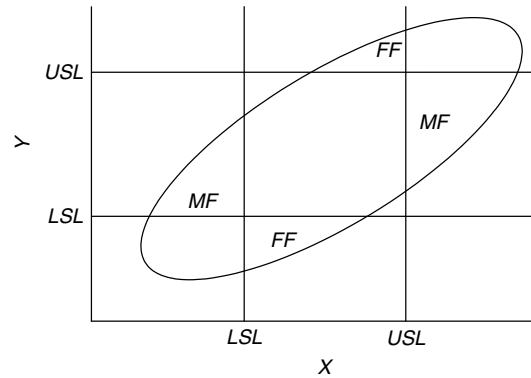


Figure 1 MF and FF regions of an equal density contour

$$\begin{aligned}
&= \int_{LSL}^{USL} \int_{-\infty}^{LSL} f(y, x) dy dx \\
&\quad + \int_{LSL}^{USL} \int_{USL}^{\infty} f(y, x) dy dx \quad (7)
\end{aligned}$$

An MF occurs when equation (5) is false and equation (6) is true. If we represent this joint probability with the symbol β ,

$$\begin{aligned}
\beta &= \Pr[(X < LSL \text{ or } X > USL) \text{ and} \\
&\quad (LSL \leq Y \leq USL)] \\
&= \Pr[(X < LSL \text{ and } LSL \leq Y \leq USL) \text{ or} \\
&\quad (X > USL \text{ and } LSL \leq Y \leq USL)] \\
&= \int_{-\infty}^{LSL} \int_{LSL}^{USL} f(y, x) dy dx \\
&\quad + \int_{USL}^{\infty} \int_{LSL}^{USL} f(y, x) dy dx \quad (8)
\end{aligned}$$

The probability δ is referred to as *producer's risk* and the probability β is referred to as *consumer's risk*. Some investigators prefer to compute the conditional probabilities

$$\delta_C = \frac{\delta}{\pi} \quad (9)$$

and

$$\beta_C = \frac{\beta}{1 - \pi} \quad (10)$$

where

$$\pi = \int_{LSL}^{USL} f(x) dx \quad (11)$$

is the probability of a good part and $f(x)$ is the marginal probability density function for X , which is normal with mean μ_P and variance γ_P . In words, δ_C is the conditional probability that a part fails given the part is good. The symbol β_C is the conditional probability that a part passes given the part is bad.

When considering whether misclassification rates are “sufficiently small”, it is useful to compare them to what one could obtain by chance. To demonstrate, suppose the value of π in equation (11) is known. A chance process for classifying parts is to take no measurements and randomly classify π of the parts as “good” and $1 - \pi$ as “bad”. For this chance process, the probability of an FF is

$$\delta^* = \Pr[(LSL \leq X \leq USL) \text{ and}$$

(Part classified as bad)]

$$\begin{aligned}
&= \Pr(LSL \leq X \leq USL) \times \Pr(\text{Part classified} \\
&\quad \text{as bad}) \\
&= \pi(1 - \pi) \quad (12)
\end{aligned}$$

Likewise, the probability of an MF for this chance process is

$$\begin{aligned}
\beta^* &= \Pr[(X < LSL \text{ or } X > USL) \text{ and} \\
&\quad (\text{Part classified as good})] \\
&= \Pr(X < LSL \text{ or } X > USL) \times \Pr(\text{Part} \\
&\quad \text{classified as good}) \\
&= (1 - \pi)\pi \quad (13)
\end{aligned}$$

Thus, for a measurement system to be useful, the misclassification rates δ and β should be less than what one could get by chance, δ^* and β^* . One can make a similar argument using the conditional probabilities δ_C and β_C . In this case, the conditional probabilities that could be obtained by chance are $\delta_C^* = 1 - \pi$ and $\beta_C^* = \pi$.

To better interpret misclassification rates, it is useful to compute ratios of the misclassification rates to the chance rates. If a measurement system is able to discriminate better than the chance process, these indexes will be less than one. These ratios are defined as

$$\delta_{\text{index}} = \frac{\delta}{\delta^*} = \frac{\delta}{\pi(1 - \pi)} \quad (14)$$

and

$$\beta_{\text{index}} = \frac{\beta}{\beta^*} = \frac{\beta}{(1 - \pi)\pi} \quad (15)$$

Notice $\delta_C/\delta_C^* = \delta_{\text{index}}$ and $\beta_C/\beta_C^* = \beta_{\text{index}}$ so these indexes can be used for both the unconditional and conditional definitions of misclassification rates.

The misclassification indexes in equations (14) and (15) can be computed for given values of μ_Y , μ_P , γ_P , γ_M , LSL, and USL by integrating bivariate normal distributions as shown in equations (7), (8), and (11). More typically, these parameters are unknown and must be estimated from the data obtained in the R&R study. If we assume $\mu_Y = \mu_P$, the overall sample mean provides an estimate of both μ_Y and μ_P . If there is a systematic bias in the system, then it is necessary to estimate μ_P from another source of information. Such an estimate might be based on historical data

4 Gauge R&R Studies, Misclassification Rates

from the process. Estimates for γ_P and γ_M will be available from the measurement data even if systematic bias is present in the measurement system.

One approach to account for the uncertainty in the estimates of μ_Y , μ_P , γ_P , and γ_M is to compute **confidence intervals** for δ_{index} and β_{index} , as well as for the misclassification rates in equations (7–10). A method to compute such intervals based on generalized inference is detailed in Burdick *et al.* [4]. Simply, the calculations are performed by replacing the parameters μ_Y , μ_P , γ_P , and γ_M with appropriate generalized pivotal quantities.

An Example

The two-factor crossed random effects ANOVA model with interaction is the classical gauge R&R design. Typically, the two factors are referred to as *parts* and *operators*. Table 2 presents the ANOVA table for an artificial data set based on the experiment described by Houf and Berman [1]. The response variable is the thermal performance of a module measured in degree celsius per watt. Each response has been multiplied by 100 for convenience of scale. The data represent measurements of 20 parts recorded by six operators. Each part is measured two times by each operator. It is assumed that parts and operators are selected at random from larger populations. The mean of all the observations is 36.9 and the specification limits are $LSL = 18$ and $USL = 58$.

The confidence intervals for the misclassification rates are shown in Table 3. These are computed using the process described in Burdick *et al.* [4].

Table 3 suggests the measurement system in this example is capable since both misclassification indexes are less than one. This suggests the misclassification rates are better than those obtained by chance, and that there is information contained in the measurements.

Table 2 ANOVA for example

Source of variation	Degrees of freedom	Mean square
Parts (modules)	19	591.5
Operators	5	13.59
$P \times O$	95	2.060
Replicates	120	0.7333

Table 3 95% confidence intervals for misclassification rates (bounds for δ , δ_C , β , and β_C are $\times 10^{-6}$)

Parameter	Lower bound	Upper bound
δ	199	8594
δ_C	199	9135
δ_{index}	0.141	0.854
β	63	6138
β_C	98 746	248 703
β_{index}	0.105	0.249

Summary

Although the two-factor random effects design is the most common gauge R&R design, more complex ANOVA designs are needed to properly model some measurement systems. Burdick *et al.* [5] provide methods for constructing confidence intervals on the misclassification rates in addition to the parameters shown in Table 1 for a wide variety of ANOVA designs (also see **Gauge Repeatability and Reproducibility (R&R) Studies, Confidence Intervals for**). Computer programs in Excel and SAS are available at the book's website.

References

- [1] Houf, R. & Berman, D. (1988). Statistical analysis of power module thermal test equipment performance, *IEEE Transactions on Components, Hybrids, and Manufacturing Technology* **11**, 516–520.
- [2] Automotive Industry Action Group (AIAG). (2002). *Measurement Systems Analysis*, 3rd Edition, AIAG, Detroit.
- [3] Burdick, R., Borror, C. & Montgomery, D. (2003). A review of methods for measurement systems capability analysis, *Journal of Quality Technology* **35**, 342–354.
- [4] Burdick, R., Park, Y.-J., Montgomery, D. & Borror, C. (2005). Confidence intervals for misclassification rates in a gauge R&R study, *Journal of Quality Technology* **37**, 294–303.
- [5] Burdick, R., Borror, C. & Montgomery, D. (2005). *Design and Analysis of Gauge R&R Studies: Making Decisions with Confidence Intervals in Random and Mixed ANOVA Models*, ASA-SIAM Series on Statistics and Applied Probability, Society for Industrial and Applied Mathematics (SIAM), Philadelphia.

RICHARD K. BURDICK

Gauge Repeatability and Reproducibility (R&R) Studies, Confidence Intervals for

Introduction: Measurement Errors in Data

Gauge repeatability and reproducibility (R&R) studies are commonly conducted in industrial settings with the goal of determining if a gauge or, broadly, a **measurement system** is adequate for monitoring a manufacturing process. Such studies are also termed as *gauge capability studies*. Variables type measurement (Y) of a manufactured part can be thought of as comprising the true value (P) of the part plus measurement error (M). Measurement errors are attributable to the gauge (repeatability) or gauge operator/operating conditions (reproducibility). A given part is regarded as nonconforming if its true value lies outside of design specification (LSL, USL). Since the true value is never directly observable, the practical approach is to check whether Y lies within (LSL, USL). This approach assumes that measurement error M is small relative to the true value P . The common statistical model for analyzing measurement errors assumes additivity of effects and normal distributed variables: $Y = P + M$. The fundamental goal of **gauge R&R** studies is to determine if the measurement system variance (σ_M^2) is practically small relative to the manufacturing process variance (σ_P^2). We will discuss common gauge R&R measures and **confidence intervals** (CIs) for them.

Common Measures of Gauge Capability

The common measures of gauge capability in R&R studies are derived from two components of the measurement variance: $\sigma_Y^2 = \sigma_P^2 + \sigma_M^2$.

Measurement Error Variance

σ_M^2 is the variance of measurement errors. It is typically decomposed into repeatability variance and reproducibility variance. Repeatability variance is attributed to the variation among repeated measurements made by the same operator on the

same part using the same gauge. Reproducibility variance is attributed to the variation among repeated measurements made on the same part by different operators with the same gauge. Rescaling σ_M , we get a commonly used measure of gauge capability called the *precision-to-tolerance ratio* (PTR), $PTR = 6\sigma_M/(USL - LSL)$. Unfortunately, as discussed in Montgomery and Runger [1], low PTR could be due to a large tolerance ($USL - LSL$) rather than small σ_M . Consequently, we are not assured that measurement errors (E) are small relative to the true values (P).

Part Variance

The part variance is denoted by σ_P^2 . If the parts used in the R&R study are randomly selected from the production process, then this variance is also taken to be the variance of the production process. However, if the parts do not represent a random sample from the process but include unusual parts that exceed specification (as suggested in [1]), then this variance cannot be regarded as a measure of production process variance.

Comparing Part and Measurement Variance

The comparison between parts and measurement variance can be made using three measures. The most direct being the ratio, $\gamma = \sigma_P^2/\sigma_M^2$. γ is a direct comparison between parts variation and measurement error variation. Another is the intraclass correlation, $\rho = \sigma_P^2/(\sigma_P^2 + \sigma_M^2) = \gamma/(1 + \gamma)$. ρ is the fraction contributed by parts variation to the variance of measurements, $\sigma_Y^2 = \sigma_P^2 + \sigma_M^2$. Finally, there is the signal-to-noise ratio (SNR), $SNR = \sqrt{2\gamma}$, which the automotive industry action group (AIAG) Measurement System Analysis Manual [2] defines SNR as the number of distinct categories that can be reliably distinguished by the measurement system. As a rule of thumb, SNR greater than 5 suggests a capable measurement system while SNR less than 2 suggests an incapable measurement system. These three measures (γ , ρ , SNR) are mathematically equivalent and therefore estimates for ρ and SNR are easily derived from that for γ .

Misclassification Probabilities

Besides these three measures, the capability of a measurement system can also be appraised by its

2 Gauge R&R Studies, Confidence Intervals for

ability to correctly or incorrectly classify good and bad parts. This is true regardless of whether the measurement system gives variables (quantitative) or attribute (pass/fail) readings. Misclassification of a good part as bad (false failure) results in wastage of resources through scrapping or reworking of good parts (*see Gauge Repeatability and Reproducibility (R&R) Studies, Misclassification Rates; Measurement Systems Analysis, Capability Measures for*). Misclassification of a bad part as good (missed fault) results in customer dissatisfaction. Therefore, a capable measurement system should have small probabilities of misclassification, where the false failure probability is defined as $\delta = \Pr(\text{a good part is classified as bad})$ and the missed fault probability is defined as $\beta = \Pr(\text{a bad part is classified as good})$. These probabilities are respectively defined by integrating the joint density of Y and P over the appropriate regions A and B which define the false failure and missed fault, that is, $\delta = \iint_A f(y, p) dy dp$ and

$$\beta = \iint_B f(y, p) dy dp \text{ with}$$

$$A = \{(y, p) : y \notin (LSL, USL), p \in (LSL, USL)\}$$

and

$$B = \{(y, p) : y \in (LSL, USL), p \notin (LSL, USL)\}.$$

Classical Gauge R&R Study

Consider the classical gauge R&R design, also known as a *balanced two-factor crossed random model with interaction*. The typical set up consists of p randomly chosen parts independently measured r times

by each of o operators. The k th measurement of the i th part by the j th operator is denoted by $Y_{ijk} = \mu_Y + P_i + O_j + PO_{ij} + E_{ijk}$ where μ_Y is a constant, $P_i, O_j, PO_{ij}, E_{ijk}$ are jointly independent, normally distributed with means as zero and variances $\sigma_P^2, \sigma_O^2, \sigma_{PO}^2$, and σ_E^2 respectively. These variances are commonly referred to as **variance components**. In terms of R&R variances, σ_E^2 is repeatability variance, and $(\sigma_O^2 + \sigma_{PO}^2)$ is the reproducibility variance. Collectively, $(\sigma_O^2 + \sigma_{PO}^2 + \sigma_E^2)$ is the measurement error variance. Estimates and CIs for R&R measures are based entirely on the analysis of variance (ANOVA) approach. The ANOVA table for the classical gauge R&R model is shown in Table 1. The point estimates are derived by first expressing the variances as linear functions of the expected mean squares (EMS) in column 2 of Table 2 and replacing each θ by its mean squares (S^2) in column 3 of Table 2. Negative estimates for $\hat{\gamma}$ and $\hat{\sigma}_P^2$ are assigned as zero.

Confidence Intervals for R&R Measures

In the classical R&R study and other R&R study designs based on the ANOVA approach, the relevant R&R measures such as σ_M^2 and γ can be expressed as functions of variance components. These measures are increasingly being studied using CIs, which incorporate the statistical uncertainty associated with the **estimation** procedure. An exact $100(1 - \alpha)\%$ CI for a measure is a random interval (L, U) based on the measurements such that there is a $(1 - \alpha)$ probability that the measure lies within it. Typical values for $(1 - \alpha)$ are 0.90, 0.95, and 0.99. In the absence of exact CI, an approximate CI is considered. An approximate $100(1 - \alpha)\%$ CI has

Table 1 ANOVA table for classical gauge R&R model

Sources of variation	Degrees of freedom	Mean squares ^(a)	Expected mean squares (EMS)
Parts O	$p - 1$	S_P^2	$\theta_P = or\sigma_P^2 + r\sigma_{PO}^2 + \sigma_E^2$
Operators P	$o - 1$	S_O^2	$\theta_O = pr\sigma_P^2 + r\sigma_{PO}^2 + \sigma_E^2$
Interaction PO	$(p - 1)(o - 1)$	S_{PO}^2	$\theta_{PO} = r\sigma_{PO}^2 + \sigma_E^2$
Replicates	$po(r - 1)$	S_E^2	$\theta_E = \sigma_E^2$

^(a)Mean squares are defined using: $S_P^2 = or \sum_i (\bar{Y}_{i**} - \bar{Y}_{***})^2 / (p - 1)$, $S_O^2 = pr \sum_j (\bar{Y}_{*j*} - \bar{Y}_{***})^2 / (o - 1)$,

$$S_{PO}^2 = r \sum_i \sum_j (\bar{Y}_{ij*} - \bar{Y}_{i**} - \bar{Y}_{*j*} + \bar{Y}_{***})^2 / (p - 1)(o - 1), \quad S_E^2 = \sum_i \sum_j \sum_k (Y_{ijk} - \bar{Y}_{ij*})^2 / po(r - 1),$$

$$\bar{Y}_{***} = \sum_i \sum_j \sum_k Y_{ijk} / por, \quad \bar{Y}_{ij*} = \sum_k Y_{ijk} / r, \quad \bar{Y}_{i**} = \sum_j \sum_k Y_{ijk} / or, \quad \bar{Y}_{*j*} = \sum_i \sum_k Y_{ijk} / pr.$$

Table 2 Estimation of R&R measures

Measure	In terms of EMS	Estimate
$\sigma_{\text{repeatability}}^2 = \sigma_E^2$	θ_E	$\hat{\sigma}_{\text{repeatability}}^2 = S_E^2$
$\sigma_{\text{reproducibility}}^2 = \sigma_O^2 + \sigma_{PO}^2$	$(\theta_O + (p-1)\theta_{PO} - p\theta_E)/pr$	$\hat{\sigma}_{\text{reproducibility}}^2 = \max\{0, (S_O^2 + (p-1)S_{PO}^2 - pS_E^2)/pr\}$
$\sigma_M^2 = \sigma_O^2 + \sigma_{PO}^2 + \sigma_E^2$	$(\theta_O + (p-1)\theta_{PO} + p(r-1)\theta_E)/pr$	$\hat{\sigma}_M^2 = \frac{S_O^2 + (p-1)S_{PO}^2 + p(r-1)S_E^2}{pr}$
σ_P^2	$(\theta_P - \theta_{PO})/or$	$\hat{\sigma}_P^2 = \max\{0, (S_P^2 - S_{PO}^2)/or\}$
$\gamma = \sigma_P^2/\sigma_M^2$	$\frac{p\theta_P - p\theta_{PO}}{o\theta_O + (p-1)o\theta_{PO} + po(r-1)\theta_E}$	$\hat{\gamma} = \hat{\sigma}_P^2/\hat{\sigma}_M^2$

approximate $(1 - \alpha)$ probability that the measure lies within it. Except for the repeatability variance, which has exact χ^2 CI, the other R&R measures do not have exact CIs. Several approximate CI methods for variance components function have been recommended for R&R studies, including for instance, the Satterthwaite's [3–5], asymptotic intervals based on restricted maximum-likelihood (REML) estimates [5], modified large sample (MLS), and generalized confidence intervals (GCIs) [6, 7]. The Satterthwaite's and REML intervals have been shown to be sometimes inadequate in coverage probabilities [5, 8] in small and moderate samples. Currently, only the MLS and GCI approaches are recommended as they have been shown through extensive simulation studies to provide coverage probabilities that are close to the stated confidence level even in small samples. Comprehensive discussion of these intervals and other analysis/design aspects for gauge R&R studies are found in Burdick *et al.* [9, 10].

Modified Large Sample Intervals

The MLS method was first introduced by Graybill and Wang [11]. Burdick and Larsen [5] recommends this MLS method for an approximate $100(1 - \alpha)\%$ CI for measurement error variance σ_M^2 :

$$(\hat{\sigma}^2 - \sqrt{V_L}/pr, \hat{\sigma}_M^2 + \sqrt{V_U}/pr) \text{ where } V_L = G_2^2 S_O^4 + G_3^2 (p-1)^2 S_{PO}^4 + G_4^2 p^2 (r-1)^2 S_E^4, \\ V_U = H_2^2 S_O^4 + H_3^2 (p-1)^2 S_{PO}^4 + H_4^2 p^2 (r-1)^2 S_E^4, \\ G_q = 1 - F_{\alpha/2; \infty, n_q}, H_q = F_{1-\alpha/2; \infty, n_q} - 1,$$

$q = 2, 3, 4, n_2 = o - 1, n_3 = (p-1)(o-1), n_4 = po(r-1)$, and $F_{\alpha; d1, d2}$ is the 100α th percentile of the F distribution with $d1$ and $d2$ **degrees of freedom**. As suggested by the given CI expression, MLS CIs have explicit but complicated expressions. In special cases where certain variance components are zero or as the dimensions o or p increase to infinity, the MLS approximate intervals reduce to exact intervals. The advantage of MLS intervals is that they are easily implemented in computer spreadsheets and always produce unique numerical values. More recommendations and a comprehensive discussion of MLS intervals in R&R studies are found in [5, 9, 10]. A complete catalogue of MLS methods in general ANOVA models is found in Burdick and Graybill [12].

Generalized Confidence Intervals

The GCI method is a computer-intensive method first suggested by Weerahandi [13] and first applied to R&R studies in Hamada and Weerahandi [6]. It uses the generalized pivotal quantity (GPQ) which is distributed free of the parameters of interest. Chiang [7] developed a similar method using surrogate variables which are comparable to the GPQ. Approximate intervals are constructed with **percentiles** of GPQ, using either numerical integration or simulation.

The procedure is a simple one using **Monte Carlo simulation**:

1. Define the GPQ for the EMS (θ) in Table 1:
 $\tilde{\theta}_P = (p-1)s_P^2/\chi_{(p-1)}^2, \tilde{\theta}_O = (o-1)s_O^2/\chi_{(o-1)}^2,$
 $\tilde{\theta}_{PO} = (p-1)(o-1)s_{PO}^2/\chi_{(p-1)(o-1)}^2,$
 $\tilde{\theta}_E = po(r-1)s_E^2/\chi_{po(r-1)}^2,$

4 Gauge R&R Studies, Confidence Intervals for

Table 3 GPQ for R&R measure

Measure	GPQ
$\sigma_{\text{repeatability}}^2 = \sigma_E^2$	$\tilde{\sigma}_{\text{repeatability}}^2 = \tilde{\theta}_E$
$\sigma_{\text{reproducibility}}^2 = \sigma_O^2 + \sigma_{PO}^2$	$\tilde{\sigma}_{\text{reproducibility}}^2 = \max\{0, (\tilde{\theta}_O + (p-1)\tilde{\theta}_{PO} - p\tilde{\theta}_E)/pr\}$
$\sigma_M^2 = \sigma_O^2 + \sigma_{PO}^2 + \sigma_E^2$	$\tilde{\sigma}_M^2 = \frac{\tilde{\theta}_O + (p-1)\tilde{\theta}_{PO} + p(r-1)\tilde{\theta}_E}{pr}$
σ_P^2	$\hat{\sigma}_P^2 = \max\{0, (\tilde{\theta}_P - \tilde{\theta}_{PO})/or\}$
$\gamma = \sigma_P^2/\sigma_M^2$	$\tilde{\gamma} = \tilde{\sigma}_P^2/\tilde{\sigma}_M^2$

where the χ_d^2 are independent χ^2 random variables and the lower case s^2 are realized values of the mean squares S^2 .

- For each R&R measure, using its definition (column 2 of Table 2), we replace each EMS by its GPQ in step 1, taking care to assign zero for negative values of GPQ. This produces the GPQ for R&R measures in column 2 of Table 3.
- Next we use simulation to obtain the 100($\alpha/2$)th and 100(1 - $\alpha/2$)th percentile of these GPQ as 100(1 - α)% GCI: simulate N independent values for each of $\chi_{(p-1)}^2$, $\chi_{(o-1)}^2$, $\chi_{(p-1)(o-1)}^2$, and $\chi_{po(r-1)}^2$. Use these to compute N values of each GPQ in Table 3. For each GPQ, we order the N values in ascending order. The $N(\alpha/2)$ th value and $N(1 - \alpha/2)$ th value in the ordered set of each GPQ is the required GCI.

Compared to MLS CI, the GCI procedure does not have explicit expression and requires computer simulation which will subject its estimates to simulation variation. Fortunately, this variation can be controlled by having a large enough value of N . It is recommended to use $N \geq 100000$. The attractiveness of GCI lies in its ability to provide CI where MLS methods are not available due to either complex functions or complex R&R study designs. For instance, Burdick *et al.* [14] propose CI for misclassification probabilities. In other situations, the GCI method can be easily applied to complex gauge R&R studies involving three factors with a nested structure such as that described by Burdick and Larsen [5].

Extension to Mixed Effects and Unbalanced Data Models

We have considered only the R&R studies where the operators and parts are regarded as random

effects as they have been randomly selected from the population at large. Daniels *et al.* [15] consider R&R studies with operators as fixed effects as they have been selected from a relatively small population. They proposed MLS and GCI intervals and showed by simulation studies that the GCI intervals are to be recommended over the MLS as they are sometimes much shorter. Another extension from our discussion concerns the balance/unbalance data nature of R&R study. A balanced data set has equal replications for each part by operator combination. Sometimes, due to practical reasons such as gauge malfunction during a study, unequal replications can occur resulting in unbalanced data. Consequently, the MLS and GCI methods described earlier, which assume balanced ANOVA and weighted **sum of squares**, cannot be applied. Instead, unweighted sum of squares (USS) are employed for the MLS and GCI approach. Gong *et al.* [16] consider CI for unbalanced two-factor gauge R&R studies.

A summary of earlier MLS related methods for mixed effects/unbalanced data models is found in Burdick and Graybill [12] while a comprehensive discussion and examples using MLS and GCI can be found in Burdick *et al.* [10].

References

- [1] Montgomery, D.C. & Runger, G.C. (1993). Gauge capability and designed experiments. Parts I: basic methods, *Quality Engineering* **6**, 115–135.
- [2] Automotive Industry Action Group. (2002). *Measurement Systems Analysis*, AIAG, Detroit.
- [3] Satterthwaite, F.E. (1946). An approximate distribution of estimates of variance components, *Biometrics Bulletin* **2**, 110–114.
- [4] Satterthwaite, F.E. (1941). Synthesis of variance, *Psychometrika* **6**, 309–316.
- [5] Burdick, R.K. & Larsen, G.A. (1997). Confidence intervals on measures of variability in R&R studies, *Journal of Quality Technology* **29**, 261–273.

- [6] Hamada, M. & Weerahandi, S. (2000). Measurement systems assessment via generalized inference, *Journal of Quality Technology* **32**, 241–253.
- [7] Chiang, A.K.L. (2001). A simple general method for constructing confidence intervals for functions of variance components, *Technometrics* **43**, 356–367.
- [8] Chiang, A.K.L. (2002). Improved confidence intervals for a ratio in an R&R study, *Communications in Statistics-Simulation* **31**, 356–344.
- [9] Burdick, R.K., Borror, C.M. & Montgomery, D.C. (2003). A review of methods for measurement systems capability analysis, *Journal of Quality Technology* **35**, 342–354.
- [10] Burdick, R.K., Borror, C.M. & Montgomery, D.C. (2005). *Design and Analysis of Gauge R&R Studies: Making Decision with Confidence Intervals in Random and Mixed ANOVA Models*; ASA-SIAM Series on Statistics and Applied Probability, SIAM, Philadelphia, ASA, Alexandria.
- [11] Graybill, F.A. & Wang, C.M. (1980). Confidence intervals on nonnegative linear combinations of variances, *Journal of the American Statistical Association* **75**, 869–873.
- [12] Burdick, R.K. & Graybill, F.A. (1992). *Confidence Intervals on Variance Components*, Marcel Dekker, New York.
- [13] Weerahandi, S. (1993). Generalized confidence intervals, *Journal of the American Statistical Association* **88**, 899–905.
- [14] Burdick, R.K., Park, Y.-J., Montgomery, D.C. & Borror, C.M. (2005). Confidence intervals for misclassification rates in a gauge R&R study, *Journal of Quality Technology* **37**, 294–303.
- [15] Daniels, L., Burdick, R.K. & Quiroz, J. (2005). Confidence intervals in a gauge R&R study with fixed operators, *Journal of Quality Technology* **37**, 179–185.
- [16] Gong, L., Burdick, R.K. & Quiroz, J. (2005). Confidence intervals for unbalanced two-factor gauge R&R studies, *Quality and Reliability Engineering International* **21**, 727–741.

ANDY KOK LEONG CHIANG

Measurement Error and Uncertainty

Introduction

Measurement systems analysis (MSA) involves understanding the measurement process in place, identifying and estimating the error present in the process or system, and determining the adequacy of the **measurement system** for use in product and **process control** (see **Engineering Process Control**). A measurement system itself consists of devices to obtain measurements (gauges), standards, methods, personnel, and so on. That is, it is the entire process for obtaining measurements on some quantity of interest (*measurand*). *Error* and *uncertainty* will be a part of any measurement system or in the results from any measurement system put into place. Unfortunately, the concepts of error and uncertainty are often confused with one another and sometimes used (incorrectly) interchangeably. There are two (if not more) schools of thought with respect to the operational definitions of error and uncertainty. In this chapter, measurement error, uncertainty, and their components are presented as they are often used and defined in the literature [1–6]. In addition, definitions and relationships for precision and accuracy with respect to measurement systems are provided.

Error

Error is by definition the difference between the measured value of the quantity of interest (*measurand*) and the true value of the measurand. In the literature, error has also been referred to and written as *measurement error*, *error of measurement*, *experimental error*, and so forth; although depending on the experimental situation, these terms may have specific and therefore, quite different meanings. In addition, issue has sometimes been taken with the use of the term *true* value of the measurand (see [3, 4, 7] for example), since a true value is unknowable. “Reference value” is often used interchangeably with true value although this use is not recommended [8]. The reference value should be considered more as the best *approximation* available for the true value.

There are numerous types of error. These may include random error, systematic error, observer

error, repeatability error, discrimination error, total error, and rounding error, just to name a few (see [6]). In this section, two important and common types of error are presented as well as their use in MSA: systematic error and random error.

Systematic Error

Systematic errors can be caused by human activities, inadequate manufacturing methods, imperfections of the measuring device, changes in temperature, pressure, and so on. This error remains the same over repeated measures taken under identical conditions. That is, the error that is induced for one measurement will be identical to the error for the next measurement. The error is *systematic* and results in values that are consistently above or consistently below the true or reference value of the measurand. Systematic error can also be defined as the difference between error and random error (discussed in the next subsection).

Random Error

Suppose that repeated measures are taken on the same item, under identical conditions (same operator, same gauge). Under ideal conditions, even when systematic errors have been identified and accounted for, error will occur. *Random errors* vary randomly over all measurements taken under identical conditions. Even in a well-designed experiment or well-designed MSA, normal random fluctuations will occur and often occur according to some distribution (usually a Gaussian distribution with zero mean). If only random errors exist in the system, then increasing the number of measurements taken will provide a better estimate or better approximation of the measurand’s true value. For example, if the average of the measurements taken is the best estimate for the true value of the measurand, then increasing the number of measurements will result in a better estimate.

Measurement Results Modeled by Systematic and Random Errors

Suppose that the following could represent a measurement result

$$Y = X + \varepsilon \quad (1)$$

2 Measurement Error and Uncertainty

where Y is the measurement result, X is the true value of the measurand, and ε represents the combination of both systematic and random error. If the error is combined to represent both the systematic and random error, and we let ν represent the mean of ε and σ^2 the variance, then it can be said that $\varepsilon \sim (\nu, \sigma^2)$. In this form, if $\nu = 0$ then the measurements are considered more *accurate* and there is no “systematic error”. Furthermore, if σ is sufficiently small, the measurements are considered more *precise*. (For more details, the reader is encouraged to see [9]). Precision and accuracy play a very important role in MSA and are discussed in the next subsection.

Measurement System Error

Measurement system error as it pertains to MSA (see **Measurement Systems Analysis, Overview**) is a combination of variability attributable to gauge *bias*, *stability*, *linearity*, *repeatability*, and *reproducibility* (see **Gauge Repeatability and Reproducibility (R&R) Studies; Gauge Repeatability and Reproducibility (R&R) Studies, Variance Components in; Gauge Repeatability and Reproducibility (R&R) Studies, Misclassification Rates; Gauge Repeatability and Reproducibility (R&R) Studies, Confidence Intervals for**). Bias, linearity, and stability are the components that make up the *accuracy* of a measurement system. Repeatability and reproducibility are the components that make up the *precision* (measurement variation).

Accuracy. *Accuracy* is defined as the difference between the measurement and the value of the measurand. Accuracy is a qualitative concept and, as a result, is not quantifiable. It does not make sense to describe a measurement system in these terms by saying, for instance, “the accuracy is 0.05 cm”. *Uncertainty* (discussed in the section titled “Uncertainty”), however, is quantifiable and can be used to describe the system: “The standard uncertainty is 0.05 cm.” In measurement system error, the three components of accuracy and their definitions are

- *Bias* – the difference between the observed average measurement and a reference or master value. It is a measure of the systematic error with respect to the measurement system.
- *Linearity* – measures how changes in the size of the measurand (part being measured) will affect

measurement system bias over the expected range of the process.

- *Stability* – measures how well the measurement system performs over time. It reflects the change in bias (if any) over time when the same parts or master value are measured.

In general, accuracy focuses on measures of location or the relationship between the measurement results and value of the measurand.

Precision. *Precision* is defined as the variation identified when the same part or item is measured repeatedly using the same measurement device or gauge. The components of precision are listed below (see **Measurement Systems Analysis, Overview; Gauge Repeatability and Reproducibility (R&R), Variance Components in; Gauge Repeatability and Reproducibility (R&R) Studies, Misclassification Rates; Gauge Repeatability and Reproducibility (R&R) Studies, Confidence Intervals for**).

- *Repeatability* – variability due to the gauge or test instrument when used by the same operator to measure the same unit. In other words, this is variability arising from the same operator (or setup) measuring the same part repeatedly using the same measuring device (gauge).
- *Reproducibility* – variability due to different operators or setups measuring the same parts using the same measuring device (gauge). This is considered the variability due to the measurement system.

Summary of Error

Chunovkina [3] among others argues that error is a model concept and as a result cannot be used to directly calculate error. Arguments have been made that error differs conceptually from uncertainty and as a result these terms should be used cautiously. Full and excellent discussions on measurement error can be found in [3, 4, 6, 9]. Measurement system error is discussed in [7].

In the next section, measurement uncertainty is presented. The discussion is based on several articles that discuss the *Guide to the Expression of Uncertainty in Measurement* (1993) also known as *GUM*, as well as the guide itself. In addition, the reader is encouraged to see *Guidelines for Evaluating*

and Expressing the Uncertainty of NIST Measurement Results [2].

Uncertainty

The International Committee for Weights and Measures (CIPM) in conjunction with several international organizations, including the International Engineering Consortium (IEC) and International Organization for Standards (ISO) began work on a project that would hopefully result in some common approach for describing the quality of measurement results among international laboratories. The result of this project is the GUM published in 1993. GUM attempts to provide a conceptual framework for assessing and reporting uncertainty in measurements. In this section, uncertainty and its extensions are presented. The reader is encouraged to see GUM as well as [2–5, 10] for full details and discussion of uncertainty as well as some critiques of GUM itself [3].

Standard Uncertainty

As stated previously, MSA as a whole is used to understand the measurement system and attempt to improve it through variation reduction techniques. *Standard uncertainty* on the other hand attempts to *quantify* the range of measurement values by providing intervals over which the measurement results would vary with a certain probability (or in some instances confidence levels). The expectation is the value of the measurand would be captured by this interval. Technically, uncertainties are characterized by standard deviations (*see Variance*).

The purpose of the GUM (1993) was to provide comprehensive and consistent information on how to report measurement uncertainty. GUM also provided a common framework for comparison of measurement results internationally. GUM states that the uncertainty of a measurement is a “... parameter, associated with the result of a measurement that characterizes the dispersion of the values that could reasonably be attributed to the measurand”. Again, standard uncertainty, u_i , is described by standard deviation. Obviously, u_i is also the positive square root of u_i^2 . Components that make up uncertainty (to be discussed in this section) are categorized into one of two categories (this is the categorization

presented by [1]). Grouping of components depends on the method used to estimate numerical values. There are two types of uncertainties, type A and type B:

- *Type A* – based on statistical analysis of a series of measurements.
- *Type B* – based on other sources of information. These “other” sources can include, for example, gauge manufacturer’s specification or published reference values.

The definitions given here are very simplistic and the reader is encouraged to see [1, 2]. In the remainder of this section, more basic definitions and descriptions for uncertainty are provided.

Combined Standard Uncertainty

Combined standard uncertainty, u_c , is the standard deviation of the combined errors. These errors include both random error and systematic error. The combined uncertainty, u_c can often be obtained using the *law of propagation of uncertainty* when the measurand, Y is a function of several independent variables. For more details, see [9] or [2, Appendix A].

In MSA for example, the combined uncertainty would be a combination of all sources of measurement variation for that particular experimental situation. This would include measurement variation in the measurement process, **calibration** issues, environment errors, and so on. (See [7] for more discussion of measurement uncertainty in MSA.) To illustrate, suppose that the sources of variation for a particular process include repeatability, reproducibility, and stability denoted by $\sigma_{\text{repeatability}}^2$, $\sigma_{\text{reproducibility}}^2$, and $\sigma_{\text{stability}}^2$, respectively. The combined uncertainty would then be

$$u_c = \sqrt{\sigma_{\text{repeatability}}^2 + \sigma_{\text{reproducibility}}^2 + \sigma_{\text{stability}}^2} \quad (2)$$

There can be many situations where the combined uncertainty can be quite complicated [2, 10] and the law of propagation of uncertainty can be very useful.

Expanded Uncertainty

Often in measurement studies it is not enough to simply report a single combined estimate of uncertainty. An interval about the measurement results that

defines the measurement uncertainty is often required. *Expanded uncertainty*, U , is the combined standard uncertainty, u_c multiplied by some “coverage factor”, k . In other words, $U = ku_c$. The coverage factor, k , is chosen such that a certain level of confidence (see **Confidence Intervals**) (for Gaussian or normally distributed measurements in a measurement system) is attained. In all instances, the factor k results in a certain coverage interval. When the normal distribution applies (see **Normal Distribution; Normality Tests**), $k = 2$ will result in an approximate 95% **confidence interval**. Suppose a measurand, Y , is estimated by some value y . The combined uncertainty would be given by $u_c(y)$, then $U = ku_c(y)$ and the expanded uncertainty would be written as $Y = y \pm U$. The expanded uncertainty provides more information about the process than a single-number statistic.

Other Considerations and Standards

To attain some level of standardization and agreement among laboratories internationally, GUM (1993) has provided guidance on how to evaluate and express measurement uncertainty. The guide is a high-level reference document with some very broad and sweeping specifications and definitions. Some very basic definitions are given in this section and the reader is encouraged to review [1–5] for more complete discussions and illustrations of the concepts outlined here.

Summary and Conclusions

In this article, some very basic definitions and concepts concerning measurement error and uncertainty are provided. There remains a great deal of room in the literature for further discussion and clarification; not only for the benefit of the measurement community, but for any practitioner or entity that must rely on measurement results, MSA, and the inferences that are to be made from the results obtained.

It is also important to note that many of the terms and concepts presented here and in the literature may

take on different meanings and interpretation based on the experimental situation at hand. For instance, a source of random error in one experimental setup may easily be a systematic error in a different setup. The roles of error and uncertainty may be quite different depending on the type of measurements actually recorded, and their use in practice. Again, the reader is encouraged to review the references provided here and the references provided within those documents.

References

- [1] ISO. (1993). *Guide to the Expression of Uncertainty in Measurement*, 1st Edition, International Organization of Standards, Switzerland.
- [2] Taylor, B.N. & Kuyatt, C.E. (1994). Guidelines for Evaluating and Expressing the Uncertainty of NIST Measurement Results, NIST Technical Note 1297.
- [3] Chunovkina, A.G. (2000). Measurement error, measurement uncertainty, and measurand uncertainty, *Measurement Techniques* **43**, 581–586.
- [4] Mari, L.P. (2005). *Explanation of Key Error and Uncertainty Concepts and Terms. Handbook of Measuring System Design*, P.H. Sydenham & R. Thorn, eds, John Wiley & Sons.
- [5] Cox, M. & Harris, P. (2005). An outline of supplement 1 to the guide to the expression of uncertainty in measurement on numerical methods for the propagation of distributions, *Measurement Techniques* **48**, 336–345.
- [6] McGhee, J. (2005). *Calculation and Treatment of Errors. Handbook of Measuring System Design*, P.H. Sydenham & R. Thorn, eds, John Wiley & Sons.
- [7] Automotive Industry Action Group. *Measurement Systems Analysis Reference Manual*, 3rd Edition, Automotive Industry Action Group, Southfield.
- [8] ASTM. E177-90a (2002). standard practice for use of the terms precision and bias in ASTM test methods, *Book of Standards 14.02.*, ASTM International West Conshohocken.
- [9] Gertsbakh, I. (2003). *Measurement Theory for Engineers*, Springer-Verlag, Berlin Heidelberg.
- [10] van der Veen, A.M.H. & Cox, M.G. (2003). Error analysis in the evaluation of measurement uncertainty, *Metrologia* **40**, 42–50.

CONNIE M. BORROR

Measurement Systems Analysis, Capability Measures for

Introduction

There are two “families” of measurement capability metrics in common use. There are those that compare measurement variation to total or part variation, and there are those that compare to the tolerance or specification width. These capability metrics are all functions of the **variance components** in the statistical model used to describe the response of interest. In addition to these capability metrics, it is also often useful to estimate the misclassification rates that arise from using the measurement system to discriminate between “good” and “bad” items being tested. An additional capability metric is obtained from the ratio of the estimated misclassification rates to the rates that would occur purely due to random chance. This section will define several capability measures, how they are related to one another, and give some common criteria for capability. Additional topics discuss the effect of the unit of measure and present a situation where none of the aforementioned metrics are particularly useful in quantifying the capability of a measurement process.

Standard MSA

The standard experiment for conducting a measurement systems analysis (**MSA**) uses a two-factor design with “parts” and “operators”. The “parts” can be almost anything for which a measurement is desired – people on whom a medical test is performed, metal rods for which the length is desired, or cell phones to determine if the functionality is working properly. “Operators” can be the people doing the measurement or can be different gauges or measurement systems. The statistical model used to describe the response is the random two-factor model

$$Y_{ijk} = \mu + P_i + O_j + PO_{ij} + E_{ijk} \quad (1)$$

$$i = 1, \dots, p; \quad j = 1, \dots, o; \quad k = 1, \dots, r$$

where μ is a constant, and P_i , O_j , PO_{ij} , and E_{ijk} are jointly independent normal random variables with means of zero and variances σ_P^2 , σ_O^2 , σ_{PO}^2 , and σ_E^2 , respectively. These variance components are used to compute the measurement capability measures. A number of authors discuss how to estimate variance components in the random two-factor model (see, e.g. [1]).

Commonly Used Capability Measures

The primary purpose of MSA is to determine whether the measurement process is adequate for its intended use. This is typically done by comparing the estimated measurement variation to either the part or total variation or to the width of any applicable tolerances or specifications.

Table 1 provides parameters of interest for assessing the capability of a measurement process. Table 2 presents some commonly used capability metrics. Perhaps the most common capability metric is the precision-to-tolerance ratio (PTR). The PTR is defined as:

$$PTR = \frac{k\sqrt{\gamma_M}}{USL - LSL} \quad (2)$$

where USL and LSL are upper and lower specification or tolerance limits and k is either 5.15 or 6. The value of $k = 6$ corresponds to the middle 99.73% of a normal probability distribution and $k = 5.15$ corresponds to the middle 99%. PTR compares the total measurement variation to the width of the specifications or tolerances the items being measured are to adhere to. The Automotive Industry Action Group (AIAG) MSA manual [2] recommends the following criteria be used for PTR:

$PTR \leq 0.1$: the measurement process is capable;
 $0.1 < PTR \leq 0.3$: the process is marginal;
 $PTR > 0.3$: the process is not capable.

It is further recommended that a confidence level be placed on the PTR to reflect the uncertainty associated with estimating variance components. For example, for the measurement process to be capable, the PTR should be less than 0.1 with at least a 90% level of confidence. Burdick *et al.* [3] give methods for computing **confidence intervals** on a variety of MSA metrics.

It has been noted in several places (see, e.g. [4]) that the PTR does not necessarily provide a

2 Measurement Systems Analysis, Capability Measures for

Table 1 Parameters of interest in MSA

Parameter	Variance component(s)	Definition
γ_P	σ_P^2	Variance among parts
γ_M	$\sigma_O^2 + \sigma_{PO}^2 + \sigma_E^2$	Total measurement variation
$\gamma_T = \gamma_P + \gamma_M$	$\sigma_P^2 + \sigma_O^2 + \sigma_{PO}^2 + \sigma_E^2$	Total variation
$\gamma_R = \gamma_P / \gamma_M$	$\sigma_P^2 / (\sigma_O^2 + \sigma_{PO}^2 + \sigma_E^2)$	Ratio of part to measurement variation
$\rho_P = \gamma_P / \gamma_T$	$\sigma_P^2 / (\sigma_P^2 + \sigma_O^2 + \sigma_{PO}^2 + \sigma_E^2)$	Ratio of part to total variation
$\rho_M = \gamma_M / \gamma_T$	$(\sigma_O^2 + \sigma_{PO}^2 + \sigma_E^2) / (\sigma_P^2 + \sigma_O^2 + \sigma_{PO}^2 + \sigma_E^2)$	Ratio of measurement to total variation

Table 2 Common capability metrics

Metric	Definition	Criteria
PTR	$\frac{k\sqrt{\gamma_M}}{USL - LSL}$	$k = 5.15$ or 6 $PTR \leq 0.1$: capable $0.1 < PTR \leq 0.3$: marginal $PTR > 0.3$: not capable
SNR	$\sqrt{2\gamma_R}$	$SNR > 10$: excellent $5 < SNR \leq 10$: acceptable $SNR < 5$: unacceptable
$PTV = \%R\&R$	$\sqrt{\rho_M} \times 100\%$	$PTV \leq 10\%$: excellent $10\% < PTV \leq 30\%$: acceptable $PTV > 30\%$: unacceptable
DR	$\sqrt{\frac{1 + \rho_P}{1 - \rho_P}}$	$DR > 4$: acceptable $2 < DR \leq 4$: marginal $DR < 2$: unacceptable

good indication of how well the measurement system performs for a particular process. This is due to the relationship between the capability of the production process and the capability of the measurement system used to monitor the production process. When the production process is highly capable of meeting the production limits, it can tolerate a higher PTR than a production process which is not capable. It is therefore useful to compute other measurement capability metrics that compare the measurement variation to total or part variation.

The signal-to-noise ratio (SNR) is defined as:

$$SNR = \sqrt{2\gamma_R} \quad (3)$$

The AIAG MSA manual defines SNR as the number of distinct data categories that can be reliably distinguished by the measurement system. A value of 5 or greater is recommended. Based on experience over a number of years, the author recommends the following criteria be met with at least a 90% level of confidence:

$SNR > 10$: the measurement process is excellent;
 $5 < SNR \leq 10$: the process is acceptable;
 $SNR < 5$: the process is unacceptable.

Another commonly used metric in the same class is the percent of total variation (PTV) or equivalently the percent of reproducibility and repeatability (%R&R). The PTV is defined as:

$$PTV = \sqrt{\rho_M} \times 100\% \quad (4)$$

AIAG recommends criteria of:

$PTV \leq 10\%$: the measurement process is excellent;
 $10\% < PTV \leq 30\%$: the process is acceptable;
 $PTV > 30\%$: the process is unacceptable.

Again, the author further recommends that the PTV be bounded by confidence intervals and the criteria be applied using at least a 90% confidence level.

Another metric in the same class is the discrimination ratio (DR) (see [5] or [6]). The DR is defined

as:

$$DR = \sqrt{\frac{1 + \rho_P}{1 - \rho_P}} \quad (5)$$

ρ_P is the intraclass **correlation** defined as the correlation between two measurements taken on the same part. Mader *et al.* [5] suggest the following criteria:

$DR > 4$: the measurement process is acceptable;
 $2 < DR \leq 4$: the process is marginal;
 $DR < 2$: the process is unacceptable.

An additional recommendation is that DR be bounded by confidence intervals and at least a 90% confidence level be used to judge the capability of the measurement process.

Interrelationships among Capability Metrics

$$PTV = \frac{100}{\sqrt{\frac{SNR^2}{2} + 1}} \quad SNR = \sqrt{2 \left(\frac{100^2}{PTV^2} - 1 \right)}$$

$$\rho_P = \frac{SNR^2}{SNR^2 + 2} \quad SNR = \sqrt{\frac{2\rho_P}{1 - \rho_P}}$$

$$DR = \sqrt{SNR^2 + 1} \quad SNR = \sqrt{DR^2 - 1}$$

Using Misclassification Rates to Determine Capability

In a “test” situation where the measurement system is used to discriminate among items such as a medical test for disease or a functional test for cell phones, Burdick *et al.* [3] have suggested a misclassification index for assessing capability. The misclassification rates or test errors are sometimes called *false failures* and *missed faults* or *false positives* and *false negatives*. The misclassification rates are directly analogous to the type I (α) and type II (β) errors in statistical **hypothesis testing**. Methods to estimate the misclassification rates have been described by several authors (see, e.g. [7]) and include numerical integration and computer simulation.

Suppose we wish to test a product to determine whether or not it meets acceptable limits.

In the case of two-sided test limits and specifications, the situation appears as in Figure 1. Following the notation in [7], define the following terms:

Y = true value $\sim N(\mu_P, \gamma_P)$
 X = measured value $= Y + E$
 E = measurement error $\sim N(0, \gamma_M)$
 $f(x, y)$ = joint pdf
 $f(x)$ = marginal pdf for X
 $f(y)$ = marginal pdf for Y

SL and SU are the lower and upper product specifications

TL and TU are the lower and upper test limits.

The test limits may or may not be equal to the product specifications.

The following probabilities are obtained by numerical integration or computer simulation.

$$P(\text{Bad}) = \int_{-\infty}^{SL} f(y) dy + \int_{SU}^{+\infty} f(y) dy = 1 - \pi \quad (6)$$

$$P(\text{Fail}) = \int_{-\infty}^{TL} f(x) dx + \int_{TU}^{+\infty} f(x) dx \quad (7)$$

$$P(\text{Pass, Bad}) = \text{Joint MF} \\ = \int_{-\infty}^{SL} \int_{TL}^{TU} f(x, y) dx dy$$

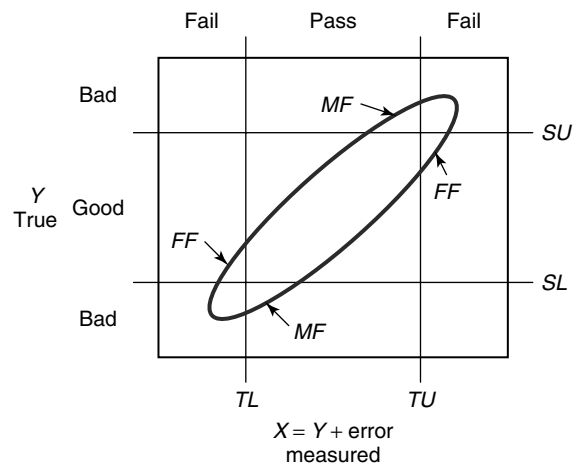


Figure 1 Misclassification rates with two-sided test limits and specifications

4 Measurement Systems Analysis, Capability Measures for

$$+ \int_{SU}^{+\infty} \int_{TL}^{TU} f(x, y) dx dy$$

$$= \beta \quad (8)$$

Both conditional and joint misclassification rates are obtained using Bayes formula and some algebra.

$$P(\text{Fail}|\text{Good}) = \text{FF}$$

$$= \frac{P(\text{Fail}) - P(\text{Bad}) + P(\text{Pass}, \text{Bad})}{1 - P(\text{Bad})}$$

$$= \alpha_c \quad (9)$$

$$P(\text{Pass}|\text{Bad}) = \text{MF} = \frac{P(\text{Pass}, \text{Bad})}{P(\text{Bad})} = \beta_c \quad (10)$$

$$P(\text{Fail}, \text{Good}) = \text{Joint FF}$$

$$= \text{FF}(1 - P(\text{Bad})) = \alpha \quad (11)$$

The joint misclassification rates, α and β are typically reported because they are relevant to the whole population of products.

When considering whether the misclassification rates are sufficiently small, Burdick *et al.* [3] have suggested comparing them to what could be obtained by chance. The proposal goes as follows. Suppose the value of π , $P(\text{Good})$, in equation 6 is known. A chance process for classifying parts is to take no measurements and randomly classify π of the parts as good and $1 - \pi$ as bad. For this chance process, the probability of a false failure is

$$\alpha^* = P(\text{Fail}, \text{Good})$$

$$= (1 - \pi)\pi \quad (12)$$

Likewise, the probability of a missed fault for this chance process is

$$\beta^* = P(\text{Pass}, \text{Bad})$$

$$= \pi(1 - \pi) \quad (13)$$

Thus, for a measurement system to be useful, the misclassification rates α and β should be less than what could be obtained by chance, α^* and β^* . A similar argument can be made for the conditional probabilities α_c and β_c . In this case, the conditional probabilities that could be obtained by chance are $\alpha_c^* = 1 - \pi$ and $\beta_c^* = \pi$.

Burdick *et al.* [3] proposed the following indices be used to assess the capability of the measurement system to discriminate good product from bad.

$$\alpha_{\text{index}} = \frac{\alpha}{\alpha^*} \quad (14)$$

$$\beta_{\text{index}} = \frac{\beta}{\beta^*} \quad (15)$$

And since $\alpha_c/\alpha_c^* = \alpha_{\text{index}}$ and $\beta_c/\beta_c^* = \beta_{\text{index}}$, these indices can be used for both the conditional and joint definitions of misclassification rates. The criteria for establishing the capability of the measurement process is for both α_{index} and β_{index} to be less than unity with at least a 90% confidence level. Burdick *et al.* [3] suggested using generalized confidence intervals and Larsen [7] proposed a different computer simulation method which provides similar results.

Effect of the Unit of Measure

Experience has shown that SNR and PTR are largely scale invariant so long as the data do not contain **outliers**. For example, decibel (db) is a common unit for electrical measurement. Decibel can be defined as the log of the ratio of two voltages. Whether the MSA is conducted directly in decibel units or with a linearizing transformation does not matter in so far as the capability measures are concerned so long as the data do not contain outliers. SNR and PTR values will be similar either way. Doing the analysis both with raw and transformed data and finding a large difference in either SNR or PTR often indicates the presence of outliers. And these outliers should be examined for validity and possible exclusion from the MSA.

A Situation Where the Usual Metrics Can Be Noninformative

It is increasingly the case in testing consumer electronic devices (cell phones, PDAs, etc.) that the measurement system is used to “align” the device rather than determine whether it performs within acceptable limits. With alignment, the value the device initially has is unimportant because the measurement system changes the value of the device to lie within a narrow range. When device parameters are aligned, the

resulting “part” variation is similar to that found in the measurement system. This implies that the SNR will tend to be well below the usual criteria of 5. Further, parameters which are aligned do not have test limits in the usual sense because the objective is not to determine pass or fail but rather to set the parameter to a predefined value or narrow range of values called *alignment limits*. The alignment limits are a trade-off between providing a sufficient margin with industry standard specifications and aligning to a nominal value that often takes excessive test time. The object of the alignment is not to center the parameter within the limits but just to get within them. Often there are no product specifications warranted to the end customer on aligned parameters of consumer electronic devices. This situation makes the PTR difficult to apply as well. Using the alignment limits in the denominator of the PTR typically provides values too high to meet the usual criteria. Finally, because the measurement system is not used to determine pass or fail, the misclassification indices are not applicable.

The question of how to assess the capability of a measurement system used to align parameters is difficult to answer. Larsen [8] makes several recommendations which include comparing the total measurement system variability found from an MSA to a tolerance stack of component tolerances in the measurement system. If the measurement variability from the MSA is less than the value from the tolerance stack, then the measurement process is at least meeting capability expectations. However, as noted, this is no guarantee that the measurement system is adequate in a particular application. It is also suggested that the measurement equipment be calibrated and traceable to National Institute of

Standards and Technology (NIST) so that system-to-system variation is lessened. Finally, the PTR might be applicable if the alignment limits are used with relaxed criteria or if product specifications are developed which reflect acceptable functional performance to the end customer.

References

- [1] Vardeman, S.B. & VanValkenburg, E.S. (1999). Two-way random-effects analyses and gauge R&R studies, *Technometrics* **41**(3), 202–211.
- [2] Automotive Industry Action Group. (1995). *Measurement Systems Analysis*, 2nd Edition, AIAG, Detroit.
- [3] Burdick, R.K., Borror, C.M. & Montgomery, D.C. (2005). *Design and Analysis of Gauge R&R Studies: Making Decisions with Confidence Intervals in Random and Mixed ANOVA Models*, SIAM-ASA Series on Statistics and Applied Probability, SIAM, Philadelphia, ASA, Alexandria.
- [4] Montgomery, D.C. & Runger, G.C. (1993). Gauge capability and designed experiments, part I: basic methods, *Quality Engineering* **6**(1), 115–135.
- [5] Mader, D.P., Prins, J. & Lampe, R.E. (1999). The economic impact of measurement error, *Quality Engineering* **11**(4), 563–574.
- [6] Wheeler, D.J. (1992). Problems with gauge R&R studies, in 46th ASQC Quality Congress Transactions, Nashville, pp. 179–185.
- [7] Larsen, G.A. (2003). Measurement system analysis in a production environment with multiple test parameters, *Quality Engineering* **16**(2), 297–306.
- [8] Larsen, G.A. (2002). Measurement system analysis: the usual metrics can be noninformative, *Quality Engineering* **15**(2), 293–298.

GREG A. LARSEN

Gauge Repeatability and Reproducibility (R&R) Studies, Destructive Testing

Gauge R&R Study

Measurement results can never simply be taken for granted. Often it's impossible to be absolutely sure, even if you are willing to investigate the measurement in depth. The practical solution for this dilemma is to examine how the measurements are obtained, and to release the procedure for future use if some well established conditions are satisfied. *Truth* and *precision* are the two aspects of measurement, dealing with the average and variation respectively. The standard method for assessing the precision of a **measurement system** is the *gauge repeatability and reproducibility study* (gauge R&R study, **Gauge Repeatability and Reproducibility (R&R) Studies**), for measurements on a numerical scale. The basis is a designed experiment with repeatedly measured objects, but the consequence is that the gauge R&R study is useless for so-called destructive measurements. Some measurement procedures are destructive in nature, e.g., when the strength of a product is obtained by recording the force needed to break it. Other measurements cannot be repeated because the measured object is constantly changing, e.g., the density of a flowing liquid in a pipeline. Both types of measurement are called *destructive*, even though the second example shows that this is not necessarily true for the measured objects.

The standard gauge R&R study uses the property that object-to-object variation can be distinguished from measurement variation, by repeating the measurements of a single object. For destructive measurements the two sources of variation are confounded. The only way to separate them is to make assumptions about the nature of the measurements, or the measured objects. For an investigation of strength measurements, you might carefully select a number of apparently equal products, and then assume that all variation in strength is measurement variation. We summarize all possible assumptions, and show how to analyze in each case.

Consequences for Destructive Measurements

The result of a gauge R&R study is an estimate of the measurement precision. If measurement variation is small compared to process variation or product tolerance width, then we decide that for our intentions the measurement is sufficiently precise. Destructive measurements usually lead to overestimation of the measurement precision. With the strength measurements above, for example, any product-to-product variation is wrongly attributed to the measurement precision. In some cases we might therefore falsely reject the measurements as too imprecise. Note that this protects us against undue optimism. The validity of the assumptions plays a crucial role, of course, whether or not measurement precision is estimated well.

Measurement

Measurement allows for mathematical study of interesting properties of objects or experimental units. A formal definition of a measurement of a property of a unit is the assignment of a numeral to that unit, reflecting the classification of the units according to the property under study. In this article we assume that the measurements use a numerical scale. But certainly in the context of destructive measurements we have to realize that objects may evolve in time. For this reason we choose to consider a single object at two different moments in time as two different experimental units. The density of the liquid in a pipeline at moment t is probably different from the density at t , when owing to the flow the liquid is not exactly the same or even completely different. This is reflected by the assumption that we are dealing with different experimental units. The mathematical definition of a measurement is the map $Y : U \rightarrow R$, with U consisting of elements u_{it} (the index i denotes experimental unit i and t refers to moment t) and R is the set (or a subset) of real numbers.

Assumptions for the Standard Gauge R&R Studies

The conditional probability distribution $F(Y|u)$ of Y given $u \in U$ defines the random measurement error. A measurement system is *accurate* when the conditional expectation of $F(Y|u)$ coincides with the true value $T(u)$. The difference is called the *bias*. Accuracy is investigated with **calibration** methods, and is

2 Gauge R&R Studies, Destructive Testing

separated from gauge R&R studies. Nevertheless we need to assume one important condition:

(a) The bias is constant for every $u \in U$ (linearity) and constant in time (stability).

The *precision* of a measurement system relates to the conditional standard deviation of $F(Y|u)$. Usually this variation is independent of the measured units. If this cannot be assumed, then a gauge R&R study yields nonsense, unless the relation between the variation and the object can explicitly be stated. Some measurements have different values for precision in different ranges, for example because different electrical components (such as resistors) are used. The solution is then to do separate gauge R&R studies for different ranges, while assuming a constant variation within one study. In this article we assume *homogeneity of measurement error*:

(b) $F(Y - \mu_{it}|u_{it}) = F(Y - \mu_{js}|u_{js})$ for all u_{it} and u_{js} (and μ is the conditional expectation).

This condition implies that the variation is independent of the measured unit.

Our formulation of measurements indicates a fundamental problem: for a fixed u_{it} normally no more than one realization of $Y(u_{it})$ is possible, prohibiting **estimation** of the variation of $F(Y|u_{it})$ without making additional assumptions. For standard gauge R&R studies two additional homogeneity assumptions are made. The first is the *temporal stability of objects*, meaning that it doesn't matter when the object is measured:

(c) For every object u_i the true value is constant in time: $T(u_{it}) = T(u_{is})$ in a relevant time interval.

The second assumption is *robustness of measurement*, meaning that the examined property of the object is not affected by the measurement:

(d) $T(u_{it})$ is equal before and after object u_i is measured.

Repeated measurements of object u_i can be used to estimate the variation of $F(Y|u_i)$ if assumptions (c) and (d) hold. For the standard gauge R&R study you normally select a few (say 10) representative objects, and pool the estimates of the variation for the separate objects (see **Gauge Repeatability and**

Reproducibility (R&R), Variance Components in for the details).

Assumptions for Gauge R&R Studies of Destructive Measurements

If at least one of the assumptions (c) or (d) does not hold, then the measurement is *destructive*. To resolve the fundamental problem of gauge capability studies, other – more stringent – assumptions are needed for estimating the measurement variation. These assumptions can be summarized according to Table 1.

Gauge R&R Studies for Destructive Measurements

In this section we identify a number of situations in which a satisfactory approach is available for conducting a gauge R&R study for destructive measurements.

Temporal Stability; Object Variation is Known

The first situation we consider is that assumption (c) holds, but assumption (d) does not: the objects u_i we want to measure do not change over a relevant time period except when they are measured. An estimate of the measurement variation requires therefore several objects, because each object can be measured only once. Unless the variation of the objects themselves is known, we still have no possibility to separate measurement variation from objects variation. Three possible solutions are given.

Table 1 Assumptions for Gauge R&R Studies

	Temporal stability (assumption (c) is true)	No temporal stability
Robust measurement (assumption (d) is true)	Standard gauge R&R study	Temporal variation can be modeled (j)
Nonrobust measurement	Objects variation known (e, f, g) Alternative measurements available (h, i)	True values known (k)

Homogeneity of Objects.

(e) There is a subset $H \subset U$ with $T(u_i) = T(u_j)$ for all u_i and u_j in H .

The first solution is applicable if you can select a number of objects with identical true value, regarding the interesting property. In a production process you may assume that products made from the same batch of raw materials, and in a small period of time, are quite alike. For the investigation of a strength test such products are obvious candidates to fulfill assumption (e). Under this condition the variation of the measurements of the objects $u_i \in H$ is the measurement precision.

Representativeness of Alternative Objects. Assumption (e) might be wishful thinking: some measured objects are inherently variable. For a gauge R&R study, which is not daily practice, we might use alternative objects instead, objects satisfying condition (e) and representative of the normal measurements. Generally speaking, this assumption is a solution when ordinary measured objects can favorably be replaced by better-defined objects. For a measurement capability study of the strength of wired diodes, it is better to replace the diodes by simple wires with strengths comparable to the diodes, because their homogeneity is much easier to control.

(f) There is a set of alternative objects $A = \{v_i\}_{i=1,\dots,k}$ for which the homogeneity condition (e) does hold, and with the same probability distribution F , i.e., $F(Y - \mu_u|u) = F(Y - \mu_v|v)$, for all $v \in A$ and any $u \in U$.

Under assumption (f) the variation of the measurements of the alternative objects is the measurement variation we are looking for.

Patterned Objects Variation. If there is no set of objects (or substitutes) for which homogeneity can be assumed, then in several cases there is a possibility to artificially create a homogeneous set by correcting for modeled heterogeneity.

(g) For all u_i in a subset $H \subset U$ the true value $T(u_i) = f(u_i)$, f is a known function with a small number of unknown parameters.

Condition (g) essentially says that the variation of the objects is not random, but that a (partly) known

pattern can be discerned. Some useful examples of the functions f are:

- $f(u_k) = \beta_0 + \beta_1 k$ for gradually increasing or decreasing successive objects $u_1, \dots, u_n$.
- $f(u_k) = B_i + P_j$ for u_k from batch i and position j .

With a suitable selection of objects u_i from the subset H the unknown parameters of the function f can be estimated from the measurements $Y(u_i)$, and from this the true value $T(u_i)$ can be inferred. Under condition (g) the variance of measurement error can be estimated as $(n - p)^{-1} \sum_{i=1}^n \{Y(u_i) - \widehat{T}(u_i)\}^2$ with $\widehat{T}(u_i)$ the fitted true value, and p the number of unknown parameters of the function f .

Temporal Stability; Alternative Measurements are Available

The second situation we consider applies when the destructive measurement is a cheap or operator-friendly alternative of either a nondestructive measurement, or destructive measurement of higher quality. Both cases are described for the situation that objects do not change over a relevant time period.

An Alternative Nondestructive Measurement. If a nondestructive alternative measurement is available (maybe at the laboratory, or at a dedicated institute), then the measurement variation of the destructive method can be determined without taking recourse to the methods of the previous paragraphs. The assumption is:

(h) An alternative nondestructive measurement exists.

The way to proceed is straightforward. Select a sample of n objects, and measure them with the alternative method for a standard gauge R&R study. Apart from an estimate of the measurement variation of the alternative method, this yields an estimate of the variation of the true values of the selected objects. Then measure the same objects destructively, with the method we are investigating. The variation of these measurements is the combined effect of the objects themselves and measurement error. The variance minus the variance

of the object's true values is an estimate of the variance of the destructive measurement method. Note that we do not assume that the alternative method yields perfect measurements: a constant bias is allowed.

An Alternative Perfect Destructive Measurement.

(i) An alternative destructive measurement X exists for which $X(u_i) = T(u_i)$ for all objects $u_i \in U$.

The meaning of condition (i) is that the method under study can be replaced with a perfect destructive alternative, i.e., the true value could be obtained from this alternative method. The case of a perfect *nondestructive* alternative is even simpler, and would in effect come down to the situation that the true values of the objects are known. In case of condition (i) a sample of objects from U is randomly split in two subgroups. One subgroup is measured with the alternative method, and the variance is an estimate of the variance of the true values of the other subgroup. The variance of the measurements of the other subgroup, using the method under study, minus the variance of the object's true values, is an estimate of the variance of this method's measurement error.

Robust Measurement; Temporal Variation is Known

The third situation we consider is that the act of measuring has no influence on the object, but the objects themselves are constantly changing. If we know how to model this, then variation due to objects can be separated from variation due to measurement. This assumption, labeled *patterned temporal variation*, is very similar to assumption (g).

(j) For all u_{it} in a subset $H \subset U$ the true value $T(u_{it}) = f(u_{it})$, with f a known function with a small number of unknown parameters.

Assumption (j) merely states that the temporal variation is not just random, but that it follows a certain pattern.

The simplest "pattern" is that some objects can be selected with a constant true value in some period of time. A tank with a completely homogeneous liquid could be pumped through a pipeline. In this

period the density may assumed to be constant, so all variation in the density measurements can be regarded as measurement variation. Another common model for some u_i in H is $f(u_{it}) = \beta_{0i} + \beta_{1i}t$.

The way to proceed is to obtain measurements $Y(u_{it(j)})$ of object u_i at k moments $t(1), \dots, t(k)$. From these measurements the unknown parameters of the function f are estimated, which in turn yield estimates of the true values $T(u_{it(j)})$. Under condition (j) the variance of measurement error can be estimated by $(k - p)^{-1} \sum_{j=1}^k \{Y(u_{it(j)}) - \widehat{T}(u_{it(j)})\}^2$

with $\widehat{T}(u_{it(j)})$ the fitted true value, and p the number of unknown parameters of the function f . For a better estimate this procedure should be repeated for several objects u_i , and the estimated variances are pooled.

True Values are Known

The final and most restrictive case is that the true values of (some) objects are known. This is the standard assumption of ordinary calibration studies. Any measurement procedure reveals its secrets when this assumption holds for a *representative* set of objects.

(k) A set H of objects u_{it} exists for which $T(u_{it})$ is known, $\forall u_{it} \in H$, and H is representative for U (i.e., the condition in (f) about the probability distribution F holds as well).

Assuming an unbiased measurement Y , the measurement variance is simply estimated by $k^{-1} \sum (Y(u_{it}) - T(u_{it}))^2$. A known bias β can be simply be corrected for by replacing $Y(u_{it})$ with $Y(u_{it}) - \beta$. And if the bias is unknown, then $k^{-1} \sum (Y(u_{it}) - T(u_{it})) = \bar{Y} - \bar{T}$ is an estimate of β , and the sample variance of the $Y(u_{it}) - T(u_{it})$ estimates the measurement error.

Evaluation

The fundamental problem of gauge capability study is that no object can be measured more than once. In many cases the problem can be solved, however, by the assumption that objects are constant in time (assumption (c) in this article). But the less severe assumption that temporal variation can be modeled (assumption (j)) may also be sufficient, at least, if measuring does not influence the object's true value

(assumption (d)). If measuring is *destructive* in the sense that assumption (d) is violated, then assumptions (e–i) or (k) could help to quantify the measurement error.

The success of the gauge R&R study depends on two conditions: the effort of doing a good experiment, and the validity of the assumptions. A bad experiment can lead to any false conclusion, but the effect of only partially true assumptions is one sided. A valid assumption is a correct model of the object-to-object variation, and variation due to measurement can be perfectly distinguished – apart from sampling error. If an assumption is not correct, however, then the estimated measurement variation will contain (part of) the object-to-object variation, and will therefore be higher than the true value.

If, for instance, assumption (e) – there is a subset of identical objects – is assumed, but in reality $T(u_i) = \mu + \varepsilon_i$ with ε_i independent drawings from the $N(0, \sigma_\varepsilon^2)$ distribution, then the measurement variance σ_m^2 is estimated by the sample variance s^2 of the measurements $Y(u_i)$. But the expected value $E\{s^2\} = \sigma_m^2 + \sigma_\varepsilon^2$ confounds the measurement variance and the object-to-object variance σ_ε^2 .

Further Reading

de Mast, J. & Trip, A. (2005). Gauge R&R studies for destructive measurements, *Journal of Quality Technology* **37**, 40–49.

ALBERT TRIP AND JEROEN DE MAST

Risk Analysis, Extremes in

General Background

In risk analysis, one often encounters situations where classical central limit theory does not extract all useful or even necessary information available in data. In such cases the extreme values in the sample might provide additional if not substantial insight into certain aspects of the data. The extremes in the sample deserve special attention, completing a thorough analysis of the available data. In specific cases, the extremes are even more vital than the central terms. For an early account on some of such issues, see the book by Gumbel [1]. More recent treatments that reflect the diversity of applications involved are given in Galambos *et al.* [2], Beirlant *et al.* [3], and Embrechts *et al.* [4].

In the section titled “Basic Extreme Value Theory” we will deal with the asymptotic theory for extremes coming from a classical statistical sample; this is by far the most encountered and best studied case. In the section titled “Statistical Aspects” we deal with some of the key statistical problems. The section titled “Discussion of Available Results” offers a discussion on adaptations and extensions of the offered material. In particular that section also refers to generalizations for which there already exists research literature. Applications are treated in **Applications of Extreme Statistics in Science and Engineering**.

Basic Extreme Value Theory

As explained, we first restrict ourselves to the case of independent and identically distributed (i.i.d.) random variables. In this way, the notational complexity can be kept to a minimum as it mimics that of traditional samples in statistics. Let $X_1, X_2, \dots, X_n$ be a sequence of independent random variables with a common distribution $F(x) = P\{X \leq x\}$ which is concentrated on an interval $-\infty \leq x_- \leq x_+ \leq \infty$ where x_+ (x_-) is the right (left) endpoint of the support of F .

With the above sequence we associate the sequence of *order statistics*

$$X_{n,1} \leq X_{n,2} \leq \dots \leq X_{n,n} \quad (1)$$

that will be crucial in what follows. To begin with, we focus our attention on the right tail behavior of F , hence to the maximum $X_{n,n}$.

The majority of properties of the extremes is controlled by the right tail behavior of F , expressed either in terms of its *quantile function* defined as the tail inverse of F by $Q(p) = \inf\{y : F(y) \geq p\}$ for $p \in (0, 1)$ or alternatively by the *tail (quantile) function* defined by $U(x) = Q(1 - \frac{1}{x})$ for $x > 1$.

De Haan’s treatise [5] lies at the origin of the revived interest in extreme value theory. There, the key challenge of extreme value theory is stated as the full solution of the *extremal domain of attraction problem*. By this we mean:

1. the derivation of all possible limit laws for $a_n^{-1}(X_{n,n} - b_n)$ where $\{a_n\}$ are positive norming constants while $\{b_n\}$ are real centering constants, and
2. the determination of the conditions on F that lead to one of the above limits.

The statement of this problem parallels that of the *central domain of attraction problem* where in the above the maximum $X_{n,n}$ is replaced by $S_n := \sum_{1 \leq i \leq n} X_i$. The solution to the first problem is offered by the class of *extreme value distributions* that can be written for $n \rightarrow \infty$ as

$$P\left(\frac{X_{n,n} - b_n}{a_n} \leq y\right) = F^n(a_n y + b_n) \longrightarrow \exp\left[-(1 + \gamma y)^{-1/\gamma}\right] =: G_\gamma(y) \quad (2)$$

for all y such that $1 + \gamma y > 0$. If Y_γ is a random variable with distribution G_γ then we write $a_n^{-1}(X_{n,n} - b_n) \xrightarrow{\mathcal{D}} Y_\gamma$. The parameter γ is called the *extreme value index* (EVI). It fully characterizes the distribution that clearly belongs to a one-parameter family of distributions.

The solution to the second problem depends on the value of the EVI: a distribution F belongs to the *domain of attraction* of the limit distribution G_γ (written in the form $F \in \mathcal{D}(G_\gamma)$) if and only if

$$\frac{\{U(xu) - U(x)\}}{g(x)} \longrightarrow h_\gamma(u) = \int_1^u v^{\gamma-1} dv \quad (3)$$

for all $u \geq 1$ as $x \rightarrow x_+$, for some *auxiliary function* $g > 0$. In general one can take $b_n = U(n)$ and $a_n = g(n)$. Equivalently, for some auxiliary function

2 Risk Analysis, Extremes in

$$b = g(1/(1 - F)),$$

$$\lim_{t \rightarrow x_+} \frac{1 - F(t + b(t)v)}{1 - F(t)} = (1 + \gamma v)^{-1/\gamma}, \text{ for all } v$$

with $1 + \gamma v > 0$ (4)

Unfortunately this streamlined result has not been at the roots of the development of extreme value theory. The one-parameter family $\{G_\gamma, -\infty < \gamma < \infty\}$ has been constructed as a reunification of three different types of extremal laws that have been separately discovered. Indeed, for the cases where $\gamma \neq 0$ one can make an additional affine transformation that simplifies the expression for the limit distribution at the cost of a loss of uniformity. More specifically

- Type I: $\gamma = 0$ is the most popular case and is called the *Gumbel case*. The domain $\mathcal{D}(G_0)$ contains most of the classical distributions like the normal, the lognormal, the gamma, and the Weibull distribution $1 - F(x) = \exp(-cx^\tau)$ with $\tau > 0$.
- Type II: $\gamma > 0$ leads to the *Pareto–Fréchet case* which is typical for heavy-tailed distributions. The prime example that actually characterizes $\mathcal{D}(G_\gamma)$ with $\gamma > 0$ is offered by the *Pareto-type distributions* where $1 - F(x) \sim x^{-1/\gamma} \ell(x)$ and ℓ is a slowly varying function. Often the extremal index γ is replaced by the *Pareto-index* defined by $\alpha = \gamma^{-1}$.
In this case the ratio $g(x)/U(x) \rightarrow \gamma$ so that also $a_n/b_n \rightarrow \gamma$. But then $X_{n,n}/U(n) \xrightarrow{\mathcal{D}} Z_\gamma$ and the limit law of Z_γ is slightly simpler than that of Y_γ .
- Type III: $\gamma < 0$ gives the *Weibull (extremal) case* which appears typically with distributions for which $x_+ = U(\infty) < \infty$ and hence can be interpreted as light tailed. Beta distributions (including the uniform) are among the most popular members of the corresponding domain of attraction. A simplification as for $\gamma > 0$ applies since in this case $g(x)/(x_+ - U(x)) \rightarrow -\gamma$.

With these alternatives we arrive at the so-called three types theorem. The determination of the limits is due to Fisher and Tippet [6] while that of the domains of attraction was derived by Gnedenko [7].

Fisher–Tippett–Gnedenko Theorem The extreme value distributions are exactly those that agree to within type with one of the following ($\alpha > 0$)

1. Pareto–Fréchet-type ($\alpha > 0$) or type II:

$$\Phi_\alpha(x) = \begin{cases} 0 & \text{if } x \leq 0 \\ e^{-x^{-\alpha}} & \text{if } x \geq 0 \end{cases} \quad (5)$$

Moreover $F \in \mathcal{D}(\Phi_\alpha)$ iff for $x \rightarrow \infty$

$$\frac{1 - F(\lambda x)}{1 - F(x)} \longrightarrow \lambda^{-\alpha}, \quad \forall \lambda > 0 \quad (6)$$

2. Gumbel-type or type I:

$$\Lambda(x) = e^{-e^{-x}}, \quad x \in \mathfrak{R} \quad (7)$$

Moreover $F \in \mathcal{D}(\Lambda)$ iff for some auxiliary function b

$$\frac{1 - F(x + tb(x))}{1 - F(x)} \rightarrow e^{-t} \quad (8)$$

3. Extremal Weibull-type ($\alpha > 0$) or type III:

$$\Psi_\alpha(x) = \begin{cases} e^{-|x|^\alpha} & \text{if } x \leq 0 \\ 1 & \text{if } x \geq 0 \end{cases} \quad (9)$$

Moreover $F \in \mathcal{D}(\Psi_\alpha)$ iff $x_+ < \infty$ and

$$\frac{1 - F(x_+ - \frac{1}{\lambda x})}{1 - F(x_+ - \frac{1}{x})} \rightarrow \lambda^{-\alpha}, \quad \forall \lambda > 0 \quad (10)$$

In reliability applications one often uses the *hazard function* that is defined by the ratio $r(x) := f(x)/(1 - F(x))$, where f is the density of F . For the three types there exist sufficient conditions in terms of this hazard function. We use the formulation of the Fisher–Tippett–Gnedenko theorem above.

von Mises Theorem Sufficient conditions for the extremal distributions are

1. for the Pareto–Fréchet-type: $x_+ = \infty$ and $x r(x) \rightarrow \alpha > 0$;
2. for the Gumbel-type: $r(x)$ is ultimately positive in the neighborhood of x_+ , is differentiable there and satisfies $r'(x) \rightarrow 0$;
3. for the Weibull-type: $x_+ < \infty$ and $(x_+ - x) r(x) \rightarrow \alpha > 0$.

The above results can be found in all treatments of extreme value analysis. We refer to Beirlant *et al.* [3], Coles [8], Embrechts *et al.* [4], Galambos [9], Leadbetter *et al.* [10], Reiss and Thomas [11], or Resnick

[12] where the reader can also find information on a number of statistical issues that are dealt with in the next section.

Statistical Aspects

The most important statistical questions related to extreme value theory are the **estimation** of the EVI γ , and the estimation of extreme **quantiles**. A further step would be to obtain confidence intervals for these quantities and to construct hypotheses tests. We briefly deal with these problems referring the reader to the aforementioned treatises for more information.

Estimation of the EVI

Two different approaches are available: either one looks at sample values that overshoot a certain threshold or one only uses a number of the largest order statistics.

- In the first case, the *peaks-over-threshold (POT) method*, a threshold u is chosen. The domain of attraction, formulated in terms of the distribution F shows that F is well approximated by the parameterized distribution

$$P\left\{\frac{X-u}{\sigma} > y | X > u\right\} \sim 1 - \left(1 + \frac{\gamma y}{\sigma}\right)^{-1/\gamma} \quad (11)$$

for all y values for which the right-hand side is feasible. The distribution on the right-hand side has been called the *generalized Pareto distribution*. The estimation of the parameter γ can be done by maximum-likelihood methods as shown by Smith [13, 14]. Alternative methods are the kernel methods of Csörgő [15], the method of probability-weighted moments of Hosking *et al.* [16], or the elemental percentile method of Castillo and Hadi [17].

- Other authors have advocated a nonparametric approach where γ is estimated using the order statistics of the sample.
 - If γ is known to be positive, then many estimation procedures exist. The most classical estimator

$$H_{k,n} := \frac{1}{k} \sum_{i=1}^k \{\log X_{n-i+1,n} - \log X_{n-k,n}\} \quad (12)$$

has been proposed by Hill [18] and is called *Hill's estimator*. This estimator and some of its kernel versions treated by Csörgő *et al.* [15] are fully covered in Beirlant *et al.* [3].

- In the general case, the choice for appropriate estimators is harder. Possible candidates are
 - * the estimator proposed by Pickands [19]

$$\frac{1}{\log 2} \log \frac{X_{n-k+1,n} - X_{n-2k+1,n}}{X_{n-2k+1,n} - X_{n-4k+1,n}} \quad (13)$$

or the refined versions by Drees [20] or Segers [21];

- * the *moment estimator* introduced by Dekkers *et al.* [22]

$$H_{k,n} + 1 - \frac{1}{2} \left(1 - \frac{(H_{k,n})^2}{H_{k,n}^{(2)}}\right)^{-1} \quad (14)$$

where in turn $H_{k,n}^{(m)} := \frac{1}{k} \sum_{i=1}^k \{\log X_{n-i+1,n} - \log X_{n-k,n}\}^m$, $m > 1$;

- * or a third alternative introduced in Beirlant *et al.* [23, 24]

$$\frac{1}{k} \sum_{i=1}^k \{\log UH_{i,n} - \log UH_{k+1,n}\} \quad (15)$$

where $UH_{i,n} = X_{n-i,n}H_{i,n}$, and $H_{i,n}$ is the Hill estimator.

In the POT method the choice of u is crucial while in the above estimators the choice of the quantity k is critical. In general, k should be small enough to avoid large **bias**, simultaneously k should be large enough to control the variance. One possibility, fully covered in Beirlant *et al.* [25,3], is to minimize a nonparametric estimator of the **mean squared error**. Other choices with their supporting arguments are presented in Danielsson *et al.* [26], Drees and Kaufmann [27], and Gomes and Oliveira [28]. In all cases k should tend to ∞ together with n but in such a way that $k/n \rightarrow 0$.

Estimation Of Extreme Quantiles

In practical situations one often likes to estimate the quantity x when the probability $p = P(X > x)$ has been fixed. In terms of the quantile function one then wants an estimate for the p th quantile, or $x = q_p := Q(1 - p)$. If we use our nonparametric approach, then we should use a random threshold like for example, the $(k + 1)$ -largest observation $X_{n-k,n}$.

From equation (11) we infer that the tail of F is well approximated by the right side of the equation. Given that we have an estimator $\hat{\gamma}$ of γ and one for $\hat{\sigma}$ of σ , q_p is estimated by

$$\hat{q}_p = X_{n-k,n} + \hat{\sigma} h_{\hat{\gamma}} \left(\frac{k}{np} \right) \quad (16)$$

where h_p is found in equation (3). This approach to estimate extreme quantiles and tail probabilities was first worked out in detail by Smith [14] and Dekkers and de Haan [29].

In the special case where $\gamma > 0$, we know from the Fisher–Tippett–Gnedenko result that

$$\lim_{t \rightarrow \infty} \frac{(1 - F(t))}{(1 - F(t))} = y^{-1/\gamma} \quad (17)$$

for every $y > 1$. Using the same thinking as above we find an alternative estimator

$$\hat{q}_p^+ = X_{n-k,n} \left(\frac{k}{np} \right)^{\hat{\gamma}} \quad (18)$$

a tail estimator for Pareto-type tails initiated in Weissman [30].

In **time series** analysis, in particular, in hydrology, an alternative way of looking at quantiles is very popular. The *return period* of an extremal event is defined by $T(x) = \{P(X > x)\}^{-1}$. It gives an indication of the time that elapses between two repetitions of events that themselves have a small probability of happening.

Note that in case $\gamma < 0$, one can let $p \downarrow 0$ to get an estimator for x_+ , the mostly unknown right-hand point of the distribution. More specifically, $\hat{x}_+ = \hat{\sigma}/|\hat{\gamma}|$.

Of course, one can also use the POT approach to get estimators for the quantiles and x_+ . For a full treatment, see Beirlant *et al.* [3].

Other Questions

- Without further assumptions, the construction of **confidence intervals** and the *reduction of the bias* in the above estimations are hardly possible. Indeed, in order to gain information about the asymptotic behavior of certain estimators, one needs *second-order properties* of the tail quantile function U or equivalently of F . We follow de Haan and Stadtmüller

[31] but other authors use different approaches. In addition to the existence of the EVI γ and the auxiliary function g , one assumes a second ultimately positive auxiliary function g_2 with $g_2(x) \rightarrow 0$ when $x \rightarrow x_+$, such that the limit

$$\lim_{x \rightarrow x_+} \frac{1}{g_2(x)} \left\{ \frac{U(ux) - U(x)}{g(x)} - h_{\gamma}(u) \right\} =: k(u) \quad (19)$$

exists and is neither identically zero nor a multiple of $h_{\gamma}(u)$. Quite generally one can then show that $g_2(u) = c h_p$ for a quantity $\rho \leq 0$ and a nonzero constant c . This second-order parameter appears in the detailed study of the bias and the variances of the classical estimators for the EVI γ . We refer to Beirlant *et al.* [3] for more explicit results and references.

- The area of **hypothesis testing** has hardly been explored. Among the few explicit results we mention the following. For an overview of what is available, see Fraga Alves and Neves [32].
 - Assume that $F \in \mathcal{C}_{\gamma}(g)$, where g is known. Then one can test the hypothesis $H_0(\gamma = 0)$ against the alternative $H_1(\gamma \neq 0)$. For two such approaches, see Hosking [33] and Segers and Teugels [34].
 - If F is strictly Pareto, that is, $1 - F(x) = cx^{-\alpha}$ then a hypothesis of the type $H(\alpha = \alpha_0)$ can be reduced to a similar test for the parameter of an exponential distribution. For explicit results, see Teugels and Van Assche [35] and its references.

Discussion of Available Results

As explained, the main goal of extreme value statistics is to estimate crucial parameters that are strongly influenced by the extremes in the sample. In particular we mentioned the EVI, high quantiles and endpoints of distributions. But overall, the assumptions made in the previous section are not always satisfied.

Classical Case

In cases where the classical theory does not apply immediately, the investigator has a number of possibilities to adapt the analysis. We mention a few.

Transformations. Sometimes a simple rewriting of the sample may already help to get a better view

of the right tail of the underlying distribution. If $\{X_i, 1 \leq i \leq n\}$ is a sample from X , then it might be advantageous to use a continuous function T to transform the data $\{X_i^* := T(X_i), 1 \leq i \leq n\}$ in such a way that the extremes of this new set have nicer statistical properties. For instance, the classical **Box–Cox transformation** $T(x) = (x^\sigma - 1)/\sigma$, ($\sigma \in \mathbb{R}$) often serves as an important example. If T is strictly increasing, then we can write for all $y > 1$:

$$\begin{aligned} P(T(X) > x) &= P(X > T^{-1}(x)) \\ &= \frac{1}{y} \iff x = T(U(y)) \quad (20) \end{aligned}$$

This means that the tail quantile function of $T(X)$ is given by the composition $T \circ U$. In Teugels and Vanroelen [36], the effect of a Box–Cox transformation on the data is fully investigated and illustrated with concrete examples. Of course, other transformations are possible as well.

Minima. Within a reliability context it is often more important to investigate the *minimum* of a sample rather than the maximum. This happens for example, in breakdown problems where the weakest link is responsible for the failure of the system. By the simple transformation $W_i = -X_i$, $1 \leq i \leq n$ on a sample, $\min_{1 \leq i \leq n} W_i = -\max_{1 \leq i \leq n} X_i$ and the classical results can be applied.

Regression Models. In classical statistics, a lot of attention has gone into methods to detect the role played by covariates on the common statistical procedures. This theme is far less popular in extreme value theory where estimations are rephrased in their conditional alternatives. Following Beirlant *et al.* [3], we can group the existing methods into four categories.

- The *method of block maxima* where the entire data set is cut in independent pieces and one or more parameters in the extremal distributions are expressed in terms of covariates.
- For the logarithms of the *spacings* between successive order statistics the classical theory has constructed exponential regression models based on second-order properties. These *quantile methods* can be extended to include **covariate** information.

- The *POT method* can be extended to allow the parameters in the generalized Pareto distribution to depend on covariates.
- A number of *nonparametric estimation procedures* such as maximum penalized likelihood or local polynomial maximum likelihood can be combined with extreme values.

Bayesian Analysis. In extreme value statistics, decisions often have to be taken on the basis of a small number of events. Moreover, such decisions are commonly of a prospective nature. For example, questions like 100-year return times are common in environmental contexts. Furthermore, prediction fits well in a Bayesian framework. Finally, some of the classical estimation approaches suffer from extra regularity conditions that need not be satisfied. For example, maximum-likelihood procedures sometimes fail if the EVI is negative. It is therefore advisable to help the investigator with expert opinions, an attitude that is typical for Bayesian statistics. Furthermore, Bayesian alternatives might help when classical methods fail.

In spite of the fact that some important attempts have been made, Bayesian Analysis for extremes is still in its infancy. Thanks to the availability of **Markov chain Monte Carlo** techniques some of the earlier computational problems can be circumvented. For a more or less full description of existing procedures and a number of case studies, see Beirlant *et al.* [3] or Coles [8] where further references are given.

Dropping the i.i.d. Assumption

In many applications the usual assumptions of independence and/or identical distribution are not realistic. Dropping the assumption of independence or/and of identical distribution leads to new problems.

- It has been shown in Leadbetter *et al.* [10] that the extremal laws are rather *robust* in the sense that under mild regularity conditions, the same laws appear for the case of weakly dependent random variables. See Deo [37] and Hsing [38] for more recent approaches for the statistical determination of an analog of the EVI.
- The incorporation of a further *time* or/and *spatial component* would of course make the above more applicable to practical data. In [39] Smith treats spatial models. Tail inference based on

high-level crossings of stationary sequences is treated by Leadbetter [40]. Time series models with extremal marginals are covered in Hall and Tajvidi [41].

- Rare events and extremes for *regenerative processes* are treated by Asmussen [42] and by Falk *et al.* [43]. For extremes in *Markov chains*, consult Bortot and Coles [44] and Bortot and Tawn [45].

Multivariate Extremes

For a number of years, extensions of univariate results to the situation of *multivariate extremes* have been developed. Take $\{(X_i, Y_i), 1 \leq i \leq n\}$ a sample from a bivariate distribution. It should be clear that even the concept of a multivariate extreme can be interpreted in a variety of ways. Still, the most common interpretation is $(X_{n,n}, Y_{n,n})$, the coordinate maximum. The basic limit theory has been started by de Haan and Resnick [46].

In multivariate extreme value theory one first makes transformations of the marginal distributions by applying a one-dimensional analysis on them. What remains can then be formalized in a *dependence function* or *copula*. Because of the high complexity of the latter, researchers fall back on parametric or semiparametric assumptions. Among the many proposals we mention a few.

- In [47], Tawn makes a number of proposals for modeling this dependence function. Nonparametric and parametric estimators of the parameters in the dependence function are discussed.
- For other representations, based on Pickands' approach, see Falk *et al.* [43].
- Hall and Tajvidi [48] suggest the use of a mixture of parametric and nonparametric methods to construct prediction regions in bivariate extreme value problems. The nonparametric part of the technique is used to estimate the dependence function (or copula) while the parametric part is employed to estimate the marginal distributions.
- Cai and Li [49] use a dependence function that is based on a multivariate phase-type distribution.

The problem of estimating multivariate *exceedance probabilities* was initiated in de Haan [50]. Also, in the multivariate setting one can break away from the i.i.d. case. For example, in [51] de Haan and Huang treat the problem of estimating extreme quantile

curves in two-dimensional distributions. Mikosch [52] models extremely large values in multivariate time series.

References

- [1] Gumbel, E.J. (1958). *Statistics of Extremes*, Columbia University Press, New York.
- [2] Galambos, J., Lechner, J. & Simiu, E. (1994). *Extreme Value Theory and Applications*, Kluwer, Dordrecht.
- [3] Beirlant, J., Goegebeur, Y., Segers, J. & Teugels, J.L. (2004). *Statistics of Extremes*, John Wiley & Sons, Chichester.
- [4] Embrechts, P.A.L., Klüppelberg, C. & Mikosch, T. (1997). *Modelling Extremal Events for Finance and Insurance*, Springer-Verlag, Berlin.
- [5] de Haan, L. (1970). *On Regular Variation and its Applications to the Weak Convergence of Sample Extremes.*, Mathematical Centre Tract 32, Amsterdam.
- [6] Fisher, R.A. & Tippett, L.H.C. (1928). Limiting forms of the frequency distribution of the largest or smallest member of a sample, *Proceedings of the Cambridge Philosophical Society* **24**, 180–190.
- [7] Gnedenko, B.V. (1943). Sur la distribution limite du terme maximum d'une série aléatoire, *Annals of Mathematics* **44**, 423–453.
- [8] Coles, S.G. (2001). *An Introduction to Statistical Modelling of Extreme Values*, Springer-Verlag, Berlin.
- [9] Galambos, J. (1978). *The Asymptotic Theory of Extreme Order Statistics*, John Wiley & Sons, New York.
- [10] Leadbetter, M.R., Lindgren, G. & Rootzén, H. (1982). *Extremal and Related Properties of Stationary Processes*, Springer-Verlag, New York.
- [11] Reiss, R.-D. & Thomas, M. (1989). *Statistical Analysis of Extreme Values*, 2nd Edition, Birkhäuser Verlag, Basel.
- [12] Resnick, S.I. (1987). *Extreme Values, Regular Variation and Point Processes*, Springer, New York.
- [13] Smith, R.L. (1985). Maximum likelihood estimation in a class of non-regular cases, *Biometrika* **72**, 67–90.
- [14] Smith, R.L. (1987). Estimating tails of probability distributions, *Annals of Statistics* **15**, 1174–1207.
- [15] Csörgő, S., Deheuvels, P. & Mason, D. (1985). Kernel estimates of the tail index of a distribution, *Annals of Statistics* **13**, 1050–1077.
- [16] Hosking, J.R.M., Wallis, J.R. & Wood, E.F. (1985). Estimation of the generalized extreme-value distribution by the probability-weighted moments, *Technometrics* **27**, 251–261.
- [17] Castillo, E. & Hadi, A.S. (1997). Fitting the generalized Pareto distribution to data, *Journal of the American Statistical Association* **92**, 1609–1620.
- [18] Hill, B.M. (1975). A simple general approach to inference about the tail of a distribution, *Annals of Statistics* **3**, 1163–1174.
- [19] Pickands III, J. (1975). Statistical inference using extreme order statistics, *Annals of Statistics* **3**, 119–131.

- [20] Drees, H. (1995). Refined Pickands estimators of the extreme value index, *Annals of Statistics* **23**, 2059–2080.
- [21] Segers, J. (2005). Generalized Pickands estimators for the extreme value index, *Journal of Statistical Planning and Inference* **128**, 381–396.
- [22] Dekkers, A.L.M., Einmahl, J.H.J. & de Haan, L. (1989). A moment estimator for the index of an extreme-value distribution, *Annals of Statistics* **17**, 1833–1855.
- [23] Beirlant, J., Vynckier, P. & Teugels, J.L. (1996). Excess functions and estimation of the extreme-value index, *Bernoulli* **2**, 293–318.
- [24] Beirlant, J., Vynckier, P. & Teugels, J.L. (1996). Tail index estimation, Pareto quantile plots and regression diagnostics, *Journal of the American Statistical Association* **91**, 1659–1667.
- [25] Beirlant, J., Teugels, J.L. & Vynckier, P. (1996). *Practical Analysis of Extreme Values*, Leuven University Press, Leuven.
- [26] Danielsson, J., de Haan, L., Peng, L. & de Vries, C.G. (2001). Using a bootstrap method to choose the sample fraction in tail index estimation, *Journal of Multivariate Analysis* **76**, 226–248.
- [27] Drees, H. & Kaufmann, E. (1998). Selecting the optimal sample fraction in univariate extreme value estimation, *Stochastic Processes and their Applications* **75**, 149–172.
- [28] Gomes, M.I. & Oliveira, O. (2001). The bootstrap methodology in statistics of extremes - choice of the optimal sample fraction, *Extremes* **4**, 331–358.
- [29] Dekkers, A. & de Haan, L. (1989). On the estimation of the extreme value index and large quantile estimation, *Annals of Statistics* **17**, 1795–1832.
- [30] Weissman, I. (1978). Estimation of parameters and large quantiles based on the k largest observations, *Journal of the American Statistical Association* **73**, 812–815.
- [31] de Haan, L. & Stadtmüller, U. (1996). Generalized regular variation of second order, *Journal Australian Mathematical Society, (Series A)* **61**, 381–395.
- [32] Fraga Alves, I. & Neves, C. (2006). Testing extreme value conditions - an overview and recent approaches, in *International Conference on Mathematical and Statistical Modeling in Honor of Enrique Castillo*, La Universidad de Castilla, La Mancha; Universidad de Cantabria, Spain.
- [33] Hosking, J.R.M. (1984). Testing whether the shape parameter is zero in the generalized extreme-value distribution, *Biometrika* **71**, 367–374.
- [34] Segers, J. & Teugels, J.L. (2000). Testing the Gumbel hypothesis using Galton's ratios, *Extremes* **3**, 291–303.
- [35] Teugels, J.L. & Van Assche, W. (1986). Sequential testing for exponential and Pareto distributions, *Sequence Analysis* **5**, 223–236.
- [36] Teugels, J.L. & Vanroelen, G. (2004). Box-Cox transformations and heavy-tailed distributions, *Journal of Applied Probability* **41A**, 213–227.
- [37] Deo, R.S. (2000). On estimation and testing goodness of fit for m -dependent stable sequences, *Journal of Econometrics* **99**, 349–372.
- [38] Hsing, T.L. (1993). Extremal index estimation for a weakly dependent stationary sequence, *Annals of Statistics* **21**, 2043–2071.
- [39] Smith, R.L. & Robinson, P.J.A. (1997). A Bayesian approach to the modeling of spatial-temporal precipitation data. in *Case studies in Bayesian statistics, Lecture Notes in statistics, Vol. 121*, Gatsonis, C. Hodges, J. Kass, R. McCulloch, R. Rossi, P. & Singpurwalla, N., eds, Vol. 3. Proceedings of a workshop Carnegie Mellon University, Pittsburgh, PA, USA. October 5–7, 1995. New York: Springer Springer-Verlag, pp. 237–269.
- [40] Leadbetter, M.R. (1995). On high level exceedance modelling and tail inference, *Journal of Statistical Planning and Inference* **45**, 247–260.
- [41] Hall, P. & Tajvidi, N. (2000). Nonparametric analysis of temporal trend when fitting parametric models to extreme-value data, *Statistical Science* **15**, 153–167.
- [42] Asmussen, S. (2003). *Applied Probability and Queues*, 2nd Edition, J.L. Teugels & B. Sundt, eds, Springer-Verlag, New York; John Wiley & Sons, Chichester, p. 654–661.
- [43] Falk, M., Hüsler, J. & Reiss, R.-D. (2004). *Laws of Small Numbers: Extremes and Rare Events*, Birkhäuser Verlag, Basel.
- [44] Bortot, P. & Coles, S. (2003). Extremes of Markov chains with tail switching potential, *Journal of the Royal Statistical Society. Series B* **65**, 851–867.
- [45] Bortot, P. & Tawn, J.A. (1998). Models for the extremes of Markov chains, *Biometrika* **85**, 851–867.
- [46] de Haan, L. & Resnick, I.S. (1977). Limit theory for multivariate sample extremes, *Zeitschrift für Wahrscheinlichkeitstheorie und Verwandte Gebiete* **40**, 317–337.
- [47] Tawn, J.A. (1990). Modelling multivariate extreme value distributions, *Biometrika* **77**, 245–253.
- [48] Hall, P. & Tajvidi, N. (2006). Prediction regions for bivariate extreme events, *Australian New Zealand Journal of Statistics* **46**, 99–112.
- [49] Cai, J. & Li, H. (2005). Conditional tail expectations for multivariate phase-type distributions, *Journal of Applied Probability* **42**, 810–825.
- [50] de Haan, L. (1994). Estimating exceedance probabilities in higher-dimensional space, *Communications in Statistics - Stochastic Models* **10**, 765–780.
- [51] de Haan, L. & Huang, X. (1995). Large quantile estimation in a multivariate setting, *Journal of Multivariate Analysis* **53**, 247–263.
- [52] Mikosch, T. (2005). How to model multivariate extremes if one must, *Statistica Neerlandica* **59**, 324–338.

Related Article

Applications of Extreme Statistics in Science and Engineering.

JOZEF L. TEUGELS

Intercalibration Studies

Introduction

The Biological Effects Quality Assurance in Monitoring Programmes (BEQUALM) self-funded scheme was initiated to develop quality standards for a range of biological effects techniques and to monitor compliance of laboratories generating data from these techniques for national and international monitoring programs (e.g. OSPAR, Oslo and Paris Commissions) and also for regulatory purposes. This self-funded scheme has now been running for 2 years. The “whole organism” component of this scheme deals with aquatic ecotoxicological techniques, which determine the inherent toxicity of a sample of water or sediment, using a variety of aquatic organisms. Each year, the lead laboratory (Cefas) organizes a program of intercalibration exercises for each of the methods included in the scheme. In this paper, the results from an intercalibration using a 48-h acute test with *Daphnia magna* (water flea) are described.

Each intercalibration exercise is divided into two parts. In the first (termed the *unknown test*), laboratories were given an unspecified concentration of a reference chemical (zinc sulphate), which they were asked to dilute by a given factor and estimate the toxicity of the substance in terms of the percentage mortality. In the second part (termed the *reference test*), a predetermined dilution series of the reference chemical was prepared by each laboratory from a provided stock and mortality at each concentration measured. From these mortalities, the medial lethal concentration (EC_{50}) was estimated for each laboratory.

In the *Daphnia* tests that are described here, six laboratories took part in the unknown tests and seven in the reference tests. However, repeat tests using different operators – particularly from laboratory 11 – meant that effectively eleven sets of results were received for both the unknown and the reference tests. To maintain confidentiality of a laboratory's results, each was randomly allocated a number.

All computing was done using the package S-plus [1].

Statistical Methods

The Unknown Test

All but one laboratory deployed a test design involving four replicates of five *Daphnia* for both controls and the unknown concentration. One laboratory used 2 replicates of 10 *Daphnia*. The controls were not used in the analysis but were used as a criterion for validity of the test, i.e. to ensure that the test animals were healthy and the environmental conditions of the test were correct. If more than 10% of the control animals died then the unknown sample data was rejected. For the *Daphnia* data described here, this did not happen for any laboratory.

The important summary statistics for the unknown samples are the mean and the standard deviation of the mortality rate estimated by each laboratory. From these, simple **confidence intervals** for the laboratory's mean mortality rate can be obtained.

The next step is to establish whether any of the results from the laboratories are inconsistent – i.e. determine if there are problems with any laboratory's procedures. The difficulty with doing this is that a **benchmark** by which to judge the performance of the laboratories needs to be established. Ideally, this should be independent of the current data. In this case it was possible to use data from the previous year, since the same concentration of the unknown sample was used. A *consensus mean* was thus calculated based on the laboratories that were deemed compliant in the 2004/2005 testing [2]. However, if the level of the unknown sample had changed – as it will in future years – the consensus mean would need to be based upon the current year's data. In such cases, some form of robust estimator that is less influenced by **outliers** would be preferable. This, therefore, means that the results from extreme laboratory results do not unduly influence the level of the consensus mean (if they did, it would mean that such a laboratory would be less likely to be deemed noncompliant). The preferred choice of such an estimator – which is used for the second stage of the study – is the robust M-estimator with bi-square down-weighting of **residuals** [1].

However the consensus mean is calculated, the assumption is effectively made that the majority of the laboratories are “correct” and that laboratories that deviate from this consensus are noncompliant. Clearly, this is where some nonstatistical ecotoxicological detective work can be used to establish why

2 Intercalibration Studies

laboratories are noncompliant or, indeed, whether they are correct and all of the others are wrong!

A common statistical technique to assess whether the results of a laboratory are consistent with the consensus mean is through the use of *Z*-scores [3]. This assumes the null hypothesis that each laboratory mean ($\bar{p}$) has a **Normal distribution** with mean equal to the consensus mean (p) and variance given by $\text{Var}(p) = p(1 - p)/N$, where N is the total number of animals over all replicates for the unknown sample ($N = 20$ for the *Daphnia* data). If the modulus of

$$Z = \frac{\bar{p} - p}{\sqrt{\text{var}(p)}} \quad (1)$$

is larger than, say, the 97.5% quantile of a standard Normal distribution (tolerance limit), then the laboratory is deemed to be noncompliant. A slight technical improvement on this can be obtained, as has been done here, by using exact tolerance limits based on Binomial distributions rather than approximations to the Normal distribution.

While the use of *Z*-score methods allows us to investigate inconsistencies of a laboratory's mean estimate of mortality, it does not allow for consideration of the variation over repeat tests. One way to do this is to use a, so-called, *P*-score defined as

$$P = \frac{s_p}{\sqrt{\text{Var}(p)}} \quad (2)$$

where s_p is the standard deviation over replicates for a particular laboratory. Effectively, a comparison is being made of the ratio of the observed laboratory standard deviation to its binomial theoretical value assuming that the mortality rate is the same for all replicates and that mortality is independent. A *P*-score much greater than unity suggests that there is additional variation in the experimental procedure. Exact tolerance limits were obtained by simulating data based on the binomial distribution with a probability of death of the consensus mean and calculating the *P*-score each time.

The Reference Test

Here, each laboratory tests across a series of contaminant concentrations and mortality is calculated at each level. The Majority of the laboratories used four replicates of five *Daphnia* for each of the contaminant concentrations stipulated to be tested by the lead

laboratory: 0, 0.1, 0.32, 1.0, 3.2, and 10 mg l⁻¹ zinc. However, laboratory 8 used additional intermediate concentration levels at 0.56 and 1.8 (which presumably was the concentration range this laboratory routinely used when conducting *Daphnia* tests), laboratory 9 used only 2 replicates but with 10 *Daphnia* in each replicate, and laboratory 12 used 6 replicates at concentration level 0. These different designs are all acceptable within current international guidelines (e.g. OECD [4]; ISO [5]).

The EC₅₀ is defined as the concentration at which the mortality is 50%. If mortality does not reach 50% for any laboratory, then it is excluded from the analysis and is deemed clearly noncompliant. The EC₅₀ was estimated by modeling the mortality proportion (m) as a logistic **generalized linear model** with $\log(\text{concentration} + 0.001)$ as the explanatory variable [6] and with binomial errors. Thus

$$m = \frac{\exp(\beta_0 + \beta_1 \log \text{Conc})}{1 + \exp(\beta_0 + \beta_1 \log \text{Conc})} \quad (3)$$

The EC₅₀ is estimated by finding the concentration for which $m = 0.5$ in equation (3). This gives $\widehat{EC}_{50} = e^{-\widehat{\beta}_0/\widehat{\beta}_1}$, where $\widehat{\beta}_0$ and $\widehat{\beta}_1$ are the maximum-likelihood estimates.

Confidence intervals for the EC₅₀ estimates could not be obtained empirically because the majority of laboratories carried out the procedure only once. One possibility would have been to construct **bootstrap** confidence intervals for the EC₅₀ based on **resampling** the data [7]. However, a profile-likelihood method was chosen to construct 95% profile-likelihood confidence intervals [6]. The method is more fully described in Appendix A.

Similar diagnostics were calculated for the values of EC₅₀ in terms of their mean value and their precision, as was done for the unknown test above. However, because repeat replicates were not available, the procedure required modification. The consensus EC₅₀ (denoted as EC₅₀^c) was estimated from a logistic model fitted to the pooled 2004/2005 data – but excluding noncompliant stations. A form of “*Z*-score” indicating the difference between the consensus mean and that from each laboratory was calculated as follows:

$$Z = \text{EC}_{50}^c - \text{EC}_{50} \quad (4)$$

The 95% tolerance limits for the EC₅₀ *Z*-scores were estimated by simulation (see Appendix B). As

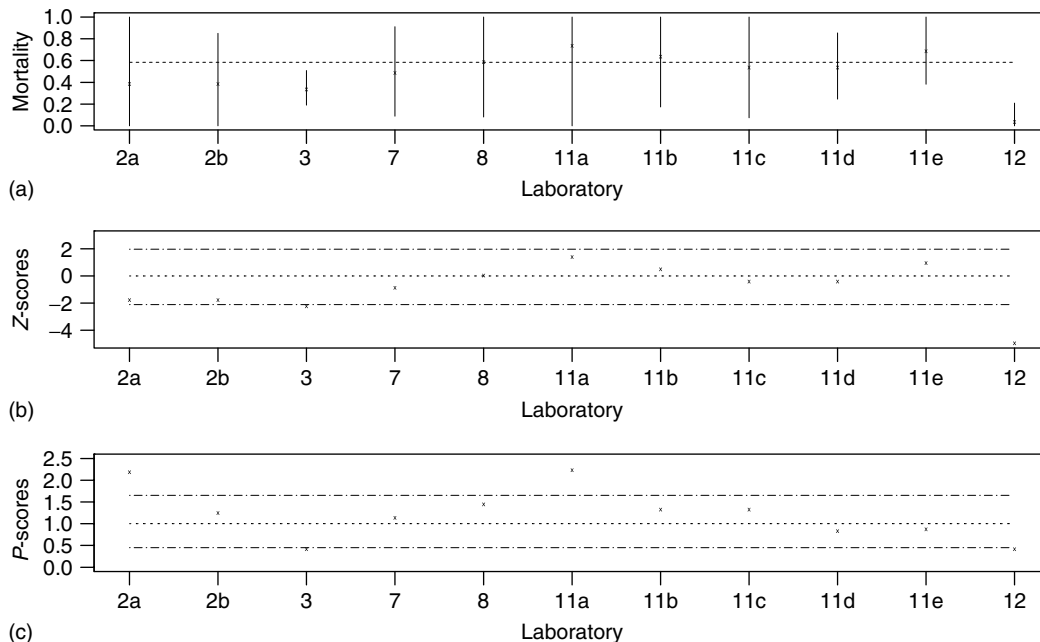


Figure 1 Unknown concentration plots for *Daphnia*. (a) Mean mortality estimates and confidence intervals; (b) Z-scores and 95% tolerance limits; (c) P-scores and 95% tolerance limits

previously stated, those laboratories where the Z-score is contained within the tolerance limits could be considered to have complied.

To assess the variance of the EC_{50} estimate obtained by each laboratory an approximate method based on a truncated Taylor series expansion was used. The ratio of this variance was compared with the empirical EC_{50} estimates from the simulations used above to obtain the tolerance limits for the EC_{50} Z-scores (further details are given in Appendix C). As done for the unknown test, an evaluation is being made of how consistent the variance of the EC_{50} from an individual laboratory is with the variance obtained under the underlying binomial model for mortality and the consensus EC_{50} .

Results

Unknown Test. The consensus mean of 0.58 was calculated from the 2004/2005 data, excluding the results of two laboratories where the mean mortality was exceptionally high or low.

The mean mortality rates and concentrations (Figure 1(a)) suggest that laboratories 3 and 12 have

low mean estimates. This is confirmed by the Z-scores plot (Figure 1(b)) which shows that laboratories 3 and 12 are noncompliant in terms of their mean estimate. The P-scores plot (Figure 1(c)) suggests that laboratories 2a and 11a have variances that are too high. However, given that all of the other results for laboratory 11 are compliant in terms of mean and variance, then, overall, this laboratory seems to have performed quite well. The variances for laboratories 3 and 12 are on the border of the lower limit. Coupled with their rejections from the Z-score tests, there seems to be some indication that there are problems with method performance.

Reference Test. Figure 2 shows the EC_{50} estimates for each laboratory together with their fitted logistic model. Laboratories 11(a) and (b) have a steep gradient from 0 to 100% deaths and the deaths for laboratory 12 do not get very high even for the highest concentration (fitting in with the low deaths for the unknown concentration for this laboratory).

Figure 3 shows the EC_{50} estimates and profile-likelihood confidence intervals. Laboratories 2b and 12 are the furthest from the consensus EC_{50} , which

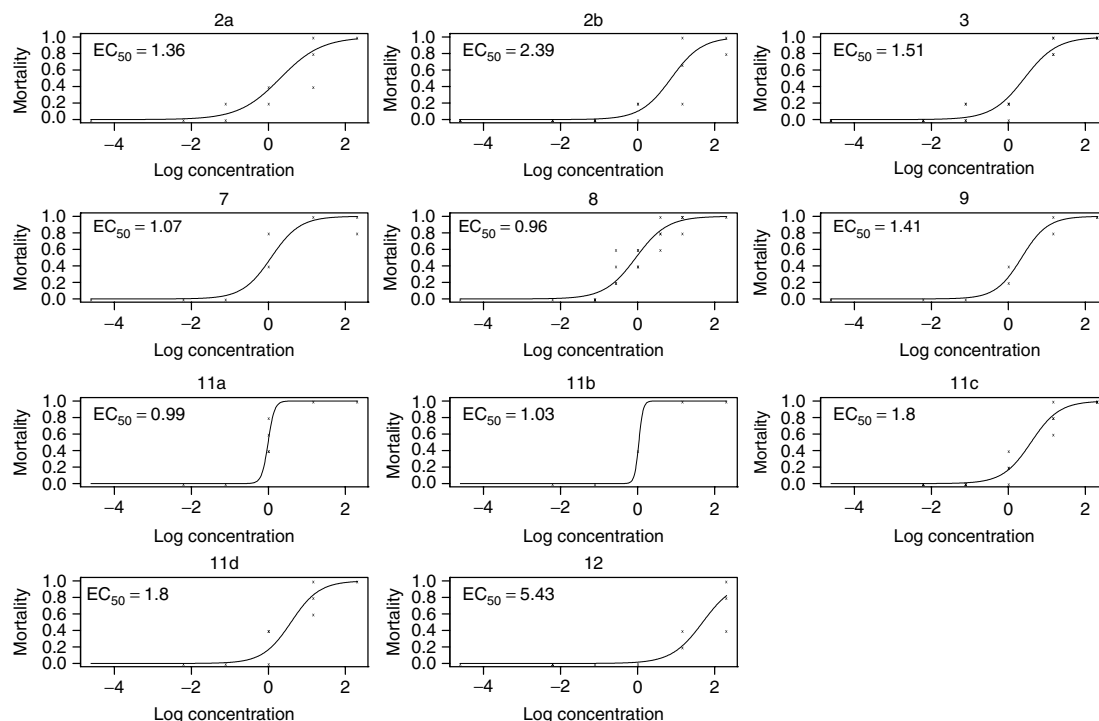


Figure 2 EC_{50} models and estimates for each laboratory

was calculated as 1.35 from the 2004/2005 data, excluding outlying results from two laboratories.

Figure 4(a) shows the EC_{50} Z-scores plot. Laboratories 2b, 11c, 11d, and 12 are all noncompliant. Figure 4(b) shows the P -scores plot. Laboratories 2b and 12 have excessively high variability, whereas 8, 11a, and 11b are all marginally too low. Overall, laboratory 12 demonstrates particularly poor performance in the test.

Discussion

We have described the procedures that are used to analyze some of the biological effects data in the BEQUALM scheme. The aim of this quality assurance scheme is to evaluate and monitor the ability of laboratories to competently perform ecotoxicity tests, the data from which are being submitted to national and international environmental monitoring programs and also for regulatory purposes (e.g. chemical registration). There are several benefits in taking part in such proficiency testing schemes. For example, satisfactory performance allows a laboratory to

provide additional confidence in the quality of their results to laboratory clients; unsatisfactory performance indicates a problem with the methodology at a laboratory, whether that be inexperience with the test and/or operator, problems with equipment **calibration**, and so on, and allows remedial actions to be initiated.

The self-funded scheme has been running for 2 years. Both practical and statistical methods have been modified as a result of the first year. It is expected that modifications will continue to be made into the second year. There are a number of improvements that could be made. For example, confirmation is needed that the sample sizes (five *Daphnia*, four replicates) used in the unknown tests are sufficient to ensure that important biological departures by laboratories can be picked up in the Z-score and P -score tests. Having said that, given that a number of laboratories are currently picked out as noncompliant by these tests, the sample sizes do seem to be reasonably adequate. It would be very beneficial if laboratories performed replicate tests so that empirical estimates of the variance of EC_{50} estimates could

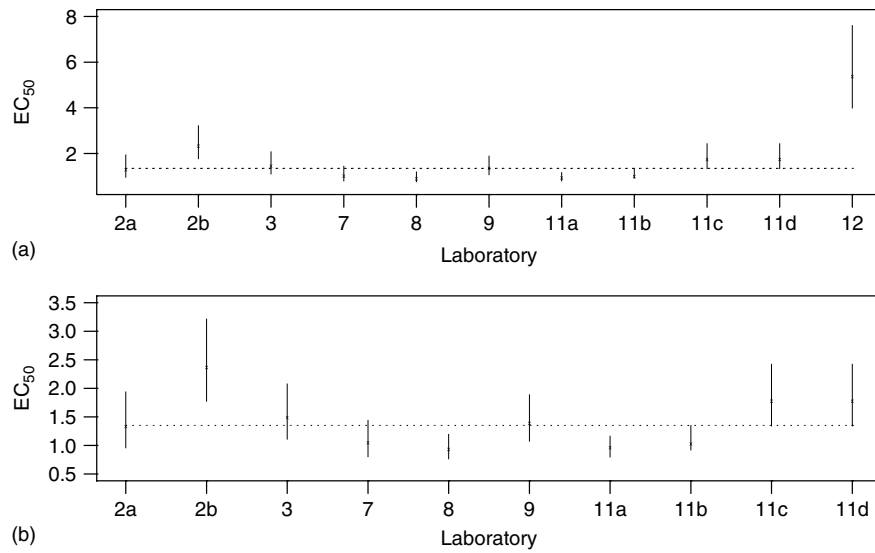


Figure 3 EC_{50} estimates and profile-likelihood confidence intervals: (a) all laboratories; (b) excluding laboratory 12

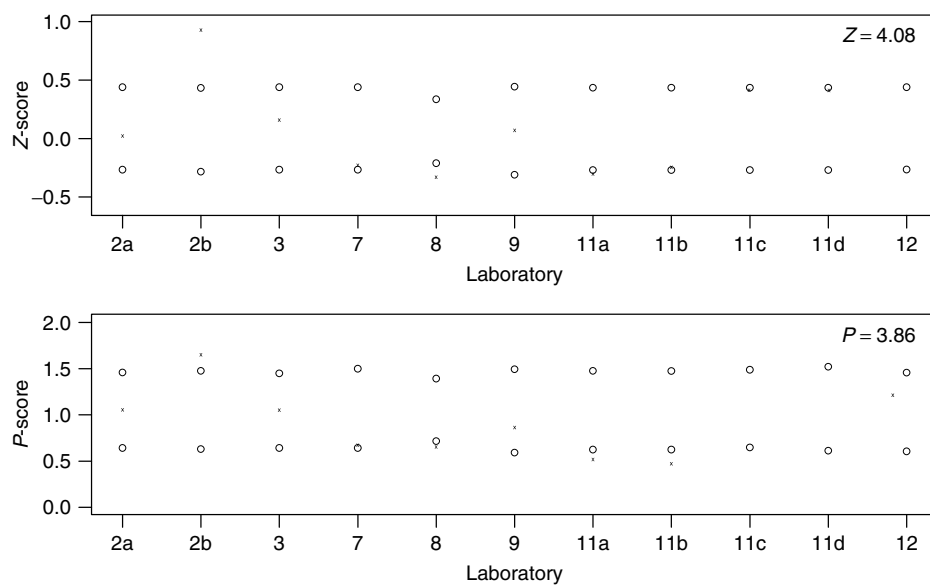


Figure 4 Z-scores and P-scores for the EC_{50} estimates. The values for laboratory 12 are not plotted because they are so extreme

be obtained. This would remove the need for some of the model-based and simulation techniques described in this article.

Many of the methods described in this article

are based on computer-intensive model fitting and simulation. These methods are not, of course, an end in themselves. The intention has always been to use these to provide evidence for potential deficiencies in

the procedures carried out by laboratories and thus discriminate between acceptable and unacceptable data for inclusion in large monitoring/regulatory programs.

Appendices

Appendix A – EC_{50} Profile-Likelihood Confidence Interval

To construct a profile-likelihood 95% confidence interval we reparameterized the logistic model as follows:

$$\begin{aligned}\log \text{it}(p) &= \beta_0 + \beta_1 \log \text{Conc} \\ &= \beta_1 (\log \text{Conc} - \log EC_{50}) \\ &= \beta_1 X\end{aligned}\quad (5)$$

To construct the profile likelihood, we vary the estimate of the EC_{50} around its original maximum-likelihood estimate $\hat{EC}_{50}$ and construct the 95% profile-likelihood confidence interval to include values of EC_{50} such that $(\text{Dev}(\hat{EC}_{50}) - \text{Dev}_0) \leq 3.84$, where Dev_0 is minus twice the log-likelihood at the maximum-likelihood estimator and $\text{Dev}(\hat{EC}_{50})$ is minus twice the log-likelihood estimator for model (5) above.

Appendix B – Tolerance Limits for EC_{50} Z-scores

The 95% tolerance limits for the Z-scores were estimated by simulation. At each concentration, we simulated binomial data with number of trials equal to the number used by the laboratory – with the probability of death taken from the value from the fitted model used in calculating the consensus mean from the 2005 data. From these simulated data, we calculated the EC_{50} estimate. We repeated this process 1000 times and the approximate 95% tolerance limits for the Z-scores are given by the 25th and the 975th ordered values.

Appendix C – Comparing the Variances of the EC_{50} Estimates

The variance of the EC_{50} estimate for a particular laboratory is obtained using an approximate method

based on the parameter estimates from the logistic model. This is derived as follows:

$$\log \hat{EC}_{50} = -\frac{\beta_0}{\beta_1} \quad (6)$$

Using a standard truncated Taylor series expansion we get

$$\begin{aligned}\text{Var}(\log \hat{EC}_{50}) &= \frac{1}{\beta_1^2} \left\{ \text{Var}(\hat{\beta}_0) \right. \\ &\quad + 2 \log \hat{EC}_{50} \text{Cov}(\hat{\beta}_0, \hat{\beta}_1) \\ &\quad \left. + (\log \hat{EC}_{50})^2 \text{Var}(\hat{\beta}_1) \right\}\end{aligned}\quad (7)$$

Thus, using a similar Taylor series truncation we obtain

$$\text{Var}(\hat{EC}_{50}) = \hat{EC}_{50}^2 \text{Var}(\log \hat{EC}_{50}) \quad (8)$$

The variances and the covariances of the parameter estimates come from the fitted logistic model for each laboratory.

Define an EC_{50} P-score by

$$P = \frac{\text{Var}(\hat{EC}_{50})}{\text{Var}(\hat{EC}_{50\text{sim}})} \quad (9)$$

where $\text{Var}(\hat{EC}_{50\text{sim}})$ is the variance of the 1000 EC_{50} s calculated from the simulations described in Appendix B.

The tolerance limits come from the 25th and 975th largest values of equation (9) when the approximate variance $\text{Var}(\hat{EC}_{50})$ is calculated from each of the simulations.

References

- [1] Venables, W.N. & Ripley, B.D. (2002). *Modern Applied Statistics with S*, 4th Edition, Springer, New York.
- [2] BEQUALM Whole Organism Component. (2005). *Daphnia magna* intercalibration exercise Year 1 2004/2005, BEQUALM report: CEFAS, Remembrance Avenue, Burnham-on-Crouch, CMO 8HA.
- [3] Wells, D.E. & Cofino, W.P. (1997). The assessment of the QUASIMEME laboratory performance studies data: techniques and approach, *Marine Pollution Bulletin* **35**, 18–27.
- [4] OECD. (2004). *Guideline for Testing of Chemicals 202 – Daphnia sp. Acute Immobilisation Test*, Organisation for Economic Co-operation and Development.
- [5] International Organisation for Standardisation. (1996). *Water Quality - Determination of the inhibition of*

- mobility of *Daphnia magna* Straus (Cladocera, Crustacea) – Acute Toxicity Test. ISO International Standard 6341:1996E, 9pp.
- [6] McCullagh, P. & Nelder, J.A. (1989). *Generalized Linear Models*, 2nd Edition, Chapman & Hall, London.
- [7] Davison, A.C. & Hinkley, D.V. (1997). *Bootstrap Methods and Their Application*, Cambridge University Press, Cambridge.

JON BARRY AND YVONNE ALLEN

Fisheries Stock Assessment

Fishing is an economically and socially important activity with around 90 million tonnes of wild fish caught each year [1]. Fisheries policy tries to balance potentially conflicting management drivers for sustainability and maintenance of ecosystem biodiversity with socioeconomic drivers of profit and employment at a variety of geographic scales. Stock assessment's role is to provide the best possible technical support to this process [2], estimating current and possible future fishery states, and the effects of management decisions. Stock is used to define a part of a fish or shellfish population. Stocks are usually identified by location, e.g., "Bay of Biscay sole", and generally treated as self-sustaining and distinct from other stocks.

The assessment process involves collecting data direct from the fishery and from "fishery-independent" sources, such as research vessel surveys. To turn these data into quantitative estimates of stock size and composition, a wide range of mathematical models is available [3]. These range from simple line fitting to complex models that integrate multiple data sources. Cotter *et al.* [4] note that stock assessments should be explicable to stakeholders; allow data sets to be weighted in relation to the independent information they contribute; use sampling data or estimates only once, and make due allowance for dependence among data. There is an art to a good assessment: to ensure that modeling assumptions are reasonable, that outputs translate sensibly to the observed situation, and information outside the mathematical model is considered, as "no fishery model can be completely trusted to capture biological reality" [5].

Data Collection

Representative **data collection** is needed to enable accurate and precise stock assessments. This falls into two categories: data referring to the fishery (catch, effort) and those referring to stock biology (length, weight, age, etc.). Details of sampling strategies and techniques should be noted in standard operating

procedures (SOPs) to ensure consistency over time and between samplers.

Data on the amount of fishing effort and fish landed are collected in two ways: census or more commonly **survey sampling**. Sampling officers visit ports and record required details for a selection of vessels, while interviews, either face to face or via telephone, may be used to collect fishery information from recreational or artisanal fisheries. See Catch Statistics in [6] for more details. Avoiding or accounting for unreported landings is an important **data quality** issue for fisheries. **Quality control** can be aided by cross-checking logbook data recorded at sea, landing declarations, and sales receipts. Observation flights and satellite-based position monitoring on larger vessels also provide useful information on fishing locations.

As fishery-based data reflect fishing practices, economics, and technological improvements as well as stock biology, "fishery-independent" data collected using standardized methods plays a vital role in many assessments. This is most commonly from a trawl survey using a research or chartered fishing vessel (see Catch Statistics and Trawl Surveys [6]). Alternatively, acoustic surveys can be used to cover large volumes of water. These surveys use equipment that must be frequently calibrated with a reference target of known target strength, as it is affected by regional changes in water salinity and temperature (affecting sound speed and absorption) [7]. Other abundance **estimation** methods include aerial surveys to identify and quantify schools of fish, use of underwater cameras, and surveys which sample fish egg distributions, leading to fish population estimates based on the number of eggs [8].

A basic quality assurance matter in all these approaches is the correct identification of species. Visual identification of common and commercial species is generally near perfect in both commercial and fishery-independent approaches, but may be inconsistent for rare species or juveniles of related species without regular training and assessment. For acoustic surveys, converting acoustic readings to species densities can work well for single species aggregations and individual animals, while for complex multispecies aggregations, there have been considerable developments in hardware and classification algorithms but automated species identification remains the goal. The ability of scientists to correctly

2 Fisheries Stock Assessment

identify the species of fish eggs has been considerably enhanced by DNA-based methods [9].

Biological information is often collected alongside fishing data, using stratified **sampling schemes**. The simplest biological data to collect are fish lengths, grouped to form length distributions. The box tally system (Figure 1) improves the quality of data recording to paper as it is clearer and provides less chance for error than the more usual tally gate. Direct-to-computer length recording, using devices such as electronic calipers and electronic data-capture fish-measuring boards, also improves quality. The boards have a ruled scale with a bar code at each length. When an infrared pen is swiped across the bar code, translation software on a computer recognizes the measurement and enters it into a database. Quality control checks on individual fish length and weight are incorporated into fisheries data entry systems using historical records of maximum lengths and published relationships between length and weight [10].

For many fish species, age can be estimated from hard body structures such as scales, spines, and otoliths (calcareous structures near the back of the skull). As fish grow, marks are deposited through which fish age can be estimated, similar to counting rings in a tree trunk (Figure 2). The marks must be validated, confirming that the start and frequency of growth patterns matches the expectations for fish of known age [11]. Common interpretation of patterns must be ensured. With a sample of known ages, the age interpretation by individual scientists can be assessed graphically with an age-**bias** plot (known age against estimated age) and statistically tested for bias through aging error matrices and matched-pair tests. As known age samples are not regularly available, samples independently aged by a number of scientists can be compared, and measures of bias and precision calculated from the most common age ([12]; www.efan.no).

Measuring a fish is much cheaper than aging it. Therefore, often a subsample of measured fish are aged, and observed age-length keys (proportions at age by length) or modeling [13] used to estimate

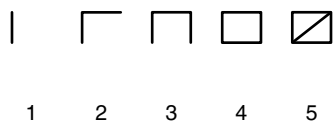


Figure 1 Box tally recording symbols

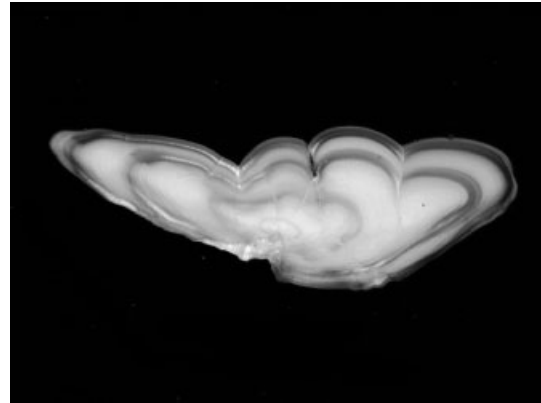


Figure 2 Cod otolith cross section showing annual growth zones. Actual width, 12.5 mm, tip to tip

the age composition of the total landings needed for many stock assessments. Recent developments incorporate otolith weight as additional information in this process [14].

As with age, fish maturity is assessed visually. By inspecting the fish's gonads, scientists can tell if the fish has the ability to spawn and its stage in the annual spawning cycle, based upon standard texts e.g., [15]. Microanalysis of gonads is usually recommended for precise assessment of the developmental stage, rather than macroanalysis, but is more time consuming. Sample exchanges to ensure consistency are problematic as they involve soft tissue. Instead, exchanges of digital photographs and scientific workshops can be used.

Fish tagging studies use simple marker tags and more recently electronic tags, which provide a detailed **time series** of observations. These studies provide information on fish migration, help define stocks, and provide growth and survival data. They can also provide the basis for mark-recapture methods to estimate population size (e.g., [16]), particularly for species such as salmon that have part of their life cycle in freshwater. Double tagging of fish can investigate the effects of tagging on mortality and growth, or estimate rates of tag shedding.

Stock Assessment Methods

There is a wide and expanding variety of stock assessment methods, the most appropriate of which will depend upon the nature of the fishery, the presence

and quality of biological and fishery data, and the skills available [17, 18]. Most operational assessments are for single stocks, although multispecies models [19] are used for scenario modeling and meta-analysis is used to study patterns across many stocks. There is also growing emphasis on assessing the effects of fishing on ecosystems as a whole.

Single stock methods develop from those using only catch and effort data, by including biological processes such as recruitment (the amount of juvenile fish entering the fishery) and by increasing the biological structure with length, age, or stage-based models. Spatial structure of the stock may also be included. Many of the methods have now been developed in a state–space framework to include observation and process error. Bayesian estimation methods are increasingly common, allowing integrated assessment of uncertainty in assessment outputs [20].

Time series of catch and effort data allow biomass dynamic models (also called *surplus production methods*) to be used. These simply model population growth as a function of total biomass. Delay-difference models extend this by including age-structured dynamics and a time lag between spawning and recruitment. Depletion methods based on a closed system estimate stock size by extrapolating decreases in catch as cumulative effort increases. Catch survey analysis extends this to an open system and has similarities to mark-recapture analysis. With detailed information on the contribution of individual cohorts to the catch (through sampling programs and age readings), catch-at-age methods can estimate the historical status of a population. By tracking catch-at-age back through time, they estimate the population that must have been present in the water to produce the observed catch. Cohort or virtual population analysis (VPA)-based methods assume catch data are exact, whereas other catch-at-age methods allow observation and process errors by making other modeling assumptions. For species that are difficult to age, length-based models can be used [3, 21]. These generally either split the length frequency data into separate cohorts based on an assumed growth model and length-at-age distributions or model the transition from one size group to the next.

With many of the methods, catch per unit effort (CPUE) is used as a measure of abundance either directly or as additional “tuning” information for **calibration**. An assumption that CPUE is proportional to abundance should be carefully checked, particularly

for fisheries-dependent data, otherwise it may cause bias. Often, **generalized linear models** and their extensions are used to improve the index of abundance [22].

After estimating current status, management considerations require short-term forecasts of catch and biomass against management objectives. These forecasts require an estimate of current population size, the likely exploitation pattern, estimates of the size of recruiting year classes, and assumptions of future overall levels of fishing mortality. The population is projected forward to identify which levels of fishing may achieve desired future stock status. These projections can be highly dependent upon the assumptions made. As a result, the longer the projection, the more uncertain the output, as projections are unlikely to be able to predict shifts in the environment or the impact of multispecies interactions. When communicating these forecasts and stock assessment results, an element of probability can be used, e.g. the estimated probability of an unwanted event occurring or cumulative probability distributions of meeting a required biological level [23].

Assessing the Assessment

Once complete, stock assessments are peer reviewed before findings are presented to managers and policy makers. Statistical methods are available to estimate uncertainty in the results [24], while reporting of uncertainty is an integral part of Bayesian approaches. The performance of stock assessments can be determined against reference data sets, although these may not reflect all the issues encountered in practice. Various diagnostics, sensitivity checks, and repeat runs to determine the robustness of assessment estimates to different weightings of data sets can be examined [3].

Stock assessments are only a part of the management regime in operation. Hence their evaluation needs to encompass whether the data and assessment methods available for management are sufficient to meet management objectives, and whether they, and the management plan itself, are robust to both the statistical uncertainty and the incomplete knowledge on factors such as stock identification and abundance, stock dynamics, and the impacts of environmental variability and trends [25]. The performance of assessments and management should therefore be simulation tested using approaches comparable

to those developed by the International Whaling Commission [26]. Where plans are inadequately precautionary, targets, constraints, the management procedure, or assessment can be modified and reevaluated. Alternatively, critical uncertainties in the process can be identified and further research initiated to reduce them before initiation of the plan.

Acknowledgments

Funding for David Maxwell has been provided by the Department for Environment, Food and Rural Affairs, UK.

References

- [1] FAO. (2004). State of World Fisheries and Aquaculture, at <http://www.fao.org/DOCREP/007/y5600e/y5600e00.htm> (accessed June 2007).
- [2] Hilborn, R. & Walters, C.J. (1992). *Quantitative Fisheries Stock Assessment*, Kluwer Academic Publishers, Boston.
- [3] Quinn II, T.J. & Deriso, R.B. (1999). *Quantitative Fish Dynamics*, Oxford University Press, New York.
- [4] Cotter, A.J.R., Burt, L., Paxton, C.G.M., Fernandez, C., Buckland, S.T. & Pan, J.-X. (2004). Are stock assessment methods too complicated? *Fish and Fisheries* **5**, 235–254.
- [5] Schnute, J.T. & Richards, L.J. (2001). Use and abuse of fishery models, *Canadian Journal of Fisheries and Aquatic Sciences* **58**, 10–17.
- [6] El-Shaarawi, A.H. & Piegorsch, W.W. (eds) (2002). *Encyclopedia of Environmetrics*, John Wiley & Sons, New York.
- [7] Simmonds, E.J. & MacLennan, D.N. (2005). *Fisheries Acoustics*, Blackwell Science, Oxford.
- [8] Stratoudakis, Y., Bernal, M., Ganas, K. & Uriarte, A. (2006). The daily egg production method: recent advances, current applications and future challenges, *Fish and Fisheries* **7**, 35–57.
- [9] Taylor, M.I., Fox, C.J., Rico, I. & Rico, C. (2002). Species-specific TaqMan probes for simultaneous identification of cod (*Gadus morhua* L.), haddock (*Melanogrammus aeglefinus* L.) and whiting (*Merlangius merlangus* L.), *Molecular Ecology Notes* **2**, 599–601.
- [10] Froese, R. & Pauly, D. (eds) (2006). *FishBase*, World Wide Web Electronic Publication, at <http://www.fishbase.org>, version (accessed May 2006).
- [11] Campana, S.E. (2001). Accuracy, precision and quality control in age determination, including a review of the use and abuse of age validation methods, *Journal of Fish Biology* **59**, 197–242.
- [12] Panfili, J., De Pontual, H., Troadec, H. & Wright, P.J. (eds) (2002). *Manual of Fish Sclerochronology*, Ifremer-IRD Coedition, Brest.
- [13] Hirst, D., Storvik, G., Aldrin, M., Aanes, S. & Bang Huseby, R. (2005). Estimating catch-at-age by combining data from different sources, *Canadian Journal of Fisheries and Aquatic Sciences* **62**, 1377–1385.
- [14] Francis, R.I.C.C., Harley, S.J., Campana, S.E. & Doering-Arjes, P. (2005). Use of Otolith Weight in Length-Mediated Estimation of Proportions at Age, *Marine and Freshwater Research* **56** pp. 735–743.
- [15] Hunter, J.R., Macewicz, B. & Sibert, J.R. (1986). The spawning frequency of skipjack tuna, *Katsuwonus pelamis*, from the South Pacific, *Fishery Bulletin* **84**, 895–903.
- [16] Polacheck, T., Eveson, J.P., Laslett, G.M., Pollock, K.H. & Hearn, W.S. (2006). Integrating catch-at-age and multiyear tagging data: a combined Brownie and Petersen estimation approach in a fishery context, *Canadian Journal of Fisheries and Aquatic Sciences* **63**, 534–548.
- [17] Hart, P.J.B. & Reynolds, J.D. (eds) (2002). *Handbook of Fish Biology & Fisheries*, Vol. 2, Blackwell Science, Oxford.
- [18] Smith, M.T. & Addison, J.T. (2003). Methods for stock assessment of crustacean fisheries, *Fisheries Research* **65**, 231–256.
- [19] Hollowed, A.B., Bax, N., Beamish, R., Collie, J., Fogarty, M., Livingston, P., Pope, J. & Rice, J.C. (2000). Are multispecies models an improvement on single-species models for measuring fishing impacts on marine ecosystems? *ICES Journal of Marine Science* **57**, 707–719.
- [20] Punt, A.E. & Hilborn, R. (1997). Fisheries stock assessment and decision analysis: the Bayesian approach, *Reviews in Fish Biology and Fisheries* **7**, 35–63.
- [21] Fournier, D.A., Hampton, J. & Sibert, J.R. (1998). MULTIFAN-CL: a length-based, age-structured model for fisheries stock assessment, with application to South Pacific albacore, *Thunnus alalunga*, *Canadian Journal of Fisheries and Aquatic Sciences* **55**, 2105–2116.
- [22] Xiao, Y.S., Punt, A.E., Millar, R.B. & Quinn, T.J. (eds) (2004). *Models in Fisheries Research: GLMs, GAMs and GLMMs*, Fisheries Research.
- [23] Caddy, J.F. & Mohn, R. (1995). *Reference Points for Fisheries Management*, FAO, Rome.
- [24] Patterson, K., Cook, R., Darby, C., Gavaris, S., Kell, L., Lewy, P., Mesnil, B., Punt, A., Restrepo, V., Skagen, D.W. & Stefansson, G. (2001). Estimating uncertainty in fish stock assessment and forecasting, *Fish and Fisheries* **2**, 125–157.
- [25] Motos, L. & Wilson, D.C. (eds) (2006). *The Knowledge Base for Fisheries Management*, Elsevier, Amsterdam.
- [26] Kirkwood, G.P. (1997). The revised management procedure of the International Whaling Commission, in *Global Trends: Fisheries Management*, E.K. Pikitch, D.D. Huppert & M.P. Sissenwine, eds, American Fisheries Society Symposium, pp. 91–99.

DAVID L. MAXWELL AND GRAHAM M.
PILLING

Sampling Sediments in the Marine Environment

Introduction

The marine environment covers more than 70% of the earth's surface and is a fundamental ecosystem which influences our weather patterns and from which we derive mineral resources and food. The seas are also a major transport system and provide significant leisure opportunities. However, over the years, we have deposited a considerable quantity of waste products into our seas and continue to use them in this context, although to a lessening extent. In addition, our marine environment is often the final location for hazardous substances brought to the seas by either rivers or through atmospheric deposition. As such, there is a clear need to understand our marine ecosystem. This requires that we make measurements – physical, chemical, and biological. In recent years, governments have become increasingly aware of the importance of our marine waters and require that we improve our understanding of how they are changing and what is causing such changes. The UK Government outlined its proposal for the marine environment in 2002 when it published *Safe-guarding Our Seas – A Strategy for the Conservation and Sustainable Development of Our Marine Environment* [1]. This presented the vision of clean, healthy, safe, productive, and biologically diverse oceans and seas and was followed, in 2005, by the devolved administration in Scotland presenting a supporting vision [2]. Also in 2005, the United Kingdom published *Charting Progress – An Integrated Assessment of the State of UK Seas* [3]. This brought together UK marine monitoring data, which was then used to describe and evaluate the current state of UK seas and also present some trends. Assessing “where we are” and “where we are going” is a feature of marine science which has been reiterated in the developing European Marine Strategy Directive [4]. However, to be able to fulfill the vision of clean, healthy, safe, productive, and biologically diverse seas and oceans, a definition of “clean” and of “safe” is required. In addition, there is a need to detect temporal trends, to be able to describe areas of the seabed, and to compare one area of the seabed with another and describe any difference such as concentrations of persistent

organic pollutants (POPs). There is also a requirement to be able to comment on whether changes in practices and procedures have resulted in improvements to our marine environment. Ultimately, this requires the collection and analysis of samples of water, sediment, and biota. Collecting sediment from the seabed is a costly process due to the need for a ship. In addition, it can be time consuming since sampling can be hindered by bad weather. Collecting a sediment sample from the seabed that is under 130 m of water requires specialized samplers, but to date we are, in effect, sampling blind since we cannot see the seabed being sampled nor how characteristic the sample is of the surrounding sediment. Constraints on sampling require that optimal sampling regimes, with a sound statistical basis, are developed. At the **Fisheries Research Services Marine Laboratory**, hazardous substances in marine sediments have been determined for some years (e.g. [5–8]). During this time, statistical tools have been used to develop sampling protocols, which have allowed differences in the concentration of hazardous substances between locations to be detected with a known statistical **power**, as well as optimize the number of samples that are required to provide data that is representative of the geographical area of interest.

Fjordic Sea Lochs

Fjordic sea lochs are characteristic features of the west coast of Scotland. They typically exhibit restricted circulation and accumulate fine-grained sediment with high organic carbon content. As such there is the potential to accumulate organic pollutants, especially those with an octanol water partition coefficient ($\text{Log } K_{ow}$) greater than 4. This includes polycyclic aromatic hydrocarbons (PAHs, e.g., pyrene, $\text{Log } K_{ow} 5.85$; benzo[a]pyrene, $\text{Log } K_{ow} 6.35$), chlorinated biphenyls (CBs, e.g., CB22, $\text{Log } K_{ow} 5.67$; CB153, $\text{Log } K_{ow} 6.92$) and polybrominated diphenyl ethers (e.g., 2,2',4-tribromodiphenyl ether (BDE 17), $\text{Log } K_{ow} 5.74$; 2,2',3,4,4',5',6-heptabromodiphenyl ether (BDE 183), $\text{Log } K_{ow} 8.27$). A study was undertaken to investigate the PAH composition of sediments from such sea lochs. This came immediately after a study into the PAH content of sediments from voes (small bays or narrow creeks) around the Shetland Islands impacted by the *Braer* oil spill of 1993. As part of this initial study, a sampling protocol was

2 Sampling Sediments in the Marine Environment

designed on the basis of previous data to ensure a 90% statistical power for detecting a 65% difference in total PAH concentrations between reference voes and those within the UK Food and Environment Protection Act 1985 (FEPA) Exclusion Zone set up following the incident [9]. The sampling protocol involved firstly generating 36 random numbers between 0 and 15. The vessel sampled the first station and then moved at 4 knots over a predetermined track. Samples were collected at the time intervals, in minutes, taken from the random number table; a 0-min time difference resulted in a duplicate sample being taken. This resulted in pseudorandom generated points. The 36 samples from each voe were later pooled into six pools of six using random selection. Each pooled sample consisted of aliquots (25 ± 1 g) from six individual samples and from this the various aliquots for the determination of particle size, total organic carbon (TOC), moisture, and hydrocarbon composition and concentration were taken. Although it was possible to draw conclusions from the resulting data using a combination of multivariate and regression techniques [10], the pooling of samples was shown to introduce considerable variation in the data. Consequently, when undertaking the more extensive survey of voes and coastal areas around the Shetland and Orkney Islands [11] and sea lochs on the west coast of Scotland [12], a similar sampling procedure was followed, but on this occasion only 14 samples were collected across each loch. However, all 14 were subject to individual analysis. Furthermore, the random number table that was generated was limited to numbers between 0 and 10. Again this sampling procedure was chosen based on previous data to ensure a 90% statistical power for detecting a 65% difference in total PAH concentration in the top 2 cm of the sediment, between different voes and sea lochs. Such a process allowed a detailed analysis of the hydrocarbons in sediments from both the voes and coastal areas around the northern isles and the west coast sea lochs. Plotting the mean log total PAH concentration (normalized for TOC) at each site, with 95% confidence limits on the mean, against the corresponding mean log TOC, showed clearly that there were significant differences between the different sea lochs and voes. These could be explained by consideration of anthropogenic inputs and thus the sampling regime was deemed to have been successful. In the case of the data from the west coast sea lochs, applying principal component analysis (PCA) to the median parent

and alkylated concentration (normalized to TOC), in each PAH group, expressed as proportions of the total median concentration, it was possible to characterize the lochs on the basis of the PAHs present and, ultimately, to come to conclusions about probable sources of these contaminants.

Open Sea Areas Associated with Offshore Oil and Gas Exploration and Production

Open sea areas represent a slightly different challenge since circulation is not restricted in the way it is within a sea loch. This is also, to some extent, the case for more open coastal locations. Historically, sediments in open sea areas have been sampled using a systematic grid approach in which several transects across an area are sampled at regular intervals. This is how the Fladen Ground, a potentially accumulative area of the North Sea which is also an area of intense oil production activity, has been surveyed [13]. Although such an approach provides maximum spatial coverage of an area for a given number of samples, is a relatively simple design to implement, provides regularly spaced or timed samples which provide for spatial and temporal correlations to be calculated, and can be implemented with little or no prior information about the site, such a design may not be as efficient as other designs if some prior information is available which can be used in the design. In addition, a grid design copes less well with missing stations (e.g., stations not collected in any year due to bad weather or because a new production platform has been constructed since the last survey). In an attempt to provide an adequate assessment of the spatial distribution of hydrocarbons in the Fladen Ground, the area was sampled using a stratified random sampling design. This allows a much simpler statistical analysis to be done. A description of the analysis can be found in [14]. In this case, the Fladen Ground was partitioned into 16 equal zones. Samples were selected by random number generation in each zone, the number of samples per zone being proportional to the far field (defined as areas more than 5 km away from multiple oil wells or 2 km from single oil wells) within each zone (Figure 1). Although used on this occasion, proportional allocation is not the only way of assigning the number of samples to a zone. An alternative is optimal allocation. This process

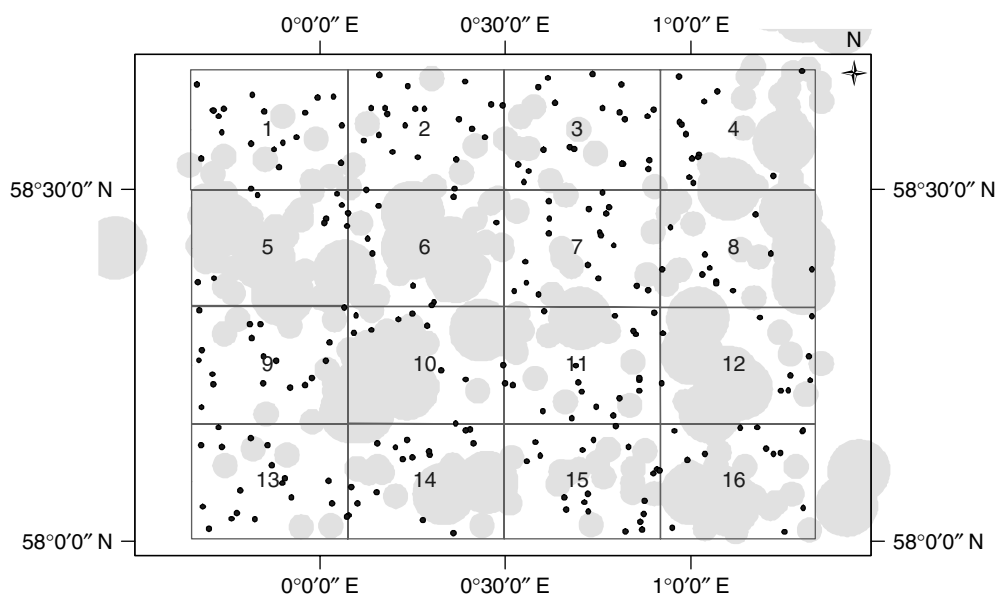


Figure 1 The 16 zones and the location of the stratified random sampling sites (black dots) in the far field areas (defined as areas more than 5 km away from multiple oil wells or 2 km from single oil wells) of the Fladen Ground, North Sea. Gray circles are near field areas (big circles are multiple oil wells and small circles are single wells)

allocates more samples to be taken from zones with higher variability in concentration between samples within a zone. Prior knowledge of the variability in the measured parameters is required. Ultimately, this method allocates more samples to zones with high variability and thereby maximizes the precision of the estimates of mean concentration for the whole far field area across the sampling area, given a fixed total **sample size**.

On the basis of proportional allocation, sediment samples were collected from the Fladen Ground by Day Grab from the FRV *Scotia* and analyzed for TOC, particle size, fluorescence, and hydrocarbons [15]. The mean total PAH concentration was calculated for each zone and this was then used to estimate the mean value over the whole far field ($\bar{y}_{st}$) using the equation:

$$\bar{y}_{st} = \frac{1}{A} \sum_{h=1}^L A_h \bar{y}_h \quad (1)$$

where A is the overall far field area, $\bar{y}_h$ is the zone mean value in zone h , A_h is the far field area in zone h and L is the total number of zones. The **standard error** for $\bar{y}_{st}$ can also be calculated. In this

particular survey, the mean PAH concentration within each of the 16 zones ranged from 53.9 ng g^{-1} dry weight to 206.6 ng g^{-1} dry weight. The mean over the whole far field was estimated to be 108.2 ng g^{-1} dry weight with a standard error for the mean over the whole far field area of 2.7 ng g^{-1} dry weight [15]. Differences in PAH concentrations between zones were also investigated using analysis of variance (**ANOVA**) and this highlighted significant differences between geographically grouped zones. Thus, this particular sampling technique provided an assessment of the spatial distribution of hydrocarbons across the Fladen Ground.

Coastal Sea Areas

The coastal zone receives direct and indirect inputs of hazardous substances. In general, estuarine and coastal areas are more highly contaminated than off-shore locations. The UK coastal zone is monitored under the UK Clean Seas Environmental Monitoring Programme (CSEMP) of the UK Marine Monitoring and Assessment Strategy, until recently known as the *UK National Marine Monitoring Programme*. This

4 Sampling Sediments in the Marine Environment

program includes the determination of the concentrations of hazardous substance in marine sediments and of any temporal trends in the concentrations of these compounds. The focus of the program has moved from repeated sampling at fixed stations to using sampling designs that are more immediately applicable to larger sea areas. The implied increase in the number of samples requiring collection and analysis has the potential to greatly increase the overall cost of the program and thus alternative strategies have been developed. Using both the Firth of Clyde and the Firth of Forth as test areas, a composite random sampling design was compared with the more traditional simple random sampling design. Composite sampling involves the physical combination of a number of homogenized samples to form a set of new (composite) samples on which the analyses are then performed. In this case, 25 sediment samples were collected by Day Grab from the FRV *Clupea* in the Firth of Clyde and the Firth of Forth (Figure 2) using a simple random sampling regime; the locations for the individual samples were selected using a random number generator. Five composite samples were prepared from the 25 samples collected using the simple random sampling design, by combining equal weights

of sediment from 5 of the individual samples. The five groupings of five samples were done on a random basis. ANOVA and multiple comparison tests at a **significance level** of 5% were used to detect significant differences among the means of % Total organic carbon (TOC), particle size analysis (PSA), PAH concentrations, and *n*-alkane concentrations. Spearman's rank **correlation** was used to determine nonparametric correlations among the parameters measured.

There was no significant difference between the mean concentration of PAHs from the 25 initial samples and the mean for the 5 composite samples for either the Firth of Clyde or the Firth of Forth. Indeed, for all parameters measured, there were no significant differences between the mean values of the composite random sampling and of the simple random sampling ($p < 0.05$; ANOVA) for both the Firth of Clyde and the Firth of Forth sediments.

The different sources of PAHs can be distinguished using specific PAH ratios. When this was done for the 25 individual sediments from both the Firth of Clyde and the Firth of Forth, a clear separation between the two locations was evident (Figure 3) with the geographical samples showing distinct clusters. When the process was repeated for the five

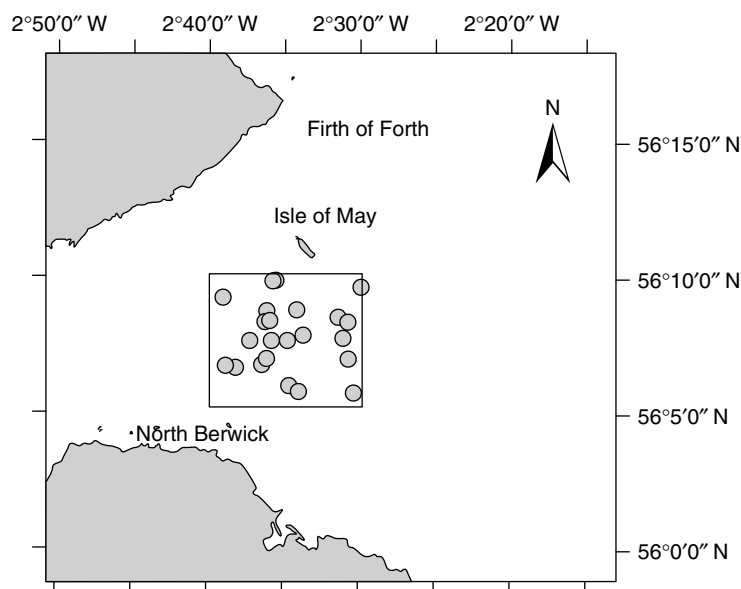


Figure 2 Map showing the Firth of Forth on the east of Scotland and the sampling area from which 25 samples were collected at random. The 25 samples were randomly grouped into five groups of five and composite samples prepared based on the grouping. All 25 samples were analyzed together with the five composite samples and the data compared

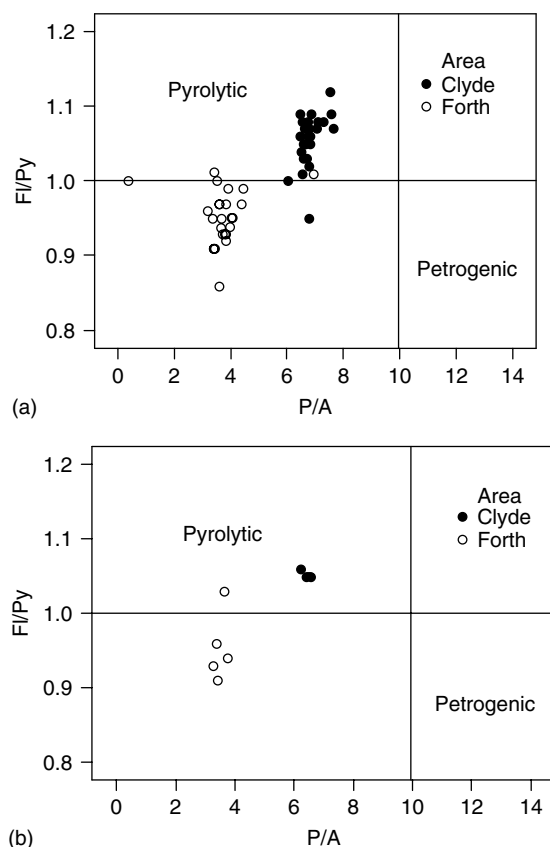


Figure 3 Fluoranthene/pyrene (Fl/Py) ratio *versus* the phenanthrene/anthracene (P/A) ratio for sediments from both the Firth of Clyde and the Firth of Forth collected using a simple random sampling scheme (a) and the composite random sampling scheme (b)

composite samples, even tighter clusters were evident (Figure 3). It was interesting to note that for some determinants, the compositing process reduced the variance between samples to such a degree that the batch-to-batch variance of the analytical methods used became an important part of the residual variance in the field data. Thus the composite random sampling design provided better representation and better precision of an estimated mean than the simple random sampling design and at less cost.

Conclusion

The use of simple statistical approaches can contribute significantly to the development of sampling regimes, which reduce the variance in the data, enable larger sea areas to be described, and detect differences in the concentrations of hazardous substances between sea lochs with a defined statistical power. Such sampling regimes are already being introduced into marine **monitoring programs** around the United Kingdom.

References

- [1] Department for Environment, Food and Rural Affairs (Defra). (2002). Safeguarding our seas – a strategy for the conservation and sustainable development of our marine environment, ISBN: 0 85521 005 1, p. 80.
- [2] Scottish Executive. (2005). Seas the opportunity – a strategy for the long term sustainability of Scotland's coasts and seas, ISBN: 0 7559 2693 5, p. 40.
- [3] Department for Environment, Food and Rural Affairs (Defra). (2005). *Charting Progress – An Integrated Assessment of the State of UK Seas*, Defra Publications, London, p. 120.
- [4] European Commission. (2006). EU marine strategy. The story behind the strategy, ISBN: 92 79 01810 8, p. 30, http://ec.europa.eu/environment/water/marine/pdf/eumarinestrategy_storybook.pdf.
- [5] Kelly, A.G. & Campbell, L.A. (1996). Persistent organochlorine contaminants in the Firth of Clyde in relation to sewage sludge input, *Marine Environmental Research* **41**, 99–132.
- [6] Webster, L., Mackie, P.R., Hird, S.J., Munro, P.D., Brown, N.A. & Moffat, C.F. (1997). Development of analytical methods for the determination of synthetic mud base fluids in marine sediments, *Analyst* **122**, 1485–1490.
- [7] Webster, L., Twigg, M., Megginson, C., Walsham, P., Packer, G. & Moffat, C. (2003). Aliphatic hydrocarbons and polycyclic aromatic hydrocarbons (PAHs) in sediments collected from the 100 mile hole and along a transect from 58° 58.32'N 1° 10.38'W to the inner Moray Firth, Scotland, *Journal of Environmental Monitoring* **5**, 395–403.
- [8] Edgar, P.J., Hursthouse, A.S., Matthews, J.E., Davies, I.M. & Hillier, S. (2006). Sediment influence on congener-specific PCB bioaccumulation by *Mytilus edulis*: a case study from an intertidal hot spot, Clyde Estuary, UK, *Journal of Environmental Monitoring* **8**, 887–896.
- [9] Topping, G., Davies, J.M., Mackie, P.R. & Moffat, C.F. (1997). The impact of the *Braer* spill on commercial fish and shellfish, in *The Impact of an Oil Spill in Turbulent*

- Waters: The Braer*, J.M. Davies & G. Topping, eds, The Stationery Office, Edinburgh, pp. 121–143.
- [10] Webster, L., McIntosh, A.D., Moffat, C.F., Dalgarno, E.J., Brown, N.A. & Fryer, R.J. (2000). Analysis of sediments from Shetland Islands voes for polycyclic aromatic hydrocarbons, steranes and triterpanes, *Journal of Environmental Monitoring* **2**, 29–38.
 - [11] Webster, L., Fryer, R.J., Dalgarno, E.J., Megginson, C. & Moffat, C.F. (2001). The polycyclic aromatic hydrocarbon and geochemical biomarker composition of sediments from voes and coastal areas in the Shetland and Orkney Islands, *Journal of Environmental Monitoring* **3**, 591–601.
 - [12] Webster, L., Fryer, R.J., Megginson, C., Dalgarno, E.J., McIntosh, A.D. & Moffat, C.F. (2004). The polycyclic aromatic hydrocarbon and geochemical biomarker composition of sediments from sea lochs on the west coast of Scotland, *Journal of Environmental Monitoring* **6**, 219–228.
 - [13] Russell, M., Webster, L., Walsham, P., Packer, G., Dalgarno, E.J., McIntosh, A.D. & Moffat, C.F. (2005). The effects of oil exploration and production in the Fladen Ground: composition and concentration of hydrocarbons in sediment samples collected during 2001 and their comparison with sediment samples collected in 1989, *Marine Pollution Bulletin* **50**, 638–651.
 - [14] Ahmed, A.S. (2005). *The development and application of statistical sampling regime to map hydrocarbons in sediments*, Ph.D. Thesis, The Robert Gordon University, Aberdeen, p. 242.
 - [15] Ahmed, A.S., Webster, L., Pollard, P., Davies, I.M., Russell, M., Walsham, P., Packer, G. & Moffat, C.F. (2006). The distribution and composition of hydrocarbons in sediments from the Fladen Ground, North Sea, an area of oil production, *Journal of Environmental Monitoring* **8**, 307–316.

COLIN F. MOFFAT, LYNDIA WEBSTER,
PAT POLLARD, ABDULWAHEED S. AHMED,
MARIE RUSSELL AND IAN M. DAVIES

Healthcare Performance Reports

Introduction

Reporting healthcare performance goes back at least to the pioneering work in the United States in 1754, when the Pennsylvania Hospital released tables of **health outcomes** tabulated by diagnostic groups [1]; and Florence Nightingale in England in 1860s, who presented analyses of hospital mortality rates [2]. Some information is relatively straightforward, such as waiting times and cleanliness; but assessing other dimensions of quality of care is typically problematic because health services are characterized by “heterogeneity, complexity, and uncertainty, and hence most performance indicators are ‘tin openers’ rather than ‘dials’ they do not give answers but prompt investigation and inquiry, and by themselves provide an incomplete and inaccurate picture” [3]. Iezzoni [4] points out that issues raised in the debate following Nightingale’s publication of hospital mortality rates echo those cited frequently about contemporary efforts at measuring provider performance, and include: risk adjustment; the need to examine mortality 30 days after admission rather than in-hospital deaths; and concerns over **data quality**, providers avoiding high-risk patients because of fears of public exposure, and the public’s ability to understand this information.

From the 1990s, technological advances in computing and web sites have resulted in an explosion in reporting of various kinds of healthcare performance. This has highlighted inadequacies and together with evidence of systems failures [5, 6] has further fuelled the demand for, and supply of, information on healthcare performance. There have been ambitious and controversial international comparisons of systems of healthcare [7, 8], assessments of countries’ own systems [9–14]. There is intense and polarized debate over the pros and cons of publication of performance data with different camps describing this as essential, desirable, inevitable, and potentially dangerous [15]. This article is about the impacts of healthcare performance reports within a country, which have been variously described as “report cards”, “consumer reports”, “public performance reports”, “provider/practice profiles”, “score cards”, “quality

assessment reports”, and “league tables” [16]; and confidential reports to individual clinicians as part of an audit process.

This article outlines issues raised by measuring performance, the models of impacts of information on healthcare performance, evidence of impacts, and concludes by discussing statistical implications.

Models of Impacts of Information on Healthcare Performance

To understand how reporting healthcare performance has an impact, we need a model of the links between the supply of this information and its impact on those who provide healthcare. Marshall *et al.* [15, p. 40] pointed out the lack of a commonly embraced rationale for the public release of performance data, which was believed to serve a wide variety of purposes and audiences. They identified three different models for professionals, public accountability, and markets which relate to three approaches [15, pp. 74–75]:

- **The “soft” approach**

In this model, clinicians are keen to improve their performance but hampered by a lack of comparative information, so a system is designed to provide information for **benchmarking** to enable those who want to improve to learn from the best. This approach is about quality improvement, focuses on success, and is voluntary.

- **The “hard” approach**

In this model the focus is on accountability and a requirement to meet minimum standards. This approach is compulsory and about quality assurance.

- **The market approach**

There are two models of the market approach because payments for healthcare are typically made not by patients, who are the consumers, but by third party payers (insurers or governments). These models are either corporate or consumers’ choices of health plans and corporate or patients’ choices of providers.

This article considers the impacts of reporting healthcare performance on four different groups: third party payers (either corporately or by consumers), patients, physicians, and providers. Reporting may also impact on the public perception of providers and inform and be part of policies for the regulation and management of systems of healthcare, in which

2 Healthcare Performance Reports

the role of the media is crucial. But if we are to understand the impact of reporting on healthcare performance, then it must be observed in these four groups: the other impacts are means to those ends.

Impacts

Third Party Payers

A problem for third party payers in the corporate model of the market approach is that for acute hospital care there is enormous heterogeneity in the scope of these services and only limited information for some services. In the United States, these payers were not always aware of available information: on clinical outcomes, consumer satisfaction data, and Health Plan Employer Data and Information Set (HEDIS) reports, that include performance measures of cancer, heart disease, smoking, asthma, and diabetes and a standardized survey of consumers' experiences. And those that were aware of this information found that it did not help them make decisions [16, 17]. Hence their early interest in this information soon waned and they focused on costs or gross indicators of quality through accreditation rather than detailed comparative information [16]. Consumers in choosing between health plans in the United States experienced further complications and difficulties: with lack of understanding of basic differences between types of insurance or how managed care plans could influence quality of care [18–21]. In the United Kingdom, the evidence is that District Health Authorities, the dominant third party payers, made little use of comparative information for contracting, but information on waiting times was used effectively by general practitioners (GPs), who had opted for their practices to become payers in the form of GP fundholders [17, 22–24]. The comparative effectiveness of GP fundholders to use simple information, illustrates that even if larger corporate payers had more adequate information on quality, they would face large transaction costs in using this for contractual purposes: it would no longer make sense to contract *en bloc* with providers: it would be necessary to go down to the level of clinical teams. The attraction of the market model of patient choice is that this promises a way of sending detailed signals at the level of clinical teams.

Patients

Reporting information to patients may be seen as part of the models that underlie the hard and the market approach. Both models assume that patients are aware of, can access, and understand this information; and the market model assumes they can use it to make choices [25].

Publication of surgical mortality rates for cardiac surgery in England [26] followed the landmark case of failure of self-regulation by the medical profession to take corrective action over high-mortality rates following pediatric cardiac surgery at the Bristol Royal Infirmary [27]. The Kennedy Report into the Bristol case observed that “little if any information was available to the parents or to the public. Such information that was given to parents was often partial, confusing, and unclear” [27, p. 3] and hence argued that “continued secrecy and anonymity, is no longer a real option”. The report in recommending that audit data be published recognized that these data are complex and hard to understand; may mislead the public, who might draw unwarranted conclusions; and may result in healthcare professionals being unfairly criticized [27, p. 398]. The report called for a new compact between the community and its hospitals in which the public must accept that the price of information is a considered and responsible reaction to it [27, p. 398]. This is different from the market approach as what is being advocated here is highly specific information relevant to patients, or, in this case, their relatives. But after that recommendation was made, when mortality rates of pediatric cardiac surgery at different centers were published [28], these were criticized as misleading and have been subsequently disputed [29].

The market model of patient choice requires designing a system in which providers' incomes depend on the numbers of patients they treat and assumes that choices made using information on quality of care will offer signals in a market which will drive providers to react. In this, model patients require timely specific information on the performance of the clinicians who supply diagnosis and treatment of their conditions. They thus encounter one of the problems of third party payers: that information is in the form of “tin openers” and coverage is partial, and this problem is exacerbated by patients' interest in timeliness. Second, a basic tenet of economics is that the consumers are skilled buyers of goods

and services that they buy repeatedly, so patients with chronic conditions can become expert patients and skilled users of the services that are available; but a patient with an acute condition will be in a very different position. The starting point of health economics is the asymmetry of information between such a patient and provider, which means that property rights over consumption are not exercised solely by the patient who relies on a provider to make a diagnosis, and the patient may from then on feel that choices are circumscribed. Such patients are often anxious and want to make a decision quickly: making choices over treatment for breast cancer is, e.g., rather different from choosing a used car. This would be difficult enough if perfect information were available, the information that is typically available is difficult to interpret, incomplete, partial, and out of date.

Hence we would expect patients to have problems in understanding and using publicly reported information on healthcare performance in either the market model or for the purposes of public accountability. And that is exactly what the evidence from the United States shows [16, 25]: when this information is published only a minority are aware of it, of those most do not understand it (including whether high or low rates of an indicator reflect good performance), trust it, or use it (with problems with timely access, and lack of genuine choice); and evaluations of later developments that addressed many of these potential barriers failed to demonstrate significant or sustained public interest. Keogh *et al.* [26] summarized findings from the publication of outcomes from cardiac surgery to conclude that what patients really want in the United Kingdom “is an operation in a hospital close to home and as soon as possible”. Furthermore, recent focus group data indicated antagonism from some members to public reporting in the accountability model, as they saw this as a punitive tool used by politicians to punish hardworking professionals [16].

Physicians

A crucial distinction in examining reporting information for clinicians is between the “hard” and “soft” approach, whether this is intended for quality improvement or assurance, which is in turn related to the issue of publication. Evidence shows that physicians were more aware than patients, of information,

which appears to have been directed at patients as consumers, but physicians too made little use of it in making referrals and distrusted and, in some cases, attempted to discredit the data [16, 24, 29]. There are concerns that publication of surgical mortality rates results in a defensive response by surgeons in which they refuse to operate on high-risk patients [15]; and one effect of reporting annual risk-adjusted mortality rates of coronary artery bypass graft (CABG) surgery in Pennsylvania was that over half the cardiologists reported increased difficulty in finding surgeons willing to perform CABG surgery in severely ill patients who required it, and over half the cardiac surgeons reported that they were less willing to operate on such patients [30].

The soft approach seeks to organize collection and publication of data in ways that enable clinicians to learn how to improve. Examples include the Northern New England Cardiovascular Disease Study Group [31], and national clinical audits in the United Kingdom, for example, for stroke [32] and myocardial infarction [33]. The Northern New England Cardiovascular Disease Study Group included feedback of outcome data, training in continuous quality improvement techniques, and site visits to other medical centers and the main outcome was a statistically significant reduction by nearly a quarter in the hospital mortality rate. In 1995, the key factors that were identified as contributing to the group’s success included: strong clinical leadership, regionally and at each institution; confidentiality of data, which were analyzed and returned in a timely fashion; and an organized forum for discussion [34].

Providers

Marshall *et al.* [16] pointed to a growing body of evidence to show that providers are the most sensitive of the four groups we have identified here to the publication of information. Despite the fact that providers’ common initial public reaction to reports of poor performance is to dispute their validity, in private they use the published information to focus on the quality of the data and care they provide. We give six important examples, of which all but one appear to have had an impact.

- Annual reporting of data on risk-adjusted mortality following CABG surgery by hospital and surgeon in New York state was the first

such program in the United States, beginning in 1989, and is now the most long lived. Chassin [35] observes that, between 1989 and 1992 risk-adjusted mortality fell by over 40% statewide in New York; which had, by 1992 the lowest risk-adjusted mortality rate of any state in the nation, and the most rapid rate of decline of any state with below-average mortality. He points out that these improvements happened because individual hospitals and cardiac surgery programs used the data to make specific changes in the way they provided care to CABG patients. Neither market forces nor managed care companies had an impact and patients did not avoid high-mortality hospitals nor seek out those with low mortality. As mentioned above, however, one reservation about the observed benefits of public reporting of surgical mortality rates is that, in private, cardiac surgeons turn away the sickest and most severely ill patients [25].

- Scotland's Clinical Resource and Audit Group (CRAG) [36] pioneered in Europe, in 1994, the publication of annual Reports of clinical outcome indicators for use by the public and doctors. CRAG's reports aimed to provide a benchmarking service for clinical staff by publishing comparative clinical outcome indicators across Scotland. Performance management was seen as a matter for local action: there were no external national pressures to act on CRAG's indicators other than publication. The main conclusions of two evaluations [36, pp. 223–229, 37] were that in Scottish trusts these indicators were rarely used to inform internal quality improvement or used externally to identify best practice with the reasons being: the information was regarded as out of date and lacking credibility, because of problems over the quality of the data and the lack of an adjustment for casemix; there was low awareness, because of little attention from the media and limited dissemination within each provider; the results were difficult to interpret, had limited coverage and were at too high a level of aggregation.
- From 1995, the Veterans Administration (VA) in the United States introduced **reengineering** to improve quality of care: routine performance measurements for high-priority conditions and performance contracts held managers accountable for meeting improvement goals. Data gathering

and monitoring were performed by an independent agency. Performance data were made public and were widely distributed within the VA, among key stakeholders such as veterans' service organizations, and among members of Congress [38].

- Hibbard *et al.* [39] undertook a controlled trial of the effects of reporting information on quality of care in 2001 between three sets of hospitals in Wisconsin: for the public-report set this information was published; the private-report set were supplied information privately (but this was not published), the control set received no report. The information included two summary indices of adverse events (deaths and complications) occurring within the broad categories of surgery and nonsurgery, and indices in three individual clinical areas. Hospitals were rated as better than expected, as expected, or worse than expected. The public-report hospitals were significantly more negative than the other two groups about the public report: its validity, value for quality improvement, and appropriateness for use by the public; but their efforts to improve quality were significantly greater than in hospitals given only private reports or provided with no information. None of the hospitals saw such information as having an impact on their market share (or "one like it" in the case of the private- and no-report hospitals), but the public-report hospitals saw this as having an impact on their public image.
- In England, from 2001 to 2005, annual performance ratings were published for NHS trusts in England [40–42]. This process gave each trust a rating from zero to three stars. Trusts that failed against a small number of key targets (dominated by waiting times) were at risk of being zero rated and their chief executives at risk of losing their job; trusts that performed well achieved three stars and were eligible for benefits from "earned autonomy". This system resulted in two kinds of outcomes: dramatic reductions in long waiting times in England, in line with the targets, which was not observed in the other countries of the United Kingdom; and various kinds of gaming, which included poor performance in domains where performance was not measured, what has been described as "hitting the target and missing the point", and exploiting ambiguity in reporting of data or fabrication in what was reported.

Statistical Implications

The survey of the various means of reporting information on healthcare performance suggests that this needs to be targeted at providers and be credible and timely; and also there ought to be widespread awareness of the results. This final section considers the statistical implications of translating data into information in these ways: the attraction and difficulties of producing individual indicators from measures of outcome and process, presentation of results including aggregation.

Outcome measures are of intrinsic interest and public appeal, and can reflect all aspects of care. But for an outcome to be a valid measure of quality, it must be closely related to a process of care that can be manipulated [43]. The data that are collected routinely on outcomes are typically limited to mortality and readmission. For hospital mortality at 30 days, indicators based on deaths in hospital may be inadequate and more reliable results are derived by linking data on hospital discharges to registry data on mortality [44]. Using outcomes to assess performance raises issues of risk rating: the data required may also not be collected routinely [45], and, even where data are available, there are serious methodological problems as it is not possible to predict with certainty which patients will recover from surgery. There have been particular problems with the high-risk population for CABG [46], which matters given marked variations in the proportion of individual surgeons' practices that are high risk. Using a single risk-adjusted number to summarize a surgeon's results runs the risk of lending "a level of spurious credibility to an analysis that does not take into account the impact of influences that are not patient related" [46]. One way of handling this problem, as the low-risk group was relatively homogeneous between surgeons, is to produce comparative analysis based on low-risk cases without the need for further risk adjustment [46]. This also has the advantage of making it clear that these indicators are "tin openers" [26]. Although measures of process lack the direct appeal of outcome measures, they can be more sensitive to differences in the quality of care [47], and are valid without the need for risk rating. There are issues over requirements for detailed clinical data that are currently found only in medical records [47], which may be more feasible with developments of electronic records [48].

Waiting times are important measures that are "gauges". There are statistical issues over how best to describe providers' performance. If, e.g., targets are in terms of the longest waiting times only, this obviously puts pressure on the providers with times that are above the target level to do better, but also creates a perverse incentive for those doing better than the target to allow their performance to deteriorate to the standard, and more generally to crowd performance towards the target [49]. Such effects can unintentionally penalize agents with exceptionally good performance but a few failures, while rewarding those with mediocre performance crowded near the target range.

Funnel plots that plot performance in terms of scale with error limits are based on sound statistical theory and are easy to understand [50] are a good way of presenting results for individual indicators. Aggregating performance of complex organizations into a single score is misleading as this fails to capture their complexity, and in the case of star ratings means it was difficult for a provider to understand why its performance changed [51] and league tables that simply rank organizational performance are easy to understand but unsound [52, 53]. What is required is developing ways of reporting data to give signals at a high level and organized to enable the investigation of those signals by drilling down.

References

- [1] Lansky, D. (1998). Perspective: measuring what matters to the public, *Health Affairs* 17(4), 40–41.
- [2] Nightingale, F. (1863). *Notes on Hospitals*, 3rd Edition, Longman Green, London, pp. 3–5, <http://books.google.com/books?id=FJhN-SqxUawC&vid=OCLC02023855&dq=florence+nightingale+hospital+notes&jtp=I>.
- [3] Carter, N., Klein, R. & Day, P. (1995). How organisations measure success, *The use of Performance Indicators in Government*, Routledge, London, p. 49.
- [4] Iezzoni, L.I. (1996). 100 apples divided by 15 red herrings: a cautionary tale from the mid-19th century on comparing hospital mortality rates, *Annals of Internal Medicine* 124(12), 1079–1085.
- [5] Abbasi, K. (1998). Butchers and grocers, *British Medical Journal* 317, 1599.
- [6] Institute of Medicine. (2001). *Crossing the Quality Chasm: A New Health System for the Twenty-first Century*, National Academy Press, Washington, DC.
- [7] WHO. (2000). *The World Health Report 2000*, WHO, Geneva.
- [8] Almeida, C., Braveman, P., Gold, M.R., Szwarcwald, C.L., Riberio, J.M., Miglionico, A., Millar J., Porto, S.,

- Costa, N. & Rubio, V. (2001). Methodological concerns and recommendations on policy consequences of the World Health Report 2000, *Lancet* **357**, 1692–1697.
- [9] RAND. (2006). *The First National Report Card on Quality of Health Care in America*, Rand, Santa Monica.
- [10] Leatherman, S. & Sutherland, K. (2003). *The Quest for Quality in the NHS: A Mid-Term Evaluation of the Ten-Year Quality Agenda*, The Stationery Office, London.
- [11] Leatherman, S. & Sutherland, K. (2005). *The Quest for Quality in the NHS: A Chartbook on Quality of Care in the UK*, Radcliffe Publishing, Oxford.
- [12] Leatherman, S. & McCarthy, D. (2004). *Quality of Care for Children and Adolescents: A Chartbook*, Commonwealth Fund, New York.
- [13] Leatherman, S. & McCarthy, D. (2005). *Quality of Health Care for Medicare Beneficiaries: A Chartbook*, Commonwealth Fund, New York, Vol. 815.
- [14] Leatherman, S. & McCarthy, D. (2002). *Quality of Health Care in the United States: A Chartbook*, Commonwealth Fund, New York.
- [15] Marshall, M., Shekelle, P., Brook, R. & Leatherman, S. (2000). *Dying to Know: Public Release of Information about Quality of Care*, Nuffield Trust, London.
- [16] Marshall, M., Shekelle, P.G., Davies, H.T.O. & Smith, P.C. (2003). Public reporting on quality in the United States and the United Kingdom, *Health Affairs* **22**(3), 134–148.
- [17] Hibbard, J.H., Jewett, J.J., Legnini, M.W. & Tusler, M. (1997). Choosing a health plan: do large employers use the data? *Health Affairs* **16**(6), 172–180.
- [18] Isaacs, S.L. (1996). Consumer's information needs: results of a national survey, *Health Affairs* **15**(4), 31–41.
- [19] Hibbard, J.H. & Jewett, J.J. (1997). Will quality report cards help consumers? *Health Affairs* **16**(3), 218–228.
- [20] Hibbard, J.H., Harris-Kojetin, L., Mullin, P., Lubalin, J. & Garfinkel, S. (2000). Increasing the impact of health plan report cards by addressing consumers' concerns, *Health Affairs* **19**(5), 138–143.
- [21] Hibbard, J.H., Slovic, P., Peters, E., Finucane, M.L. & Tusler, M. (2001). Is the informed-choice policy approach appropriate for Medicare beneficiaries? *Health Affairs* **20**(3), 199–203.
- [22] Audit Commission. (1996). *What the Doctor Ordered*, HMSO, London, pp. 19–22.
- [23] Glennerster, H., Matsaganis, M., Owens, P. & Hancock, S. (1994). *Implementing GP Fundholding*, Open University Press, Buckingham.
- [24] Bernard, D. (1997). Effect of fundholding on waiting times: database study, *British Medical Journal* **315**, 290–292.
- [25] Werner, R.M. & Asch, D.A. (2005). The unintended consequences of publicly reporting quality information, *Journal of the American Medical Association* **293**, 1239–1244.
- [26] Keogh, B., Spiegelhalter, D., Bailey, A., Roxburgh, J., Magee, P. & Hilton, C. (2004). The legacy of Bristol: public disclosure of individual surgeons' results, *British Medical Journal* **329**, 450–454.
- [27] Secretary of State for Health. (2001). *Learning from Bristol – Report of the Public Inquiry into Children's Heart Surgery at the Bristol Royal Infirmary (the Kennedy Report)* (CM 5207(1)), The Stationery Office, London, <http://www.bristol-inquiry.org.uk/final-report/>.
- [28] Aylin, P., Bottle, A., Jarman, B. & Elliott, P. (2004). Paediatric cardiac surgical mortality in England after Bristol: descriptive analysis of hospital episode statistics 1991–2002, *British Medical Journal* **329**, 825.
- [29] Keogh, B. (2005). Surgery for congenital heart conditions in Oxford, *British Medical Journal* **330**, 319–320.
- [30] Schneider, E.C. & Epstein, A.M. (1996). Influence of cardiac-surgery performance reports on referral practices and access to care – a survey of cardiovascular specialists, *New England Journal of Medicine* **335**, 251–256.
- [31] Royal College of Physicians. (2005). *National Sentinel Stroke Audit 2004*, Royal College of Physicians, London.
- [32] Royal College of Physicians. Fourth Public Report Prepared on behalf of the MINAP Steering Group. (2005). *How the NHS Manages Heart Attack*, Royal College of Physicians, London.
- [33] O'Connor, G.T., Plume, S.K., Olmstead, E.M., Morton, J.R., Maloney, C.T., Nugent, W.C., Hernandez Jr, F., Clough, R., Leavitt, B.J., Coffin, L.H., Marrin, C.A., Wennberg, D., Birkmeyer, J.D., Charlesworth, D.C., Malenka, D.J., Quinton, H.B. & Kasper, J.F., The Northern New England Cardiovascular Disease Study Group. (1996). A regional intervention to improve the hospital mortality associated with coronary artery bypass graft surgery, *Journal of the American Medical Association* **275**(11), 841–846. <http://jama.ama-assn.org/cgi/content/abstract/275/11/841>.
- [34] Malenka, D.J. & O'Connor, G.T., Northern New England Cardiovascular Study Group. (1995). A regional collaborative effort for CQI in cardiovascular disease, *Joint Commission Journal on Quality Improvement* **21**(11), 627–633.
- [35] Chassin, M.R. (2002). Achieving and sustaining improved quality: lessons from New York State and cardiac surgery, *Health Affairs* **21**(4), 40–51.
- [36] CRAG. (2002). *Clinical Outcome Indicators*, A report from the Clinical Outcomes Working Group, Scottish Executive, Edinburgh, www.show.scot.nhs.uk/crag.
- [37] Mannion, R. & Goddard, M. (2001). Impact of published clinical outcomes data: case study in NHS hospital trusts, *British Medical Journal* **323**, 260–263.
- [38] Jha, A.K., Perlin, J.B., Kizer, K.W. & Dudley, R.A. (2003). Effect of the transformation of the veterans affairs health care system on the quality of care, *New England Journal of Medicine* **348**(22), 2218–2227.
- [39] Hibbard, J.H., Stockard, J. & Tusler, M. (2003). Does publicizing hospital performance stimulate quality improvement efforts? *Health Affairs* **22**(2), 84–94.
- [40] Bevan, G. & Hood, C. (2006). Have targets improved performance in the English NHS? *British Medical Journal* **332**, 419–422.

-
- [41] Bevan, G. (2006). Setting targets for health care performance: lessons from a case study of the English NHS, *National Institute Economic Review* **197**, 67–79.
- [42] Bevan, G. & Hood, C. (2006). What's measured is what matters: targets and gaming in the English public health care system, *Public Administration* **84**(3), 517–39.
- [43] Chassin, M. (1996). Academic quality improvement: new medicine in old bottles, *Quality Management in Health Care* **4**(4), 40–46.
- [44] Seagroatt, V. & Goldacre, M.J. (2004). Hospital mortality league tables: influence of place of death, *British Medical Journal* **328**, 1235–1236.
- [45] Craig, M., Price, C., Backhouse, A. & Bevan, G. (1990). Medical audit and resource management: lessons from hip fractures, *Financial Accountability and Management* **6**, 285–294.
- [46] Bridgewater, B., Grayson, A.D., Jackson, M., Brooks, N., Grotte, G.J., Keenan, D.J.M., Millner, R., Fabri, B.M. & Jones, M. (2003). Surgeon specific mortality in adult cardiac surgery: comparison between crude and risk stratified data, *British Medical Journal* **327**, 13–17.
- [47] Mant, J. (2001). Process versus outcome indicators in the assessment of quality of health care, *International Journal for Quality in Health Care* **13**, 475–480.
- [48] Palmer, R.H. (1997). Process-based measures of quality: the need for detailed clinical data in large health care databases, *Annals of Internal Medicine* **127**(5 Part 2), 733–738.
- [49] Bird, S.M., Cox, D., Farewell, V.T., Goldstein, H., Holt, T. & Smith, P.C. (2005). Performance indicators: good, bad, and ugly, *Journal of the Royal Statistical Society A* **168**(1), 1–27.
- [50] Spiegelhalter, D.J. (2004). Funnel plots for comparing institutional performance, *Statistics in Medicine* **24**(8), 1185–1202.
- [51] Spiegelhalter, D.J. (2005). The mystery of the lost star, *Significance* **2**, 150–153.
- [52] Goldstein, H. & Spiegelhalter, D.J. (1996). League tables and their limitations: statistical issues in comparisons of institutional performance, *Journal of the Royal Statistical Society A* **159**, 385–443.
- [53] Marshall, E.C. & Spiegelhalter, D.J. (1998). Reliability of league tables of in vitro fertilisation clinics: retrospective analysis of live birth rates, *British Medical Journal* **316**, 1701–1704.

RICHARD G. BEVAN

Disease Mapping

Disease Mapping basics

Disease mapping is the analysis of epidemiological or public health data which is geo-referenced. Typically the data arises in two forms: either (a) the residential addresses of people affected by disease are known, or (b) counts of disease observed within arbitrary small areas such as census tracts, zip codes, or postcodes are available. The locational information is used in the analysis usually to make inferences about spatial health effects.

Often hypotheses of interest in disease mapping focus on whether the residential address of people affected by disease yields insight into the etiology of the disease or, in a public health application, whether adverse environmental health hazards exist locally within a region (as exemplified by local increases in disease risk). For example, in a study of the relationship between malaria and diabetes in Sardinia a strong negative relationship has been found. This relation had a spatial expression and the geographical distribution of malaria was important in generating explanatory models for the relation [1]. In public health practice it is of considerable importance to be able to assess whether localized areas that have larger than expected numbers of cases of disease have in common any underlying environmental cause. Here, spatial evidence of a link between cases and a source is fundamental in the analysis. Evidence such as a decline in risk with distance from the putative source of hazard or elevation of risk in a preferred direction is important in this regard (see e.g., Lawson *et al.* [2], Lawson [3, Chapter 7], and Elliott *et al.* [4]).

There are four main areas where statistical methods have seen development: *disease map reconstruction*, *disease clustering*, *ecological analysis*, and *disease map surveillance*. The term *disease mapping* is used in different contexts to either signify the group of methods that are applied to disease maps, or more specifically the **estimation** of relative risk from disease maps (*disease map reconstruction*). Hence there can be some confusion in the use of this term. It is appropriate to consider some common themes or issues, which arise in all areas of the subject.

Basic Definitions and Models

A fundamental feature of data available for analysis is that it is usually discrete (either in the form of a point process or counting process), and the cases of concern arise from within a local human population, which varies in spatial density and in susceptibility to the disease of interest. Hence any model or test procedure must make allowance for this background (nuisance) population effect. The background population effect can be allowed for in a variety of ways. For count data, it is commonplace to obtain expected rates for the disease of interest based on the age–sex structure of the local population, and some crude estimates of local relative risk are often computed from the ratio of observed to expected counts (e.g., standardized mortality ratios (SMR's) or standardized incidence ratios). For case event data, expected rates are not available at the resolution of the case locations and the use of the spatial distribution of a control disease has been advocated. In that case, the spatial variation in the case disease is compared to the spatial variation in the control disease. A major issue in this approach is the correct choice of control disease. It is important to choose a control that is matched to the age–sex structure of the case disease but is unaffected by the feature of interest. For example, in the analysis of cases around a putative health hazard, a control disease should not be affected by the health hazard. Counts of control disease cases could also be used instead of expected rates when analyzing count data. Figures 1 and 2 display case event and control data maps for a region of the United Kingdom for a fixed time period. Figure 3 displays a typical count data example.

Case Event Data

Case event locations often represent residential addresses of patients, and the patients are from a heterogeneous population that varies in both spatial density and susceptibility to disease. A heterogeneous **Poisson process** model is often assumed as a starting point for further analysis. The focus of interest for making inference regarding parameters describing excess risk lies in a relative risk function, which is included in the first-order intensity of the Poisson process.

It is possible that population or environmental heterogeneity may be unobserved in the data set.

2 Disease Mapping

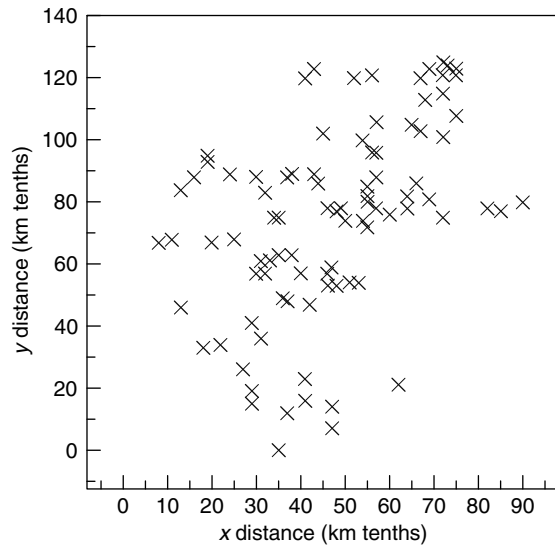


Figure 1 Distribution of death certificate addresses of gastric and esophageal cancer in Arbroath, Scotland, UK (1966–1976)

This could be because either the population background hazard is not directly available or the disease displays a tendency to cluster (perhaps due to unmeasured covariates). The heterogeneity could

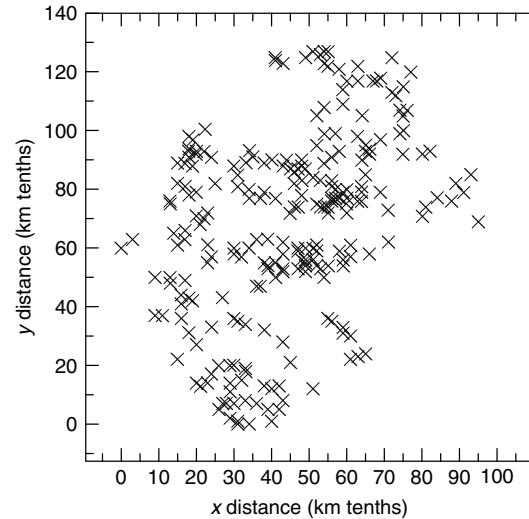


Figure 2 Control distribution: distribution of death certificate addresses for lower body cancers Arbroath, UK (1966–1976)

be spatially correlated, or it could lack **correlation**, in which case it could be regarded as a type of “overdispersion”. One can include such unobserved heterogeneity within the framework of conventional models as a random effect.

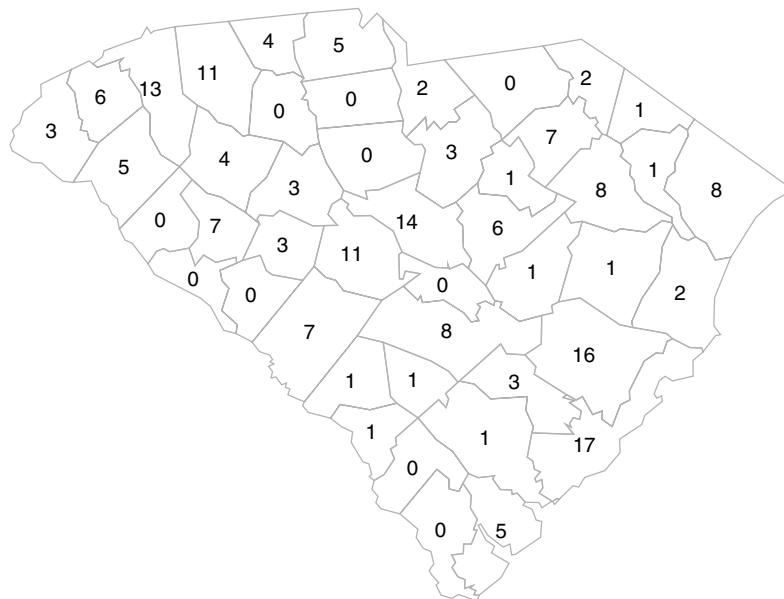


Figure 3 Distribution of counts of congenital deaths within the counties of South Carolina, USA 1990

Count Data

Considerable literature has developed concerning the analysis of count data in spatial analysis of disease (e.g., see reviews in [3–6]).

The usual model adopted for the analysis of region counts assumes that the counts are independent Poisson random variables with parameter 2_i in the i th region. This model may be extended to include unobserved heterogeneity between regions by introducing a prior distribution for the log relative risks ($\log 2_i$). Incorporation of such heterogeneity has become a common approach and the Besag, York, and Mollié (BYM) model is now a standard model. A full **Bayesian analysis** using this model is available on WinBUGS. A recent overview of Bayesian methods in disease mapping is found in [7].

References

- [1] Bernardinelli, L., Pascutto, C., Montomoli, C., Komakec, J. & Gilks, W.R. (1999). Ecological regression with errors in covariates: an application, in *Disease Mapping and Risk Assessment for Public Health*, A.B. Lawson, A. Biggeri, D. Boehning, E. Lesaffre, J.-F. Viel & R. Bertollini, eds, John Wiley & Sons.
- [2] Lawson, A.B., Bohning, D., Lesaffre, E., Biggeri, A., Viel, J.F. & Bertollini, R. (1999). *Disease Mapping and Risk Assessment for Public Health*, John Wiley & Sons.
- [3] Lawson, A.B. (2006). *Statistical Methods in Spatial Epidemiology*, 2nd Edition, John Wiley & Sons, New York.
- [4] Elliott, P., Wakefield, J., Best, N. & Briggs, D. (2000). *Spatial Epidemiology: Methods and Applications*, Oxford University Press, London.
- [5] Lawson, A. B. & Williams, F. L, R. (2001). *An Introductory Guide to Disease Mapping*, Wiley, New York.
- [6] Lawson, A. b., Browne, W. & Vidal Rodeiro, C. L. (2003). *Disease Mapping with WinBUGS and MLwiN*, Wiley, New York.
- [7] Lawson, A.B. & Banerjee, S. (2007). Bayesian spatial analysis, in *Handbook of Spatial Analysis*, S. Fotheringham & P. Rogerson, eds, Sage Publications, New York.

ANDREW B. LAWSON

Disease Surveillance

Disease Surveillance Basics

Disease surveillance is the detection of aberrations in health or disease, usually as they arise. This definition stresses both the unusual nature of the disease event and the importance of temporal change in surveillance. How “unusual” an event or a sequence of events is, is of course an issue in the design of any surveillance system. The Centers for Disease Control (CDC) defines surveillance as:

the ongoing, systematic collection, analysis, and interpretation of health data essential to the planning, implementation, and evaluation of public health practice, closely integrated with the timely dissemination of these data to those who need to know. The final link of the surveillance chain is the application of these data to prevention and control. A surveillance system includes a functional capacity for **data collection**, analysis, and dissemination linked to public health programs [1].

This definition stresses the collection, analysis, and dissemination of data in a timely manner. Hence, it is a very broad definition that places an emphasis on public health needs. However it is possible to distinguish two basic types of surveillance that play different roles in public health activities.

First, *retrospective* surveillance concerns the collection of historical data on disease occurrence and its examination. The purpose of such analysis may be to inform decision makers as to temporal or spatial trends and other features of disease behavior. This form of surveillance is closely associated with classical epidemiological analysis, and differs mainly in its focus on public health needs.

Second, *prospective* surveillance is the on-line or active examination of disease data to discover changes in disease at or close to the time of occurrence. In this case, monitoring of disease occurrence is done “as data arrive” so that decisions can be made concerning outbreaks of disease. The importance of the latter form of surveillance has been heightened with the recent rise in terrorism and the potential threat to health from bioterrorism.

The release of toxic or highly infectious agents into a population would be of grave concern in this context and so it is now important that fast and

accurate surveillance of diseases be undertaken to detect changes as early as possible.

The design of disease surveillance systems must consider some of the following issues:

1. Early detection

It is important to detect changes in disease incidence as early as possible. For example, for certain infectious diseases it may take up to 7 days to receive laboratory confirmation of cases. However such a delay may be unacceptable if a serious event had occurred. Hence ways to speed up detection could be important.

2. Syndromic methods

The use of ancillary information to speed up detection of population health aberrations. A formal definition is given by Sosin [2]

...the public health term Syndromic Surveillance has been applied to systematic and ongoing collection, analysis and interpretation of data that precede diagnosis (e.g., laboratory test requests, emergency department chief complaint, ambulance response logs, prescription drug purchases, school or work absenteeism, as well as signs and symptoms recorded during acute care visits) and that can signal a sufficient probability of an outbreak to warrant public health investigation.

Hence, often, **covariate** information could be useful in helping to establish the character of an outbreak. Figure 1 displays the **time series** of pharmaceutical sales and gastrointestinal disease reporting for a Canadian example (see for further examples Lawson and Kleinman [3]).

3. Sensitivity and specificity of detection methods

The **calibration** of a surveillance system is very important in that false aberration alarms (false positives) could lead to unnecessary public alarm, while false aberration negatives could lead to health disasters. Farrington and Andrews [4] and Le Strat [5] provide more details on these issues.

4. Which disease and what to look for?

In retrospective surveillance, the disease is usually known and the features of interest are also known (e.g., trends). In prospective surveillance, particularly where bioterrorism may play a role, there could be little *a priori* knowledge of the disease and the

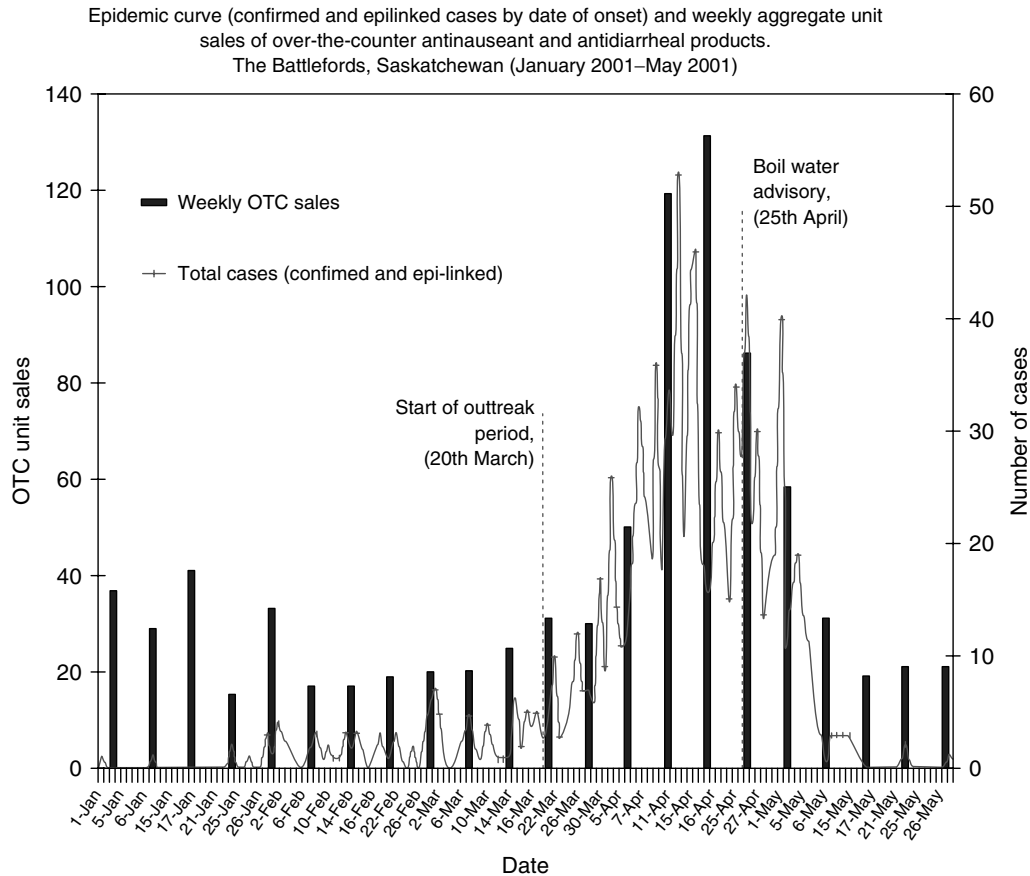


Figure 1 Battlefords Saskatchewan: Epidemic curve and time series of over-the-counter (OTC) pharmaceutical sales from January to May 2001

aberration to look for. Hence, detection systems must have the capability to deal with multiple diseases and possibly multiple forms of aberration. This is known as *multivariate-multifocus* surveillance. Because of the potentially huge database searching problem that results from this, *data mining* approaches have been adopted (see, e.g., Wong *et al.* [6] and Madigan [7]).

Aberration Types

1. Temporal

In temporal surveillance, a time series of events is available (usually either as counts of disease in intervals or as a point process of reporting or notification times). With time series of discrete counts,

the aberrations that might be of interest (suggesting “important” disease changes) are highlighted in Figure 2. In this figure, there are three main changes (A, B, C). The first aberration (A) is a sharp rise in risk (jump) or change point in level. The second aberration (B) that might be of interest is a cluster of risk (an increase then decrease in risk). This can only be detected retrospectively of course. The third aberration (C) is an overall process change where in addition to the change in the level of the process the variability is also increased or decreased.

When a point process of event times is observed, action may need to be taken whenever a new event arrives. Aberrations that are found in point processes can take the form of unusual aggregations of points (temporal clusters), sharp changes in rate of

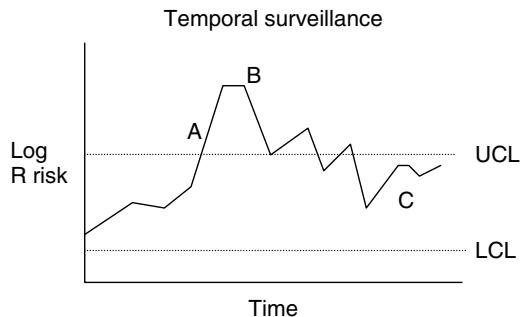


Figure 2 Schematic features of changes/aberrations found in temporal surveillance

occurrence (jumps), and overall process change (as above).

2. Spatio-temporal

If a spatial domain (disease map) is to be monitored, then spatial and spatio-temporal aberrations must be considered. Spatial aberrations could consist of discontinuities in risk between regions (jumps in risk across region boundaries), spatial clusters of risk (localized aggregations of cases of disease), and long-range trend in risk.

When disease maps are monitored in time, *changes* in the above features might be of interest. The development of new clusters of (say) infectious disease in time might signal localized outbreaks, for example.

Multivariate Issues

In health surveillance systems designed for prospective surveillance, many of the above features would have to be detectable across a wide range of diseases. Not only would whole disease streams be monitored, but subsets of disease streams may also need to be examined. For example, for early detection of effects it may be important to monitor frail subsets (such as the old or young age groups), and correlations between subsets may be important to making detection decisions. In addition, if disease maps are to be monitored over time as well as time series, then the problem grows considerably. The *Biosense* system developed by CDC (www.cdc.gov/phn/component-initiatives/biosense/index.html) has a multivariate and mixed-stream (spatial and temporal) capacity.

Models

One of the current concerns about large-scale surveillance systems is the ability to correctly model the variation in disease and to properly calibrate sensitivity and specificity with such a multivariate and multifocused task. The application of Bayesian hierarchical modeling has been advocated for surveillance purposes and this may prove to be useful in its ability to deal with large-scale systems evolving in time in a natural way. For example, recursive Bayesian learning could be important. Clearly computational issues could be paramount here given the potentially multivariate nature of the problem, and the need to optimize computation is paramount. Sequential **Monte Carlo** methods have been proposed to speed up computation [8–10] and faster algorithms are also available [11].

References

- [1] Thacker, S. (1994). Historical development, in *Principles and Practice of Public Health Surveillance*, S. Teusch & R. Churchill, eds, Oxford University Press.
- [2] Sosin, D. (2003). Draft framework for evaluating syndromic surveillance systems, *Journal of Urban Health* **80**, i8–i13.
- [3] Lawson, A.B. & Kleinman, K. (2005). *Spatial and Syndromic Surveillance for Public Health*, John Wiley & Sons, New York.
- [4] Farrington, P. & Andrews, N. (2004). Outbreak detection: application to infectious disease surveillance, in *Monitoring the Health of Populations: Statistical Principles & Methods for Public Health Surveillance*, R. Brookmeyer & D. Stroup, eds, Oxford University Press, New York.
- [5] Le Strat, Y. (2005). Overview of temporal surveillance, in *Spatial and Syndromic Surveillance for Public Health*, A. Lawson & K. Kleinman, eds, Oxford University Press, New York.
- [6] Wong, W., Moore, A., Cooper, G. & Wagner, M. (2002). *Rule-based Anomaly Pattern Detection for Detecting Disease Outbreaks*, MIT Press.
- [7] Madigan, D. (2005). Bayesian data mining for health surveillance, in *Spatial and Syndromic Surveillance for Public Health*, A. Lawson & K. Kleinman, eds, John Wiley & Sons, New York.
- [8] Kong, A., Lai, J. & Wong, W. (1994). Sequential imputations and Bayesian missing data problems, *Journal of the American Statistical Association* **89**, 278–278.
- [9] Doucet, A., De Freitas, N. & Gordon, N. (2001). *Sequential Monte Carlo Methods in Practice*, Springer, New York.

4 Disease Surveillance

- [10] Doucet, A., Godsill, S. & Andrieu, C. (2005). On sequential Monte Carlo sampling methods for Bayesian filtering, *Statistics and Computing* **10**, 197–208.
- [11] Neill, D. & Moore, A.W. (2005). Efficient scan statistic computations, in *Spatial and Syndromic Surveillance for*

Public Health, A.B. Lawson & K. Kleinman, eds, John Wiley & Sons, New York.

ANDREW B. LAWSON

Healthcare Policy and Management of Stroke, Breast Cancer, and Mental Health

Introduction

Usually the main stated goal of healthcare policy makers and managers is to use available resources to achieve the greatest possible improvement to the health and well-being of the people they serve. Health is often a cumulative result of a variety of influences or services, either provided or not provided, by a variety of organizations over time. A healthcare system is one of many players. In addition to a direct role using its own resources, the healthcare system also has a role in acting as an advocate for health and as a partner with other organizations concerned with health. Any serious attempt at monitoring the achievement of stated health goals should reflect this more complex reality. This article attempts to show the kinds of monitoring questions that policy makers and managers are often faced with and the role of statistics in such monitoring. Statistics on **health outcomes** that are available within the National Health Service (NHS) in England for stroke, breast cancer, and mental health related services, on the basis of work undertaken by the National Center for Health Outcomes Development (NCHOD), are used to illustrate the issues raised. Gaps in existing data and their implications for policy making and management are discussed briefly. The generic issues raised and illustrated here should be applicable to other healthcare systems and other types of health problems. In a short article it is only possible to provide an overview of concepts, methods and issues. Further details on these are available in the *Clinical and Health Outcomes Knowledge Base* published by NCHOD [1].

Policy and Management Questions Concerning Health Services

Figure 1 is a population health overview for stroke related services, developed by NCHOD [1], as an illustrative way of drawing a variety of linked health

states and services together to provide a summary overview of policy and management issues. It shows further categorization of “health improvement” to reflect more specific aspects of outcome, i.e., success in:

- reducing the risk of developing high blood pressure (HBP) and stroke;
- reducing the level of HBP to prevent stroke;
- reducing adverse consequences of blood pressure treatment such as drug side effects;
- reducing adverse consequences (other than death) of lack of timely and appropriate health care, such as preventable strokes and their complications, pressure sores;
- restoring function and improving quality of life; and
- reducing premature deaths.

Similarly, the services needed to achieve such success are divided into more precise categories, i.e.,

- proactive services to avoid the risk of developing HBP;
- timely detection of risk and removal/reduction of existing risk to health;
- timely services to detect and treat HBP; and
- late services to minimize consequences of stroke.

This presentation helps to clarify and highlight what we are trying to achieve, covering positive aspects such as improved quality of life as well as negative aspects such as avoidable disease and ill health. It reflects aspects such as clinical signs and disease as well as those of more immediate concern to patients such as handicap and impact on quality of life. It shows how lack of action may lead to a chain of potentially avoidable consequences, from risk, to disease, complications of disease, poor quality of life, and premature death and shows the kinds of action that may alter these consequences for the better. It also shows feedback loops, for example, a patient developing a stroke being at risk of developing a further stroke and needing all the services for reduction of risk and treatment of HBP. It should be noted that presentation of data in the context of such a “health outcomes overview” describes mismatching outcomes at a point in time. The data on risk show risk to future health. The data on levels of disease and death show consequences of missed opportunities to reduce risk (if applicable) in the past.

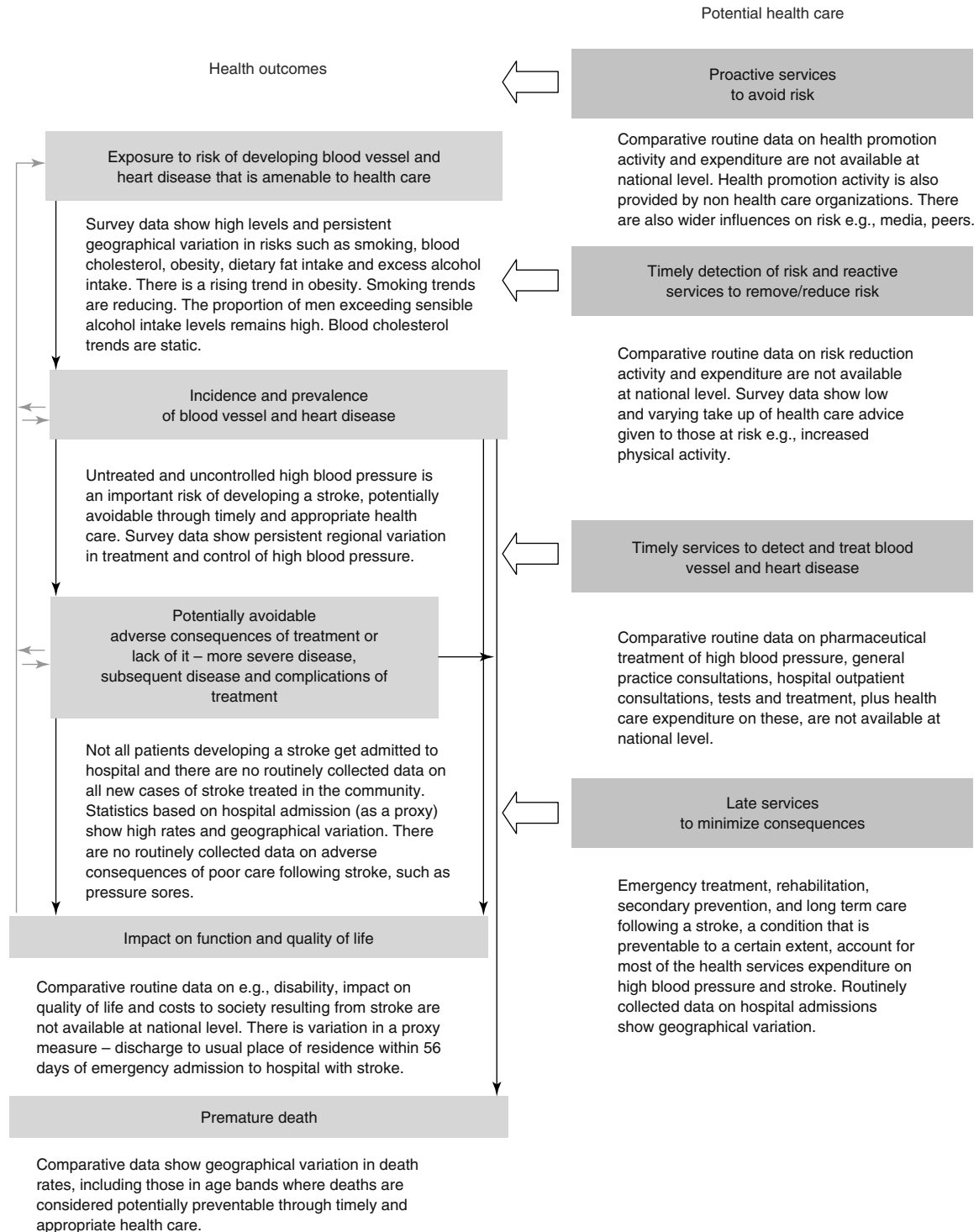


Figure 1 Overview of health care and outcomes: high blood pressure/stroke

It is also not possible to draw direct inferences on “attributable effects” of the services, as there may be different populations involved and a time lag between the services and the outcomes. However, such an approach does bring disparate but related items of information together in a meaningful way.

The same approach may be used to consider policy and management questions concerning breast cancer and mental health services. NCHOD’s Clinical and Health Outcomes Knowledge Base [1] contains health overviews with indicators for both conditions in its section on concepts and frameworks.

Statistics for Monitoring Achievements – Stroke

The greatest proportion of healthcare expenditure on HBP and stroke [2] is on stroke treatment, rehabilitation and long-term care, even though stroke is to a certain extent preventable through health service intervention. In assessing what is being achieved by healthcare, single indicators in isolation would present an incomplete picture of a complex reality.

Table 1 contains a set of indicators, drawn from NCHOD’s *Compendium of Clinical and Health Indicators* [3] that provide a more detailed insight as follows:

- A statistically significant rising trend in obesity. This implies an increasing future risk of HBP and stroke, and hence of future healthcare resource needs. Obesity is only partially amenable to health care.
- A statistically significant rising trend in levels of treated and controlled HBP (the single most important known preventive health service for stroke), although there is still plenty of room for improvement. In 2002, out of 10 653 adults of all ages with blood pressure measurements in the health survey for England (HSE), 25.1% had untreated HBP, and a further 8.4% had HBP that was treated but not controlled [4].
- A statistically significant falling trend in population levels of hospitalization for stroke. In the absence of suitable community based data, emergency hospital admissions may be used as a rough proxy for incidence, although not everyone with a stroke is admitted to hospital, and this will vary geographically. New guidelines on the development of dedicated stroke units were expected to

lead to an increase in hospital admission rates, so a falling trend may indicate falling incidence.

- A statistically significant rising trend in emergency readmissions after live discharge from hospital following treatment for stroke. A proportion of these readmissions may be due to potentially avoidable adverse events.
- A statistically significant falling trend in deaths after admission to hospital following a stroke.
- A statistically significant rising trend in timely discharge to usual place of residence. The timing is a proxy for completed treatment and rehabilitation, and a discharge to a place other than the usual place of residence may be indicative of an important life change and poor outcome.
- A statistically significant falling trend in population mortality from stroke. It is suggested that 70% of stroke mortality represents failure of secondary prevention and treatment by healthcare organizations, and 30% represents failure of primary prevention, some but not all of which is amenable to health care [5]. There is some overlap with the indicator on deaths after admission to the hospital.

These indicators provide a mixed picture. The NHS is improving on several fronts, but many of its resources are used for managing problems that are to some extent preventable. There is a concern about how stroke hospital activity (and its commensurate expenditure) should be used in the assessment of what is achieved. On the one hand it reflects services provided by a healthcare system when needed, but on the other, is a consequence of past failures of the healthcare system. This set of indicators also raises the issue of a time lag between resource use and its impact.

Variation in What is Achieved Through Stroke Services

Alongside reviewing national data and trends, policy makers and managers also need to review data at the level of geographic, administrative, or provider organizations. Variation in indicator values between such organizations may be either due to special local issues and circumstances or a reflection of inefficiency in the healthcare system. Table 1 shows variation between geographic administrative organizations in England (minimum and maximum values for each indicator

Table 1 Stroke indicator values over six NHS financial (FY)^(a) or calendar years^(b,c)

Outcome ^(c)		Year 1	Year 2	Year 3	Year 4	Year 5	Year 6 (latest) (95% confidence interval)	Average annual improvement (–deterioration) (trend ^(d))		
								%	R ²	Statistical significance
(1) Obesity	Number in sample obese England rate (%) Minimum SHA rate Maximum SHA rate	37 093 18.1 14.8 21.3	29 094 18.8 14.1 21.9	28 122 19.7 17.3 23.4	27 484 20.6 17.1 24.0	27 057 21.4 17.3 27.3	Data awaited	–4.4%	1.0	Significant
(2) Treated and controlled high blood pressure	Number in sample with treated and controlled HBP England rate (%) Minimum SHA rate Maximum SHA rate	322 6.1 4.3 9.0	No measurements	194 7.8 3.6 14.4	417 8.3 6.5 13.4	226 9.8 5.3 16.1	Data awaited	+12.1%	0.94	Significant
(3) Hospitalization for stroke – proxy for incidence	Number of admissions England rate/100 000 Minimum LA rate Maximum LA rate	69 457 144.7 41.6 240.2	67 738 140.3 31.3 242.2	64 930 132.9 36.6 235.5	64 826 131.1 50.1 240.4	66 006 132.1 28.0 210.2	65 108 129.3 46.0 208.5	+2.1%	0.84	Significant
(4) Stroke emergency readmissions	Number of readmissions England rate (%) Minimum PCO rate Maximum PCO rate	3284 7.7 0.0 30.3	3485 8.2 1.7 36.3	3401 8.2 1.1 24.7	3391 8.2 1.8 19.8	3517 8.6 2.4 20.5	3806 9.1 1.0 17.9	–2.7%	0.85	Significant

(continued overleaf)

Table 1 (continued)

Outcome ^(c)		Year 1	Year 2	Year 3	Year 4	Year 5	Year 6 (latest) (95% confidence interval)	Average annual improvement (-deterioration) (trend ^(d))		
								%	R ²	Statistical significance
(5) Stroke hospital case-fatality	Number of deaths	21 300	19 523	18 163	17 658	17 773	17 150			
	England rate/100 000	30 497.1	28 711.6	27 500.8	26 842.4	26 684.3	25 974.3			
	Minimum	16 032.9	16 823.0	11 134.6	13 425.9	15 303.7	12 846.6			
	PCO rate	45 406.8	46 243.6	43 839.6	40 764.8	44 054.4	42 241.5			
(6) Timely discharge to usual place of residence following stroke admission	Maximum PCO rate						(29 581.1–58 482.1)			
	Number of discharges	30 145	30 192	29 589	29 558	30 126	30 548			
	England rate (%)	47.9	48.9	49.8	50.1	50.8	52.0			
	Minimum PCO rate	26.2	34.3	30.2	33.5	34.7	28.0			
(7) Stroke population mortality	Maximum PCO rate	63.2	66.4	63.8	61.2	67.3	66.4			
	Number of deaths	59 367	57 850	54 297	54 861	55 351	54 022			
	England rate/100 000	73.9	71.3	65.9	65.3	65.0	63.1			
	Minimum LA rate	47.3	47.1	42.8	34.0	35.6	37.5			
	Maximum LA rate	124.5	109.1	105.2	98.2	102.3	92.2			
							(70.1–114.2)			
								+3.0%	0.9	Significant
								+1.5%	0.97	Significant
								+2.9%	0.9	Significant

- (a) Financial year; FY.
- (b) Organizations: SHA = NHS Strategic Health Authorities, PCO = NHS Primary Care Organizations, LA = Local Authorities (local government) – boundaries as at April 2003.
- (c) Sources of data: HES = Hospital Episode Statistics, ONS = Deaths data from the Office for National Statistics, HSE = Health Survey for England, Continuous inpatient spells (CIPS) combine all contiguous episodes across all NHS hospitals for treatment, rehabilitation and care within the year of analysis for individual patients.
- (1) Directly age and sex standardized percent of adults (16+) who are obese (body mass index over 30) for 1998–2002, standardized using the England 2001 census noninstitutional adult population count (*Source*: HSE).
 - (2) Those on prescribed blood pressure drugs with a blood pressure reading below 140/90 mm Hg during the survey, as a directly age standardized percent of all with high blood pressure, for calendar years 1998–2002, standardized using the England 2001 census noninstitutional adult population count (*Source*: HSE).
 - (3) Indirectly age and sex standardized rates of emergency admissions of patients to hospital with a stroke (primary diagnosis international classification of diseases (ICD) version 10 (ICD-10) codes I61–I64) per 100 000 residents of all ages, for FY 1998/1999–2003/2004, standardized to 2001/2002. The denominators are mid-year population estimates for 1998–2002 (*Source*: HES, ONS).
 - (4) Indirectly age and sex standardized percent of emergency readmissions of patients of all ages to any hospital between 0 and 27 days (inclusive) of live discharge following admission with a stroke (primary spell diagnosis ICD-10 I61–I64), for FY 1998/1999–2003/2004, standardized to 2001/2002 (*Source*: HES).
 - (5) Indirectly age and sex standardized rates of deaths occurring in hospital and after discharge among patients of all ages, between 0 and 29 days (inclusive) of an emergency admission with a stroke (primary spell diagnosis ICD-10 I61–I64) per 100 000 CIPS, for FY 1998/1999–2003/2004, standardized to 2001/2002 (*Source*: HES, ONS).
 - (6) Indirectly age and sex standardized percent of continuous inpatient spells (CIPS) for patients of all ages with an emergency admission with a primary diagnosis of stroke (ICD-10 codes I61–I64), where the patient was discharged to the preadmission category of accommodation between 0 and 55 days (inclusive) of admission, for FY 1998/1999–2003/2004, standardized to 2001/2002 (*Source*: HES).
 - (7) Directly age standardized mortality rate from stroke (ICD-10 codes I60–I69, ICD-9 codes 430–438) per 100 000 residents of all ages for 1998–2003, standardized to the European standard population. These ICD codes are different from those used for stroke in the HES based indicators. HES has ICD-10 coding across all years while ONS deaths data do not, and there is no equivalent for ICD-10 I61–I64 in ICD-9. Years 1998 and 2000 use ICD-9, and years 1999, 2001 and 2002 use ICD-10. Data based on ICD-9 are adjusted for equivalence to ICD-10 (*Source*: ONS).

See the *Compendium of Clinical and Health Indicators*, www.ncho.d.nhs.uk for further details on indicators and methods.

(d) Fitting a linear trend to log transformed rates.

and for each year) at the level of NHS Strategic Health Authority (SHA) (populations 1–2.5 million), NHS Primary Care Organization (PCO) (populations 65 000–375 000) or Local Authorities (LA, local government, populations 25 000 to 990 000). The data for the latest year are accompanied by **confidence intervals**, showing that some of the minimum and maximum values are statistically significantly different from the England value. For example, the maximum PCO rate of deaths after admission to hospital with stroke, 42 241.5/100 000, is statistically significantly higher than the England average of 25 974.3/100 000. Sometimes, confidence intervals at organizational level may be wide due to errors in **estimation** when small numbers of cases are involved and differences in rates may not be statistically significant (see methods, annex 3, in the *Knowledge Base*) [6]. The number of stroke deaths at PCO level is sufficiently large for this not to be an issue.

Figure 2 shows geographical variation between PCOs in England, grouped in Office for National Statistics (ONS) Area Groups. Each dot represents a PCO. The line around each dot represents 95% confidence intervals around the death rate. The standardised death rate is an estimate and the true value is likely to be within the interval with 95% probability. In visual terms, if the confidence intervals of a PCO value do not overlap those for England, the PCO value is considered statistically significantly different from the England one. When comparing geographical areas, it is important to compare “like” with “like”, hence the grouping of PCOs by ONS Area Groups. The ONS has used cluster analyses to create an area classification for grouping PCOs that are most similar in terms of 42 demographic, socioeconomic, housing, and other census 2001 variables [7]. Figure 2 shows that there is substantial variation between “like” PCOs, suggesting that factors other than population or environment may influence these rates, such as quality of health care. There may also be differences in the quality of data submitted by different organizations that need to be taken into account when interpreting organizational variations [8].

Table 2 shows statistically significant differences in indicator values at England level by gender and social deprivation. The latter is based on the index of multiple deprivation [7]. These are aspects that managers and policy makers need to take into account.

Some may reflect varying risk, others poorer quality of service delivery. Where variation reflects different levels of risk that is not amenable to health care, (e.g., older people being at higher risk of developing a stroke), indicators that involve comparison between organizations, or over time need to be adjusted to reflect such differences in risk between the populations being assessed. The indicators in Table 1 have, in general, been standardized by age and sex.

Statistics for Monitoring Achievements – Breast Cancer

There have been several recent commitments to improving cancer services and outcomes (such as the NHS cancer plan). The overall target for reducing the disease burden is expressed in terms of mortality, but mortality trends reflect both the incidence of cancer and the survival of cancer patients. Mortality is not the best indicator to measure the progress of health services in treating those with the disease.

Table 3 contains a set of indicators, drawn from NCHOD’s *Compendium of Clinical and Health Indicators* [3] that provide a more detailed insight as follows:

- Large numbers of new patients diagnosed with breast cancer each year, with no statistically significant change in rates over the years. There is substantial variation in incidence rates between Local Authorities.
- On average, around three quarters of all eligible women in England who are offered screening for breast cancer take it up. In a recent year, this take-up varied from 27% to 86% between PCOs, as shown in Figure 3.
- A statistically significant falling trend (improvement) in deaths from breast cancer.
- Around a fifth of women dying from breast cancer die at home. This varies from 4% to 46% between PCOs and may reflect patient and family choice, availability of hospices and community support and so on. There are quality of dying implications.
- A statistically significant rising trend (improvement) in five-year survival after diagnosis of breast cancer.

Relative survival for a particular period is an estimate of the proportion of cancer patients who

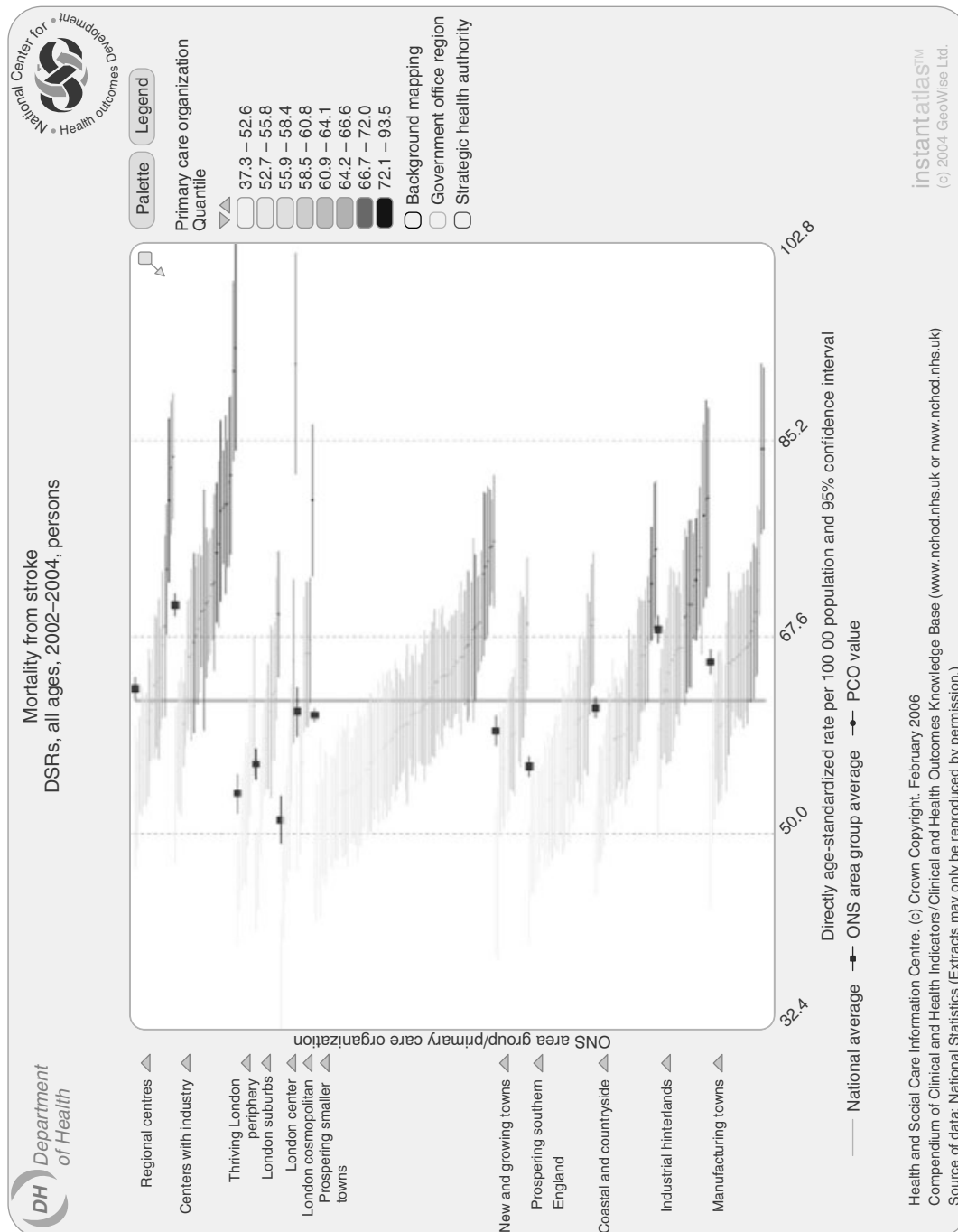


Figure 2 Geographical variation in deaths from stroke

Table 2 Stroke variation by gender and deprivation^(a)

Return to preadmission category of accommodation within 56 days (inclusive) of emergency admission to hospital with a stroke, England, financial year 2003/2004, all ages					
	Number of completed inpatient spells	Number of discharges	Indirectly standardized percent	Lower limit of 95% confidence interval	Upper limit of 95% confidence interval
Persons (age and sex standardized to persons 2001/2002)	58 922	30 548	52.0	51.4	52.6
Males (age standardized to persons 2001/2002)	28 034	16 332	54.1	53.3	55.0
Females (age standardized to persons 2001/2002)	30 888	14 216	49.8	49.0	50.6
Emergency readmissions of patients of all ages to any hospital within 28 days (inclusive) of live discharge following admission with a stroke, England, financial year 2003/2004					
	Number of discharges	Number of readmissions	Indirectly age and sex standardized (to 2001/2002) %	Lower limit of 95% confidence interval	Upper limit of 95% confidence interval
All groups (including those not classified)	41 623	3806	9.1	8.8	9.4
Deprivation level 1 (most deprived)	9491	944	10.1	9.5	10.8
Deprivation level 2	8474	790	9.3	8.7	10.0
Deprivation level 3	8470	775	9.1	8.4	9.7
Deprivation level 4	7732	666	8.5	7.9	9.2
Deprivation level 5	7016	603	8.5	7.8	9.2

^(a)Sources of data: hospital episodes statistics. The Information Centre for health and social care

survive their disease after adjustment for death from other causes (that is, the ratio of the survival rate observed to that which would have been expected if cancer patients had the same overall mortality as the general population). Relative survival can be directly standardised for age to take into account variation in survival by age and changes in the age distribution of cancer patients over time or between areas. At the time of their publication, cancer survival indicators will reflect the impact of past as well as current health care, over the five years from diagnosis. Trends and geographical variation may be due to several factors including the quality of primary care, the speed of referral, and the quality of treatment

services. The survival rates may also be a proxy for other factors not readily measured, such as the local population's level of understanding of cancer symptoms and what to do about them, variations in the extent of disease at diagnosis (stage) and in the histology and grade of tumours, and artefacts in the data.

Not enough is known about the causes of breast cancer and the potential to prevent it. However, early detection and timely, high quality treatment can make a difference to survival. The substantial variation between geographical areas and comparisons with other countries show that there remains scope for improvement.

Table 3 Breast cancer indicator values over six NHS financial (FY)^(a) or calendar years^(b,c)

Outcome ^(c)		Year 1	Year 2	Year 3	Year 4	Year 5	Year 6 (latest) (95% confidence interval)	Average annual improvement (–deterioration) (trend ^(d))		
								%	R ²	Statistical significance
(1) Incidence of breast cancer	Number of patients with breast cancer registered England	35 190	34 466	35 043	34 600			+0.2%	0.04	Not significant
	rate/100 000	119.2	115.8	116.6	114.9					
	Minimum LA rate	46.9	38.5	0.0	33.3					
	Maximum LA rate	230.9	308.3	181.4	186.9					
(2) Breast cancer screening programme coverage	Number of women screened				2 600 704	2 649 806	2 651 551			
	England rate (%)	Data not available	Data not available	Data not available	76.13	75.28	74.95	–0.8%	0.94	Not significant
	Minimum PCO rate				35.95	26.12	27.44			
	Maximum PCO rate				90.58	92.22	85.56			
(3) Deaths from breast cancer	Number of deaths England	3793	3823	3726	3638	3554	3543			
	rate/100 000	70.1	69.9	67.5	65.1	62.9	61.2	+2.9%	0.98	Significant
	Minimum LA rate	0.0	0.0	0.0	0.0	0.0	0.0			
	Maximum LA rate	149.8	433.8	203.4	203.6	355.1	334.2			

(continued overleaf)

Table 3 (continued)

Outcome ^(c)		Year 1	Year 2	Year 3	Year 4	Year 5	Year 6 (latest) (95% confidence interval)	Average annual improvement (-deterioration) (trend ^(d))		
								%	R ²	Statistical significance
(4) Deaths at home from breast cancer	Number of deaths	Data not available	Data not available	Data not available			6456			
	England rate (%)						20.5 (20.0–20.9)			
	Minimum PCO rate						4.2 (1.1–15.3)			
	Maximum PCO rate						45.5 (33.4–57.8)			
(5) Five-year survival after diagnosis of breast cancer	England rate (%)	77.1	78.5	79.9	81.6		81.6	+1.5%	0.95	Significant
	Minimum SHA rate	68.2	69.6	73.3.9	78.4		76.5			
	Maximum SHA rate	83.6	82.2	85.1	85.6		84.7			
	SHA rate									

(a) Financial year; FY.

(b) Organisations: SHA = NHS Strategic Health Authorities, PCO = NHS Primary Care Organisations, LA = Local Authorities (local government) – boundaries as at April 2003.

(c) Sources of data: ONS = Office for National Statistics,

(d) fitting a linear trend to log transformed rates.

(1) Directly age standardised incidence rate of breast cancer (ICD-9 code 174 adjusted for equivalence to ICD-10, ICD-10 code C50) per 100 000 resident women of all ages for 1999–2002, standardised to the European Standard Population (*Source*: ONS).(2) Percent of women, aged 53 to 64 years and eligible for breast cancer screening, who have had a test with a recorded result at least once during the previous 3 years, for 2002–2004 (*Source*: Department of Health breast screening programme).(3) Directly age standardised mortality rate from breast cancer (ICD-9 code 174 adjusted for equivalence to ICD-10, ICD-10 code C50) per 100 000 resident women of ages 50–69 years, for 1999–2004, standardised to the European Standard Population (*Source*: ONS).(4) Indirectly age standardized percent of deaths from breast cancer (ICD-10 code C50) that occur at home, among women of all ages for 2002–2004 pooled (*Source*: ONS).(5) Relative survival rate (%) of women diagnosed 5 years previously with breast cancer (age at diagnosis 15–99 years), for 2000–2004. The relative survival rate is the ratio of the survival rate actually observed among the cancer patients and the survival that would have been expected if they had only had the same overall mortality rates as the general population (*Source*: ONS).See the *Compendium of Clinical and Health Indicators*, www.nchod.nhs.uk for further details on indicators and methods.

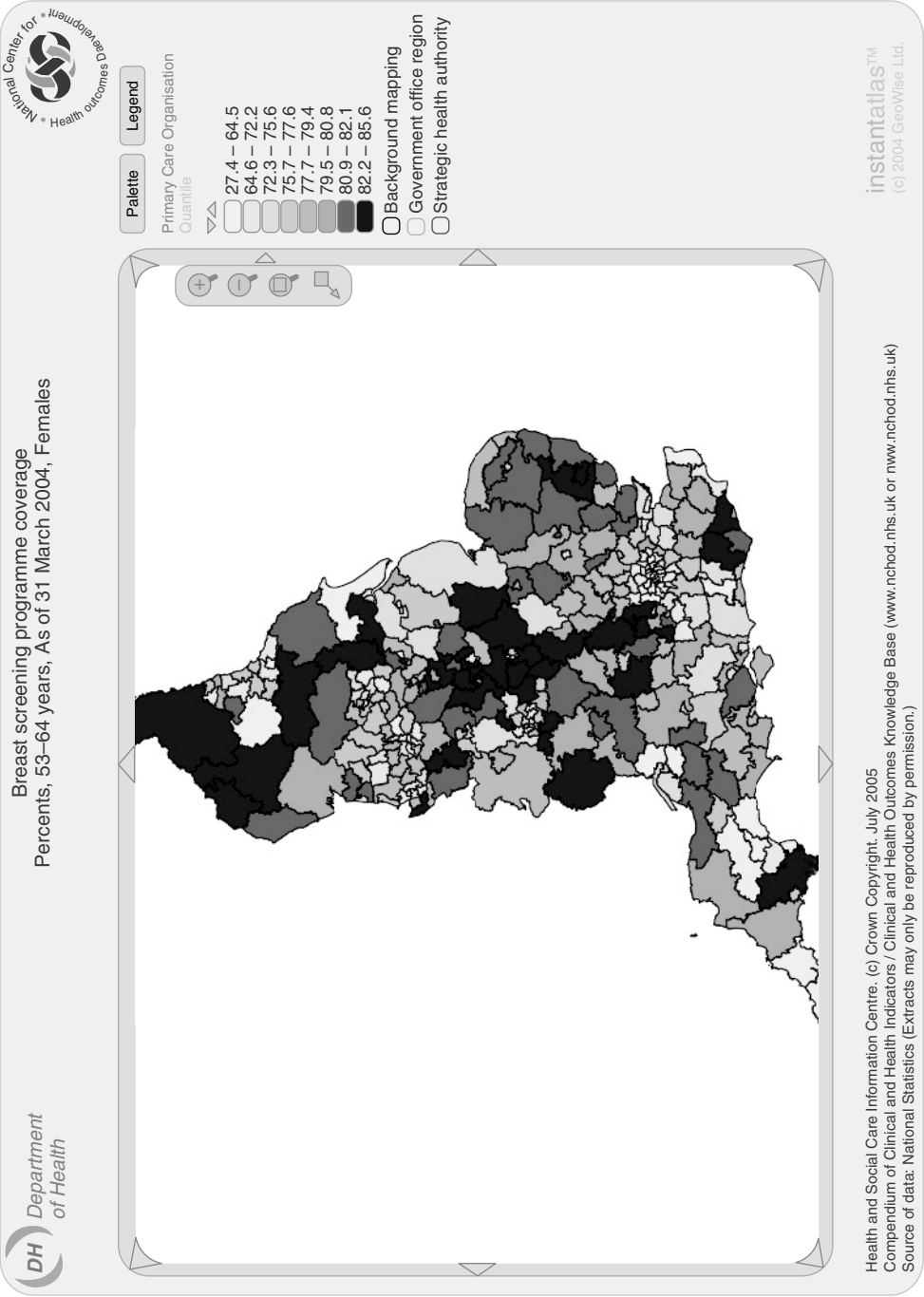


Figure 3 Geographical variation in the coverage of a national breast cancer screening programme

Table 4 Mental health indicator values over six NHS financial (FY)^(a) or calendar years^(b,c)

Outcome ^(c)		Year 1	Year 2	Year 3	Year 4	Year 5	Year 6 (latest) (95% confidence interval)	Average annual improvement (–deterioration) (trend ^(d))		
								%	R ²	Statistical significance
(1) Readmission to hospital following psychiatric discharge – adults	Number of readmissions	45 264	44 011	44 454	44 687	42 145	41 416			
	England rate (%)	2.11	2.10	2.09	2.13	2.14	2.22	–1.04%	0.78	Significant
	Minimum PCO rate	1.00	1.00	1.00	1.00	1.29	1.34			
	Maximum PCO rate	6.04	5.46	5.47	4.11	4.33	4.68			
(2) Readmission to hospital following psychiatric discharge – older people	Number of readmissions	28 501	25 061	24 048	22 735	20 916	19 989			
	England rate (%)	2.75	2.63	2.62	2.68	2.66	2.68	+0.29%	0.14	Not significant
	Minimum PCO rate	1.00	1.00	1.00	1.00	1.00	1.00			
	Maximum PCO rate	8.84	8.75	10.31	10.75	8.53	10.06			
(3) Total stay in hospital during one year following a psychiatric admission – adults	Total days in hospital	1 433 096	1 397 809	1 548 007	1 486 155	1 378 242	1 296 549			
	England rate /100 000	66.80	66.86	72.85	70.87	70.02	69.43	–1.16%	0.44	Not significant
	Minimum PCO rate	0.50	1.00	24.00	3.00	26.44	32.63			
	Maximum PCO rate	219.89	300.50	365.00	140.00	127.44	163.67			
(4) Total stay in hospital during one year following a psychiatric admission – older people	Total days in hospital	855 624	803 175	851 680	796 332	721 393	686 406			
	England rate (%)	82.41	84.15	92.84	93.77	91.72	92.10	–2.54	0.75	Significant
	Minimum PCO rate	26.18	11.00	9.00	0.00	30.50	6.00			
	Maximum PCO rate	232.50	365.00	365.00	730.00	273.00	160.83			

(continued overleaf)

Table 4 (continued)

Outcome ^(c)		Year 1	Year 2	Year 3	Year 4	Year 5	Year 6 (latest) (95% confidence interval)	Average annual improvement (–deterioration) (trend ^(d))		
								%	R ²	Statistical significance
(5) Combined score for total admissions and stay in hospital over one year following a psychiatric admission – adults	England composite z-score	0.57	0.57	0.76	0.77	0.76	0.89	–12.68	0.85	Significant
	Minimum PCO score	–3.23	–4.11	–3.27	–4.04	–1.87	–1.95			
	Maximum PCO score	8.90	7.82	9.14	5.56	7.42	4.84			
(6) Combined score for total admissions and stay in hospital over one year following a psychiatric admission – older people	England composite z-score	0.90	0.85	1.07	1.13	1.07	1.10	–5.11%	0.73	Significant
	Minimum PCO score	–1.72	–1.95	–2.33	–2.56	–1.52	–2.41			
	Maximum PCO score	6.27	6.61	7.14	16.59	6.72	6.88			
(7) Deaths from suicide	Number of deaths England rate /100 000	3398	3316	3123	3082	3004	3169	+2.6%	0.80	Significant
	Minimum LA rate	0.00	0.00	0.00	0.00	0.00	0.00			
	Maximum LA rate	15.30	20.79	18.21	17.79	16.06	18.16			
	Number of deaths	5006	4743	4560	4450	4452	4518			
(8) Deaths from suicide and injury undetermined	England rate /100 000	9.80	9.26	8.84	8.60	8.51	8.57	+2.8%	0.90	Significant
	Minimum LA rate	0.00	0.00	0.00	0.00	0.00	0.00			
	Maximum LA rate	22.84	24.54	22.55	19.42	23.82	30.30			

- (a) Financial year; FY.
- (b) Organisations: SHA = NHS Strategic Health Authorities, PCO = NHS Primary Care Organisations, LA = Local Authorities (local government) – boundaries as at April 2003.
- (c) Sources of data: HES = Hospital Episode Statistics (The Information Centre for health and social care), ONS = Deaths data from the Office for National Statistics, Continuous inpatient spells (CIPS) combine all contiguous episodes across all NHS hospitals for treatment, rehabilitation and care within the year of analysis for individual patients
- (d) Fitting a linear trend to log transformed rates.
- (1) Average number of admissions to any hospital in England over 12 months, for patients aged 17–64 years first admitted for psychiatric care during April to June of each of FY1999/2000 to 2004/2005 (*Source*: HES).
 - (2) Average number of admissions to any hospital in England over 12 months, for patients aged 65+ years first admitted for psychiatric care during April to June of each of FY1999/2000 to 2004/2005 (*Source*: HES).
 - (3) Total stay in hospital across all admissions to any hospital in England over 12 months, for patients aged 17–64 years first admitted for psychiatric care during April to June of each of FY1999/2000 to 2004/2005 (*Source*: HES).
 - (4) Total stay in hospital across all admissions to any hospital in England over 12 months, for patients aged 65+ years first admitted for psychiatric care during April to June of each of FY1999/2000 to 2004/2005 (*Source*: HES).
 - (5) Composite modified z score for average number of admissions and total stay in hospital over a 12 month period, for patients aged 17–64 years first admitted for psychiatric care during April to June of each of FY1999/2000 to 2004/2005, compared with the best performance among SHAs in FY 2001/02 (*Source*: HES).
 - (6) Composite modified z score for average number of admissions and total stay in hospital over a 12 month period, for patients aged 65+ years first admitted for psychiatric care during April to June of each of FY1999/2000 to 2004/2005, compared with the best performance among SHAs in FY 2001/02 (*Source*: HES).
 - (7) Directly age standardised mortality rate from suicide (ICD-9 codes E950–E959, ICD-10 codes X60–X84) per 100 000 residents of all ages for 1999–2004, standardised to the European Standard Population (*Source*: ONS).
 - (8) Directly age standardised mortality rate from suicide and injury undetermined (ICD-9 codes E950–E959, E980–E989 exc. E988.8; ICD-10 codes X60–X84, Y10–Y34 exc. Y33.9) per 100 000 residents of all ages for 1999–2004, standardised to the European Standard Population (*Source*: ONS).

See the *Compendium of Clinical and Health Indicators*, www.ncho.d.nhs.uk for further details on indicators and methods.

Statistics for Monitoring Achievements – Mental Health

Monitoring the quality of mental health care is problematic. Deaths are low, so indicators based on the number of deaths are inadequate. Table 4 shows statistically significant falling death rates from suicides and undetermined injury but this is an inadequate reflection of all of mental health care. Much of mental health care occurs outside the hospital and there is little by way of routine **data collection** on activity or outcomes. Also, there are few explicit standards. The challenge is to find a way of using routine data from hospitals to make some inferences about both hospital and community care and to identify aspects of care that are materially below optimum. National service guidelines for mental health highlight the preference for community over hospital care. With this policy being expected to produce better outcomes, assessing the way the service for mental health is delivered may be used as a proxy for quality of care.

Mentally ill patients vary in the number of readmissions to hospitals and time spent in hospital, and these may reflect variations in the quality and availability of care and support in the community. Too many readmissions and long cumulative lengths of stay may reflect inadequate community care, but too few readmissions and short cumulative lengths of stay may reflect inadequate provision of necessary hospital care. A combination of the total number of admissions and the total time spent in hospital by patients during a year could be used to assess this. There may well be a trade-off between time spent in hospital and frequency of admission, and it is therefore important to consider the balance between the two variables as well as the individual measures. Studies of observed variation between populations could be used to derive target ranges for optimum patterns of care.

Table 4 shows sets of indicators for adults and older people that have been constructed by linking hospital data for individual patients over 1 year. The starting point for each patient is a first admission to hospital for psychiatric care between April and June during the year of assessment, with follow up for 12 months after the date of the first admission. The averages for all patients within a geographical area for the total number of days spent by patients in hospital and the total number of admissions per person during 12 months have been calculated.

Readmissions could be to any NHS hospital in England and for any condition, such as injury, not just mental illness.

Figure 4 shows that there is substantial variation between primary care organisations for each of the two variables. To give each component equal weight, they were transformed into a modified version of a statistic called a “z score” (by measuring the distance from a defined reference point and dividing by the standard deviations of the individual distributions). Conventionally, the mean of a distribution is the reference point for z scores. However, the mean is not necessarily a suitable target for a performance score because it is partly a reflection of “poor performance” at the tail end of the distribution. Reference points that are assumed to reflect an “optimum” may be more appropriate. Modified z scores, in which the reference points were 1.96/1.94 admissions per person and 59.76/71.99 days total stay in hospital during 12 months (for adults/older people respectively), were calculated. These reference points were selected as a realistic joint target because they had been actually achieved by a Strategic Health Authority in 2001/02 (the financial year selected as standard). This “best achieved combination” of a low rate of admissions per person and a low total stay per person, giving each the same level of importance, equates to the lowest composite z score at the level of strategic health authority, although at the level of primary care organisation, some achieved even better z scores. These reference points were used to calculate modified z scores across all 6 years.

The z scores measure how far each primary care organization is from the defined optimum for each of total admissions and total stay. A composite score was then produced for each organisation by adding its two z scores, equating to a summary of that organisation’s mix of experience on two fronts. A high composite score would be considered undesirable as it would indicate more hospitalisation than expected. Two organisations with a different mix of values for admissions and length of stay may have the same composite z score, reflecting the same performance but different trade-offs. Improvement over time should lead to lower overall composite z scores and a reduction in variation. Table 4 shows statistically significant worsening trends in composite z scores at the England level for both, adults and older people over six financial years. There is also substantial variation between primary care organizations.

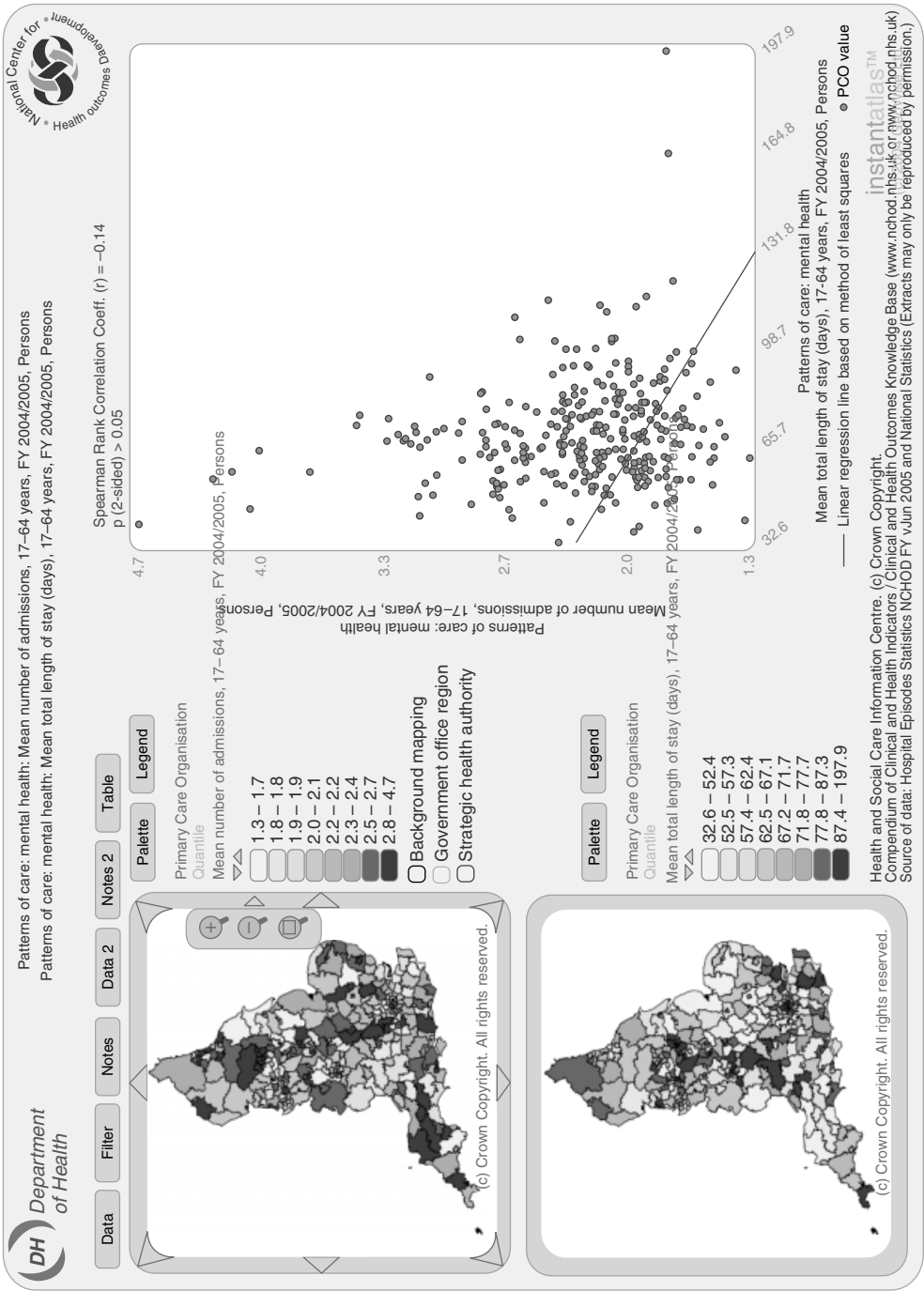


Figure 4 Geographical variation in the relationship between frequency of hospital admission and total time spent in hospital over one year by mental health patients

A possible cause of variation between primary care organisations might be differences in the prevalence of illness and in the level of support provided in the community, both by professional bodies and informal networks such as families. A full assessment of the effect of such factors is difficult, however, a comparison of similar geographical areas may help. For example, the populations of big cities can be more transient, have a greater incidence of some mental illnesses, and have fewer informal networks. The Office for National Statistics has used cluster analyses to create an area classification for grouping primary care organisations that are most similar in terms of 42 demographic, socioeconomic, housing, and other census 2001 variables. Figure 5 shows the composite modified z score for each primary care organisation grouped within its area group. Although it shows some differences between groups, there is much greater residual variation within each group, suggesting that influences other than demography and socioeconomic conditions are largely responsible for the variation, such as service availability and clinical practice.

The overall conclusion is that while there is a policy (with implied benefit) for care for mental health patients to take place in the community, there is a growing trend in hospitalisation and substantial variation between geographical areas. These may reflect both the availability and quality of health and social care for such patients and need further investigation. This example also shows how existing routine data may be used creatively for policy and management purposes.

Gaps in Data

Figure 1 shows a number of gaps in existing stroke related data. For example, there are no routinely collected data, at national level, on new patients with stroke that are treated in the community, making measurement of the incidence of stroke and hence what is achieved through preventive services, incomplete. Similarly, there are no data on levels of disability, impact on quality of life, and consequences of poor quality care after a stroke, for example, pressure sores, aspects which are highly important to patients. A multi-disciplinary working group, involving clinicians, policy makers, managers, researchers, representatives of patients, and others, established to

advise on data that should be collected, made recommendations for 29 indicators that could be produced to monitor all aspects of clinical and health outcomes [9]. Many of these indicators will need new data collection.

Two further working groups, on breast cancer and severe mental illness have also recommended indicators along the same lines [10,11]. While these reports are now eight years old and need updating, they illustrate how achievements of health care may be monitored more comprehensively. Recommendations for indicators related to breast cancer, for example, included:

- Indicators related to improvement of early detection and treatment of breast cancer:
 - incidence of primary breast cancer by cancer stage at diagnosis
 - detection rate of small primary breast cancers
 - incidence of interval breast cancers
 - incidence of invasive breast cancer following treatment for ductal carcinoma in situ.
- Indicators related to reduction of death and complications from breast cancer and its treatment
 - survival rates by cancer stage at diagnosis
 - population-based death rates from breast cancer in women
 - percentage of patients, diagnosed as having breast cancer, who had received a diagnostic triple assessment
 - percentage of patients with breast cancer under the care of a breast cancer specialist
 - percentage of patients receiving treatment for breast cancer whose care was planned jointly by a multi-disciplinary group
 - percentage of patients who, having undergone potentially curative surgery for breast cancer, were given chemotherapy as part of their primary treatment
 - recurrence rates of breast cancer by site and type of primary surgery
 - Rates of specific complications detected, within one year of discharge from hospital, among patients having undergone in-patient treatment for breast cancer
- Indicators related to maintenance of well-being during and following treatment for breast cancer.

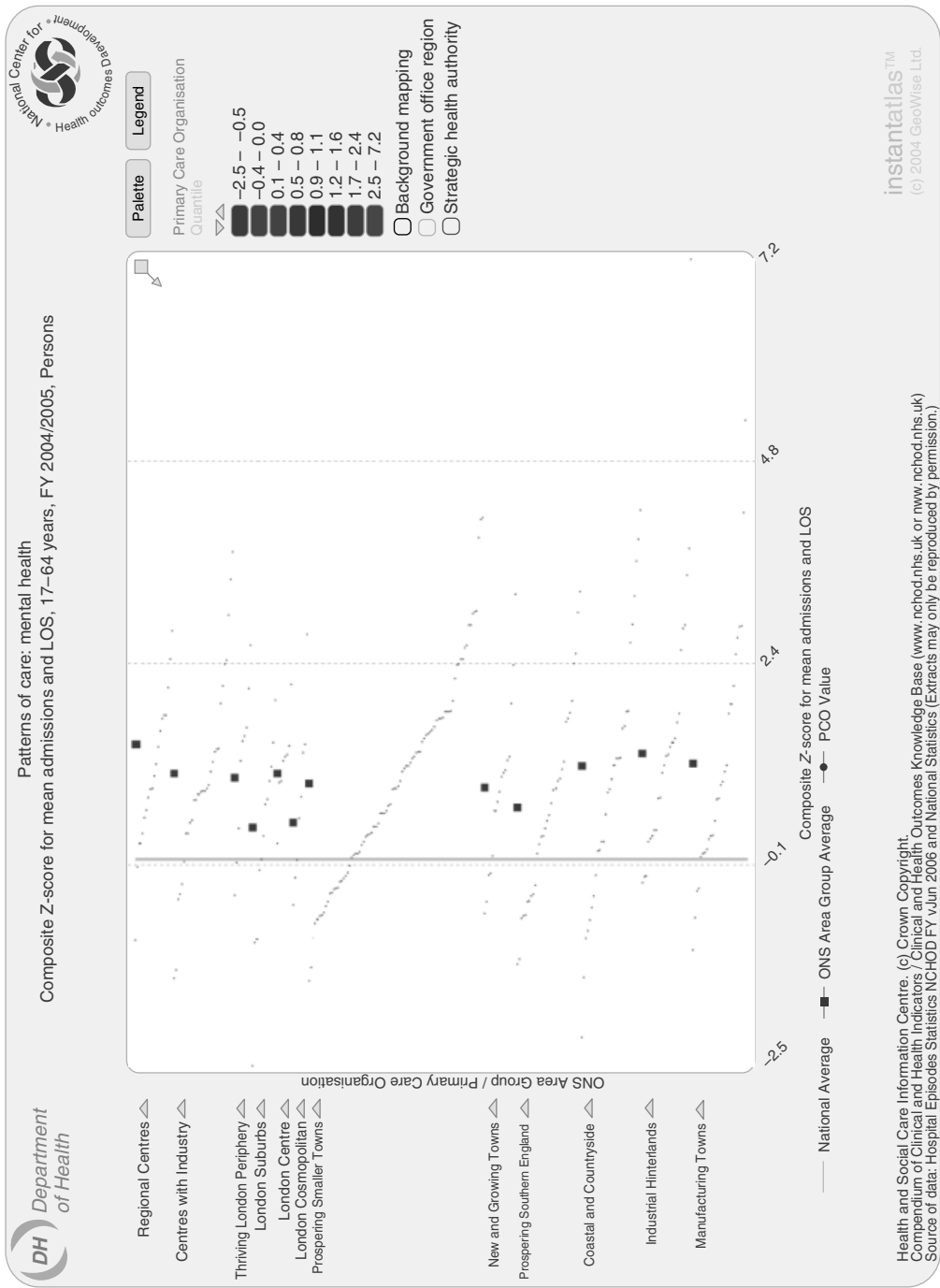


Figure 5 Geographical variation, within and between clusters of populations with similar socio-demographic characteristics, in a hospitalisation score for mental health patients

- delays from patient identification of symptoms to start of treatment for breast cancer
- patient satisfaction with specific areas of the management of their breast cancer care
- assessment of health-related quality of life within a population of patients one year after diagnosis of breast cancer
- assessment of psychological distress in the patient following treatment for breast cancer.

Where routine data are not available at the national level, policy makers and service providers may need information on existing independent datasets and guidance on how to undertake their own measurements and data collection to fill in gaps. NCHOD provides a list and selective review of the quality of existing databases in the UK (*Directory of Clinical Databases (DoCDat)*), some of which are voluntary with only part coverage across the country [12]. Such databases, for example, the Royal College of Physicians national stroke audit, can be alternative sources of data for such monitoring. NCHOD also provides a bibliography of the research literature on the measurement of outcomes from a patient's perspective, plus guidance on how to select and use relevant measurement instruments [13].

Conclusions

This brief article has shown the kinds of outcome related questions that policy makers and managers ought to be asking about health care. The focus has been on what is achieved through resource use. The role of statistics in monitoring this, and constraints in the availability and use of statistics have been illustrated using data on stroke, breast cancer, and mental health from the NHS in England. Information on outcomes is one part of wider monitoring, covering the organization of services; resources available; access to and quality of service delivery; and patient perspectives, including their satisfaction with the services provided. This article is of necessity, an oversimplification of what are in reality, extremely complex systems and issues.

References

- [1] Lakhani, A., Olearnik, H. & Eayres, D. (eds) (2006). *Clinical and Health Outcomes Knowledge Base*, The Information Centre for Health and Social Care and the National Centre for Health Outcomes Development, London (<http://www.nchod.nhs.uk>, accessed May 2007).
- [2] Central Health Monitoring Unit. (1996). *Burdens of Disease*, NHS Executive, Wetherby (Cat. No 96CC0036).
- [3] Lakhani, A., Olearnik, H. & Eayres, D. (eds) (2006). *Compendium of Clinical and Health Indicators*, The Information Centre for Health and Social Care and the National Centre for Health Outcomes Development, London (<http://www.nchod.nhs.uk>, accessed May 2007).
- [4] The Information Centre. (2006). *Health Survey for England – Trend Data* (www.ic.nhs.uk/pubs/hse05trends, accessed August 2007).
- [5] Nolte, E. & McKee, M. (2004). *Does Healthcare Save Lives? – Avoidable Mortality Revisited*, Nuffield Trust, London.
- [6] Eayres, D. (2006). *Explanation of Statistical Methods: Methods Annex 3 – Compendium of Clinical and Health Indicators*, The Information Centre for Health and Social Care and the National Centre for Health Outcomes Development, London (<http://www.nchod.nhs.uk>, accessed May 2007).
- [7] Greeno, D. & Eayres, D. (2006). *Measuring Deprivation: Methods Annex 1 – Compendium of Clinical and Health Indicators*, The Information Centre for Health and Social Care and the National Centre for Health Outcomes Development, London (<http://www.nchod.nhs.uk>, accessed May 2007).
- [8] Spence, C., Lakhani, A. & Coles, J. (2006). *Hospital Episodes Statistics – Assessment of Data Quality: Methods Annex 4 – Compendium of Clinical and Health Indicators*, The Information Centre for Health and Social Care and the National Centre for Health Outcomes Development, London (<http://www.nchod.nhs.uk>, accessed May 2007).
- [9] Rudd A., Goldacre M., Amess M., Fletcher J., Wilkinson E., Mason A., Fairfield G., Eastwood A., Cleary R., Coles J., (eds) (1999). *Health Outcome Indicators: Stroke. Report of a Working Group to the Department of Health*, National Centre for Health Outcomes Development, Oxford, (<http://www.nchod.nhs.uk>, potential indicators, accessed May 2007).
- [10] Haward R., Goldacre M., Mason A., Wilkinson E., Amess M., (eds) (1999). *Health Outcome Indicators: Breast Cancer. Report of a Working Group to the Department of Health*, National Centre for Health Outcomes Development, Oxford, (www.nchod.nhs.uk, potential indicators, accessed May 2007).
- [11] Charlwood P., Mason A., Goldacre M., Cleary R., Wilkinson E., (eds) (1999). *Health Outcome Indicators: Severe Mental Illness. Report of a Working Group to the Department of Health*, National Centre for Health Outcomes Development, Oxford, (www.nchod.nhs.uk, potential indicators, accessed May 2007).

- [12] Lakhani, A., Buxton, A. & Black, N. (eds) (2006). *Directory of Clinical Databases (DoCDat) – Clinical and Health Outcomes Knowledge Base*, The Information Centre for Health and Social Care and the National Centre for Health Outcomes Development, London (<http://www.nchod.nhs.uk>, accessed May 2007).
- [13] Garratt, A.M., Schmidt, L., Mackintosh, A. & Fitzpatrick, R. (2002). Quality of life measurement: bibliographic study of patient assessed health outcome measures, *British Medical Journal* **324**, 1417–1419.

AZIM D. H. LAKHANI

Statistical Issues in Vaccine Safety Evaluation

Introduction

Mass vaccination programmes have greatly reduced the incidence of many common infections of childhood. However, since vaccines are often administered to healthy children, there is heightened public interest in their safety. Because vaccines are used on a large scale, rare reactions can result in substantial numbers of adverse events, with potentially disastrous consequences for public confidence in vaccination programmes. Evidence of safety, and rapid identification of potential problems, are thus essential to the success of vaccination programmes. In this article, some of the statistical methods used for monitoring vaccine safety are described.

Prelicensure Evaluation of Vaccine Safety

Vaccines are evaluated rigorously for safety and efficacy prior to being licensed. At the developmental stage, an experimental vaccine is evaluated extensively in the laboratory. Once sufficient data have been accumulated in the laboratory, the new vaccine is evaluated in humans. Since the 1940s, such evaluations have been based on the methodology of **clinical trials** [1].

Phase I and Phase II Studies

Phase I studies are the first trials in which the experimental vaccine is used in humans. They are usually conducted in small numbers of adult volunteers. In view of their small size, the safety data available from phase I trials is limited to documenting common reactions and identifying serious or unusual side effects.

The next stage is to evaluate the vaccine in phase II studies, which are conducted in those individuals for whose benefit the vaccine was developed. Phase II studies are usually comparative, involving one or several vaccine groups and a control group, and are undertaken using the methodology of double-blind, randomized controlled trials [2]. Phase II trials are of appreciable size, typically including 50–500

participants. However, sample sizes of this order are sufficient only to compare rates of common adverse reactions.

Phase III Trials

The primary purpose of a phase III vaccine trial is to assess the protective efficacy of the vaccine, and to provide further evidence of safety. Phase III trials typically involve 1000–50 000 participants. They are usually double blind and randomized. A phase III trial provides a unique opportunity to assess the safety of the vaccine in a large study under experimental conditions. However, even the largest phase III vaccine trials are too small, and follow-up periods are too short, to provide robust evidence of safety in respect of rare adverse events, or delayed events. This must be achieved by surveillance after the vaccine is licensed.

Postlicensure Surveillance of Vaccine Safety

Once a vaccine has been shown to be effective, is licensed and is recommended for routine use in mass vaccination programmes, it is necessary to monitor its safety using epidemiological surveillance.

Two Examples

Surveillance is necessary both to identify rare or delayed adverse events that cannot be picked up in clinical trials, and to monitor safety more generally.

Example 1 The Cutter Incident. The 1954 field trial of the Salk inactivated polio vaccine remains to date one of the largest vaccine trials ever undertaken. The trial involved 1.8 million children in the United States, of whom 400 000 were randomly assigned to vaccine or placebo, and the remainder enrolled in an open study. On April 12, 1955 the Salk vaccine was declared to be safe and effective. The vaccine was licensed and mass immunization began a few days later. On April 26, reports were received of children who had become paralyzed after receiving polio vaccine from batches made by Cutter Laboratories. It was found that a production fault had caused live virus to enter these vaccines, causing 163 cases of paralysis and 10 deaths, and undermining public

confidence in an otherwise safe and effective vaccine [3, 4].

Example 2 Swine Flu Vaccine and Guillain-Barré Syndrome. In 1976 concern grew in the United States that an epidemic of influenza A virus of a type commonly known as *swine flu* was about to occur. Accordingly the authorities organized a mass vaccination campaign using swine flu vaccines. Between October 1 and mid-December, some 45 million doses of vaccine were administered. However, on December 16 the immunization programme was suspended following reports of Guillain-Barré syndrome (GBS), an uncommon disorder resulting in paralysis of the limbs, occurring in close temporal relation with swine flu vaccination [5].

Note that in both the Cutter incident and the swine flu episode, the association between vaccination and the adverse event became apparent from spontaneously reported instances of adverse events by medical personnel, owing to the close temporal relationship between vaccination and the adverse event.

Temporal Association and Spontaneous Reporting Systems

Spontaneous reporting systems rely on individuals reporting adverse events which they believe may be causally associated with particular vaccines. Such systems have been formalized in several countries, for example through the “Yellow Card” system in the United Kingdom and the Vaccine Adverse Event Reporting System (VAERS) in the United States [6].

Spontaneous reporting systems, however, suffer from two main shortcomings. First, it is seldom possible to attribute **causality** or estimate effect parameters such as relative risks, since the data are not controlled. It is sometimes possible, however, to identify changes in reporting patterns, using measures such as proportional reporting ratios [7]. A second shortcoming is that reporting may be incomplete or inaccurate. This may mean that some adverse events are missed, or that their importance is underestimated. The experience of Urabe mumps vaccines and aseptic meningitis is a case in point.

Example 3 Urabe Mumps Vaccine and Aseptic Meningitis. Measles, mumps, and rubella (MMR) vaccines were introduced in the United Kingdom in 1987. Spontaneous surveillance methods had

suggested that there was a small risk (between 0.4 and 10 per million doses) of aseptic meningitis after mumps vaccine. A subsequent study found that the true risk of aseptic meningitis in the period 15–35 days after receipt of an MMR vaccine containing the Urabe mumps strain was much higher, about 1 in 11 000 doses [8]. As a result of this finding, Urabe-containing vaccines were replaced by other types of vaccines for which there was no evidence of an association with aseptic meningitis.

Following this experience, active surveillance methods based on record linkage between hospital and vaccination databases were put in place, to provide more reliable estimates of postvaccination risks than those available from spontaneous reporting systems [9].

In most instances, adverse events may occur for reasons unconnected with vaccination. Thus, observing a temporal association provides no evidence of causal association. To demonstrate a causal association, it is necessary to estimate the rate at which events occur in temporal association with vaccination, relative to the rate at which they would have occurred in the absence of vaccination, and provide convincing statistical evidence that this relative rate is greater than 1. The following controversies exemplify the risks of placing too much reliance on close temporal association as an indicator of causality.

Example 4 DTP Vaccine and Brain Damage. The diphtheria-pertussis-tetanus (DTP) vaccine was introduced in the United Kingdom in the 1950s, leading to a large reduction in whooping cough. By 1974, 81% of children in the United Kingdom were vaccinated with the DTP vaccine. However, in 1974, public confidence in the vaccine was undermined following intense media coverage of a report describing 36 cases of neurological complications occurring in close temporal association with DTP vaccination [10]. The idea took hold that DTP vaccine caused brain damage, and as a result the coverage of the vaccine dropped sharply to 31%, resulting in the reemergence of whooping cough epidemics. However, subsequent epidemiological investigations and close scrutiny of the evidence led to the conclusion that severe side effects of the DTP vaccine, including brain damage, are so rare that they defy measurement [11].

Example 5 MMR Vaccine and Autism. In 1998, a report was published suggesting that MMR vaccine might cause autism. This report was based on 11 children who were reported to have suddenly experienced serious behavioral problems within 2 months of receiving the MMR vaccine [12]. The report received wide publicity in the media, and generated a controversy that threatened to undermine the MMR vaccination campaign. MMR vaccine coverage dropped from 92 to 60% in some areas, measles outbreaks occurred resulting in the first death from measles in a decade. Several large epidemiological studies were undertaken, none of which confirmed the association between MMR vaccination and autism [13–15].

In both examples, much publicity was given to serious adverse events occurring shortly after vaccination. However, subsequent epidemiological investigations did not confirm a causal link. It is reasonable to conclude that the observed temporal associations were chance events, the incidence of the events coincidentally peaking at the recommended ages for vaccination.

Postlicensure Study Designs for Evaluating Vaccine Safety

Surveillance systems based on spontaneous reporting play an important role in monitoring vaccination programmes and generating hypotheses. However, they can seldom be used for testing hypotheses and providing reliable estimates of vaccine-associated risks, owing to reporting biases and the difficulty in establishing appropriate baseline rates. Instead, epidemiological studies are required.

Cohort Studies

In **vaccine safety** evaluation, cohort studies are usually retrospective, based on preexisting databases such as the Vaccine Safety Datalink (VSD) database in the United States [16] and the General Practitioner Research Database (GPRD) in the United Kingdom [17], or using several databases combined by record linkage techniques. Two basic types of cohort studies are used: parallel group and risk-interval cohort studies.

Example 6 Organic Mercury in Vaccines and Developmental Disorders. Some vaccines contain

thiomersal, a preservative made from ethyl mercury. Concern has been expressed that the administration of vaccines containing thiomersal might cause developmental disorders in later life. A retrospective parallel group cohort study was undertaken [18]. Exposure groups were defined according to the cumulative dose of ethyl mercury received through the administration of thiomersal-containing vaccines by age 6 months. The subsequent occurrences of developmental disorders in the different exposure groups were then analyzed using **survival analysis** techniques, including Cox regression [19]. Little evidence was found that exposure to thiomersal-containing vaccines cause neurodevelopmental disorders.

In many cases, the biological mechanism through which a vaccine may cause adverse events is short lived. Thus it is appropriate to represent the putative vaccine effect by a step function taking a value β within a predefined risk period after vaccination, and 0 outside this risk period. The value β is the logarithm of the hazard ratio (or incidence ratio) over the risk period. In some cases, it is appropriate to use several risk periods, to represent different effects, or to obtain a more accurate impression of how the vaccine-associated risk varies with time since vaccination. Exposure is thus a time-varying **covariate**.

Example 7 Sudden Infant Deaths and DTP Vaccine. The possible association of DTP vaccine and sudden infant death syndrome (SIDS) was investigated in a risk-interval cohort study [20]. Four post-vaccination risk periods were used: 0–3, 4–7, 8–14, and 15–30 days postvaccination. Other times were allocated to the baseline or control period. There was no evidence of a positive association between DTP vaccination and SIDS: for the 0–3 day risk period, for example, the hazard ratio was 0.18, with 95% **confidence interval** (CI) (0.04, 0.8).

The apparent protective effect in this study may be due to a selection **bias**, immunizations being postponed if the child is in poor health. Such biases potentially affect all risk-interval methods [21].

Case–Control Studies

For rare adverse events, cohort studies require access to very large databases. Case–control studies thus represent an attractive alternative, and are commonly used for the evaluation of vaccine safety [22]. In view of the importance of age effects, especially

for the evaluation of pediatric vaccines, case-control studies are usually individually matched on date of birth, and analyzed by conditional logistic regression [23]. Exposures within each case-control set are determined according to vaccination history prior to the index date, namely the date of onset for the case. Case-control studies for vaccine safety evaluation, like cohort studies, can be broadly classified into two categories according to whether or not the putative vaccine-associated risk is believed to be time limited after vaccination.

Example 8 MMR Vaccination and Autism. A large case-control study was undertaken within the GPRD to investigate the association, if any, between MMR vaccination and autism [15]. One thousand two hundred and ninety-four cases were included, individually matched to up to five controls per case on year of birth, sex, and general practice. The primary exposure variable was ever having received MMR vaccine: thus, in this analysis, the postvaccination risk period is assumed to be unlimited. The adjusted odds ratio was 0.86, 95% CI (0.68, 1.09).

Example 9 Hepatitis B Vaccination and Multiple Sclerosis. Mass vaccination against hepatitis B was introduced in France in 1994, targeted at infants and older children, with a catch-up vaccination programme in adults. By July 1996, 200 cases of multiple sclerosis occurring within 2 months of a hepatitis B vaccination had been reported. Several studies were subsequently undertaken to assess the risk of multiple sclerosis following hepatitis B vaccination.

One case-control study involved 236 new cases of multiple sclerosis and two or three controls per case individually matched on age, gender, date, and center of referral [24]. Two exposures were used: hepatitis B vaccination 0–2 months or 2–12 months prior to the index date. These choices reflect the belief that the postvaccination risk is time limited, and is highest in the 0–2-month risk period after vaccination. A raised but nonsignificant odds ratio was found for the 0–2-month period: 1.8, 95% CI (0.7, 4.6), though when the analysis was restricted to individuals with vaccination certificates the odds ratio dropped to 1.4, 95% CI (0.4, 4.5).

As shown in this example, it is important in all vaccine safety studies, but particularly in risk-interval studies, to use documented (rather than self-reported) data on vaccination histories.

Studies Using Only Cases

A variety of statistical methods based only on cases have been used for evaluating vaccine safety, often exploiting the natural experiments resulting from changes in vaccination schedules [25]. Two generic methods are particularly attractive because they automatically control fixed confounders. These are the self-controlled case series method [26, 27] and the case-crossover method [28]. The self-controlled case series method is derived from a Poisson risk-interval cohort model by conditioning on the total number of events experienced by each individual in the cohort, from which only the cases are informative. The method allows for adjustments for age. In contrast, the case-crossover method is a version of a matched case-control study in which control periods are sampled from the life history of the case prior to the event. In this method it is assumed that there is no time trend in the probability of vaccination.

Example 10 Rotavirus Vaccine and Intussusception. In 1998 a new vaccine against rotavirus infection was introduced in the United States. After receiving reports of intussusception soon after vaccination, the vaccination programme was suspended and an epidemiological study was undertaken [29]. A case series analysis of 432 cases was conducted, with observation period 1–12 months of age. The risk periods 3–7, 8–14, and 15–21 days after each dose of vaccine were used. The analysis found a relative incidence of 58.9, 95% CI (31.7, 110), for the 3–7-day risk period after the first dose, thus confirming a significant association between rotavirus vaccination and intussusception in infants.

Self-controlled cases series work best when the risk period or periods used are short compared to the observation period, namely the period over which cases are ascertained. However, in some circumstances, the method may be used with indefinite risk periods [25, 30].

The self-controlled case series method controls for age effects. In the parametric version of the model these effects are represented by a step function. In the semiparametric model, they are unspecified [31]. Allowing for age effects is essential when dealing with adverse events relating to pediatric vaccines, since the incidence of such events is usually highly age dependent. Since pediatric vaccines are usually administered according to a schedule, vaccination

rates are also highly age dependent, precluding the use of case-crossover methods. However, in adult populations, age effects may be less important, and the case-crossover method may be used.

Example 11 Hepatitis B Vaccine and Relapses in Multiple Sclerosis. A case-crossover study was set up to investigate the possible association between a range of vaccinations, including hepatitis B vaccination, and relapses in multiple sclerosis [32]. Vaccination histories were obtained from 643 individuals with multiple sclerosis who suffered a relapse in 1992–1997. The case period was defined as the 2 months preceding the relapse. The 8 months prior to the case period were split into four 2-month control periods. Hepatitis B vaccination in the case and control periods were documented, and analyzed as in a case–control study under the assumption that vaccination rates did not vary appreciably over the 10-month age range spanned by the control and case periods. The odds ratio for the association between hepatitis B vaccination and relapse was 0.67, 95% CI (0.20, 2.17).

Wider Statistical Issues

Vaccines are to some extent the victims of their own success. After an effective vaccine is introduced, the perceived balance of risks and benefits shifts as memories of the disease fade, to be replaced by more immediate concerns about potential risks. This raises some wider issues. First, evidence of safety or lack of it is often required quickly. Thus, timeliness is an important consideration, as well as statistical **power**. Second, evidence of safety can never be absolute, since it is not possible scientifically to prove that administration of a vaccine carries zero risk. Thus statisticians working on vaccine safety issues must be prepared to engage in a public evidence-based discussion about risks and benefits.

References

- [1] Farrington, P. & Miller, E. (2003). Clinical trials, in *Methods in Molecular Medicine: Vaccine Protocols*, A. Robinson, M.J. Hudson & M.P. Cranage, eds, Humana Press, New Jersey, pp. 335–351.
- [2] Pocock, S.J. (1983). *Clinical Trials: A Practical Approach*, John Wiley & Sons, Chichester.
- [3] Nathanson, N. & Langmuir, A.D. (1963). The Cutter incident: poliomyelitis following formaldehyde-inactivated poliovirus vaccination in the United States during the spring of 1955, I, *American Journal of Hygiene* **78**, 16–18.
- [4] Offit, P.A. (2005). The Cutter incident, 50 years later, *New England Journal of Medicine* **352**, 14.
- [5] Langmuir, A.D., Bregman, D.J., Kurland, L.T., Nathanson, N. & Victor, M. (1984). An epidemiologic and clinical evaluation of Guillain-Barré syndrome reported in association with the administration of swine influenza vaccines, *American Journal of Epidemiology* **119**, 841–879.
- [6] Chen, R.T., Rastogi, S.C., Mullen, J.R., Hayes, S.W., Cochi, S.L., Donlon, J.A. & Wassilak, S.G. (1994). The Vaccine Adverse Event Reporting System (VAERS), *Vaccine* **12**, 542–550.
- [7] Evans, S.J.W., Waller, P.C. & Davis, S. (2001). Use of proportional reporting ratios (PRRs) for signal generation from spontaneous adverse drug reaction reports, *Pharmacoepidemiology and Drug Safety* **10**, 525–526.
- [8] Miller, E., Goldacre, M., Pugh, S., Colville, A., Farrington, P., Flower, A., Nash, J., MacFarlane, L. & Tettmar, R. (1993). Risk of aseptic meningitis after measles, mumps and rubella vaccine in UK children, *Lancet* **341**, 979–982.
- [9] Farrington, P., Pugh, S., Colville, A., Flower, A., Nash, J., Morgan-Capner, P., Rush, M. & Miller, E. (1995). A new method for active surveillance of adverse events from diphtheria tetanus pertussis and measles mumps rubella vaccines, *Lancet* **345**, 567–569.
- [10] Kulenkampff, M., Schwartzman, J.S. & Wilson, J. (1974). Neurological complications of pertussis inoculation, *Archives of Disease in Childhood* **49**, 46–49.
- [11] Gangarosa, E.J., Galazka, A.M., Wolfe, C.R., Phillips, L.M., Gangarosa, R.E., Miller, E. & Chen, R.T. (1998). Impact of anti-vaccine movements on pertussis control: the untold story, *Lancet* **351**, 356–361.
- [12] Wakefield, A.J., Murch, S.H., Anthony, A., Linnell, J., Casson, D.M., Malik, M., Berelowitz, M., Dhillon, A.P., Thomson, M.A., Harvey, P., Valentine, A., Davies, S.E. & Walker-Smith, J.A. (1998). Ileal-lymphoid-nodular hyperplasia, non-specific colitis, and pervasive developmental disorder in children, *Lancet* **351**, 637–641.
- [13] Taylor, B., Miller, E., Farrington, C.P., Petropoulos, M.C., Favot-Mayaud, I., Li, J. & Waight, P.A. (1999). Autism and measles, mumps and rubella vaccine: no epidemiological evidence for a causal association, *Lancet* **353**, 2026–2029.
- [14] Madsen, K.M., Hviid, A., Vestergaard, M., Schendel, D., Wolfahrt, J., Thorsen, P., Olsen, J. & Melbye, M. (2002). A population-based study of measles, mumps and rubella vaccination and autism, *New England Journal of Medicine* **347**, 1477–1482.
- [15] Smeeth, L., Cook, C., Fombonne, E., Heavey, L., Rodrigues, L. C., Smith, P. G. & Hall, A. J. (2004). MMR vaccination and pervasive developmental disorders: a case-control study, *Lancet* **364**, 963–969.

- [16] Chen, R.T., DeStefano, F., Davis, R.L., Jackson, L. A., Thompson, R. S., Mullooly, J. P., Black, S. B., Shinefield, H. R., Vadheim, C. M., Ward, J. I., Marcy, S. M., Glasser, J., Rhodes, P. H., Wassilak, S., Kramarz, P., Swint, E., Walker, D., Okoro, C., Gargiullo, P., Baughman, D., King, D., Benson, P., Immanuel, V., Barlow, W., Bohlke, K., Poly, P. L., Rebolledo, V., Rubanowice, D., Mell, L., Lafrance, M., Huggins, L., Dhillon, J., Zavitskovsky, A., Salazar, A., Adams, E., Shafer, M., Sunderland, M., Habanangas, J., Drake, J., Kramer, J., Drew, L., Olson, K., Mesa, J., Pearson, J. A., Allison, M., Bauck, A., Redmond, N., Lieu, T. A., Ray, P., Lewis, E., Fireman, B. H., Lee, H., Jing, J., Goff, N., March, S. M., Lugg, M., Braun, M. M., Wise, R. P., Salive, M. E., & Caserta, V., (2000). The Vaccine Safety Datalink: immunization research in health maintenance organizations in the USA, *Bulletin of the World Health Organization* **78**, 186–194.
- [17] Walley, T. & Mangtani, A. (1997). The UK general practice research database, *Lancet* **350**, 1097–1099.
- [18] Andrews, N., Miller, E., Grant, A., Stowe, J., Osborne, V. & Taylor, B. (2004). Thimerosal exposure in infants and developmental disorders: a retrospective cohort study in the United Kingdom does not support a causal association, *Pediatrics* **114**, 584–591.
- [19] Kalbfleisch, J.D. & Prentice, R.L. (2002). *The Statistical Analysis of Failure Time Data*, 2nd Edition, John Wiley & Sons, New Jersey.
- [20] Griffin, M.R., Ray, W.A., Livengood, J.R. & Schaffner, W. (1988). Risk of sudden infant death syndrome after immunization with the diphtheria tetanus pertussis vaccine, *New England Journal of Medicine* **319**, 618–623.
- [21] Fine, P.E.M. & Chen, R.T. (1992). Confounding in studies of adverse reactions to vaccines, *American Journal of Epidemiology* **136**, 121–135.
- [22] Rodrigues, L.C. & Smith, P.G. (1999). Use of the case-control approach in vaccine evaluation: efficacy and adverse effects, *Epidemiologic Reviews* **21**, 56–72.
- [23] Collett, D. (1991). *Modelling Binary Data*, Chapman & Hall, London.
- [24] Touzé, E., Fourrier, A., Rue-Fenouche, C., Rende-Oustau V., Jeantaud I., Begaud B. & Alperovitch A (2002). Hepatitis B vaccination and first central nervous system demyelinating event: a case-control study, *Neuroepidemiology* **21**, 180–186.
- [25] Farrington, C.P. (2004). Control without separate controls: evaluation of vaccine safety using case-only methods, *Vaccine* **22**, 2064–2070.
- [26] Farrington, C.P. (1995). Relative incidence estimation from case series for vaccine safety evaluation, *Biometrics* **51**, 228–235.
- [27] Whitaker, H.J., Farrington, C.P., Spiessen, B. & Musonda, P. (2006). Tutorial in biostatistics: the self-controlled case series method, *Statistics in Medicine* **25**, 1768–1798.
- [28] Maclure, M. (1991). The case-crossover design: a method for studying transient effects on the risk of acute events, *American Journal of Epidemiology* **133**, 144–153.
- [29] Murphy, T.V., Gargiullo, P.M., Massoudi, M.S., Nelson D. B., Jumaan A. O., Okoro C. A., Zanardi L. R., Setia S., Fair E., LeBaron C. W., Wharton M. & Livingood J. R. (2001). Intussusception among infants given an oral rotavirus vaccine, *New England Journal of Medicine* **344**, 564–572.
- [30] Farrington, C.P., Miller, E. & Taylor, B. (2001). MMR and autism: further evidence against a causal association, *Vaccine* **19**, 3632–3635.
- [31] Farrington, C.P. & Whitaker, H.J. (2006). Semiparametric analysis of case series data (with discussion), *Journal of the Royal Statistical Society. Series C* **55**, 553–594.
- [32] Confavreux, C., Suissa, S., Saddier, P., Bourdès, V. & Vukusic, S. (2001). Vaccinations and the risk of relapse in multiple sclerosis, *New England Journal of Medicine* **344**, 319–326.

C. PADDY FARRINGTON, HEATHER J. WHITAKER AND MOUNIA N. HOCINE

Self-Controlled Case Series Analysis

Introduction

The self-controlled case series method or case series method for short, was developed by Farrington [1] to investigate rare adverse events associated with routine childhood immunizations. It has been used in several areas of epidemiology, for example, in **vac-****cine** safety studies by Farrington *et al.* [2], Miller *et al.* [3], and Andrews [4]; in nonvaccine pharmacoepidemiology by Mullooly *et al.* [5], Hubbard *et al.* [6], and Hocine *et al.* [7, 8]; and Becker *et al.* [9] used the method to investigate the association between long-haul flights and thromboembolism. However, the case series method can be used more widely, to study associations between adverse events and any time-varying exposures as shown by Farrington and Whitaker [10].

The main advantages of this method are that it is based on data from cases only, it implicitly adjusts for all fixed confounders, and in some circumstances it has high efficiency relative to the cohort method. The case series method combines aspects of the case-control and cohort methods, using retrospectively ascertained exposure histories in cases to estimate the relative incidence (RI) in intervals during or in some defined period after exposure, relative to a control period, as described below.

Case Series Method

The self-controlled case series method is based on data on cases only, but may be derived from an underlying Poisson cohort model as follows. Consider a cohort of individuals who may experience the adverse event of interest over a defined observation period $[c_i, d_i]$. We assume that data are available on event times and exposure histories over the entire observation period (note that exposures occurring after an event are required as well as those occurring prior to an event). For each individual i , the observation period $[c_i, d_i]$ is split into intervals for age-groups j , $j = 0, 1, 2, \dots$ and risk periods k , where $k = 0$ for control periods and $k = 1, 2, \dots$ during exposure risk periods, when an individual is deemed to

be at increased (or reduced) risk compared to control periods. Let n_{ijk} and e_{ijk} denote respectively, the number of events experienced and the time spent by individual i in age-group j and risk period k . Within each such interval, the incidence

$$\lambda_{ijk} = \exp(\varphi_i + \alpha_j + \beta_k) \quad (1)$$

is assumed to be constant and the number of events n_{ijk} is assumed to be Poisson with rate $e_{ijk}\lambda_{ijk}$. In equation (1), φ_i represents an effect for individual i , α_j the effect associated with age-group j , and β_k the effect associated with risk period k , with $\alpha_0 = \beta_0 = 0$. The incidence function (1) during the baseline period is thus $\lambda_{i00} = \exp(\varphi_i)$. The exponentiated quantities $\exp(\beta_k)$ are referred to as *relative incidences*.

Conditioning on the number of events $n_i = \sum_{j,k} n_{ijk}$ observed for individual i during $[c_i, d_i]$, the likelihood is product multinomial with kernel:

$$L(\alpha, \beta) = \prod_{i=1}^N \prod_{j,k} \left(\frac{e_{ijk} \exp(\alpha_j + \beta_k)}{\sum_{r,s} e_{irs} \exp(\alpha_r + \beta_s)} \right)^{n_{ijk}} \quad (2)$$

where N is the number of cases. Note that the individual effects φ_i (which encompass any time-independent confounders) cancel out. Only age (or other time-dependent covariates) need to be modeled, though it is possible to include fixed confounders as interactions with the exposure effect. It should be noted that individuals i with no events ($n_i = 0$) in $[c_i, d_i]$ do not contribute to the conditional likelihood in equation (2), and hence only cases within the underlying cohort need be sampled.

The case series method may also be applied to rare nonrecurrent events. In this case, λ_{ijk} in equation (1) is a hazard function representing the incidence in individuals who have not experienced the event, rather than a Poisson rate. However the log likelihood (2) is valid in the limit $\varphi_i \rightarrow -\infty$. This is because recurrent and nonrecurrent events effectively become indistinguishable except in very large samples. Thus, rare nonrecurrent events may be analyzed using the case series model. For more details see Farrington and Whitaker [10].

2 Self-Controlled Case Series Analysis

Example: ITP and MMR Vaccine

Idiopathic thrombocytopenic purpura (ITP) is a rare, potentially recurrent adverse event, in which abnormal bleeding into the skin occurs due to low blood platelet count. In rare instances, measles, mumps, and rubella (MMR) vaccine is thought to cause ITP. The data in this example are described in Farrington and Whitaker [10], and relate to the first ITP in 35 children aged 1–2 years with MMR histories. The observation periods ran from age 366 to 730 days. In this example, three age-groups are used: 366–487 days ($j = 0$), 488–609 days ($j = 1$),

As an illustration, consider the first child ($i = 1$), who was vaccinated at age 670 days and suffered a first ITP at age 691 days. Figure 1 shows the observation period of this child which was split into seven consecutive intervals according to age and exposure group. The details, together with the corresponding Poisson rates λ_{1jk} are given in Table 1.

The likelihood contribution for child 1 is based on the multinomial probability that the event occurred in the fifth interval (at age 691 days), given that it occurred during his or her observation period:

$$L_1(\alpha_2, \beta_2) = \frac{14e^{\alpha_2 + \beta_2}}{34 + 122e^{\alpha_1} + 60e^{\alpha_2} + 15e^{\alpha_2 + \beta_1} + 14e^{\alpha_2 + \beta_2} + 14e^{\alpha_2 + \beta_3} + 18e^{\alpha_2}} \quad (3)$$

and 610–730 days ($j = 2$). Three 2-week-long post-MMR risk groups were defined: 0–14 days after vaccination ($k = 1$), 15–28 days after vaccination ($k = 2$), and 29–42 days after vaccination ($k = 3$). Multiple risk periods are used here to differentiate between phases of varying severity postvaccination. The control period for each child ($k = 0$) includes all other times within its observation period.

Table 1 Age and exposure intervals and Poisson rates for child 1

Interval (days)	Age level (j)	Risk level (k)	e_{1jk}	λ_{1jk}
[453 – 487]	0	0	34	e^{ϕ_1}
[487 – 609]	1	0	122	$e^{\phi_1 + \alpha_1}$
[609 – 669]	2	0	60	$e^{\phi_1 + \alpha_2}$
[669 – 684]	2	1	15	$e^{\phi_1 + \alpha_2 + \beta_1}$
[684 – 698]	2	2	14	$e^{\phi_1 + \alpha_2 + \beta_2}$
[698 – 712]	2	3	14	$e^{\phi_1 + \alpha_2 + \beta_3}$
[712 – 730]	2	0	18	$e^{\phi_1 + \alpha_2}$

The log likelihood for all 35 children is obtained in a similar way as for child 1.

The RIs for the postvaccination risk periods compared to the control period, which included all days in the observation period before and 43 days or more days after MMR vaccination are given in Table 2. The risk of hospital admission for first ITP was significantly raised 15 to 28 days after MMR vaccination. The overall significance of the vaccine effect may be tested using the likelihood-ratio test. The observed test statistic is 15.49 on 3 **degrees of freedom**;

Table 2 First ITP and MMR vaccine: relative incidence by risk period

Risk period (days after MMR)	Number of events	Relative incidence (RI)	95% confidence interval
Control	22	1	–
0–14	2	1.61	(0.36, 7.16)
15–28	8	7.18	(2.95, 17.49)
29–42	3	3.07	(0.86, 10.92)

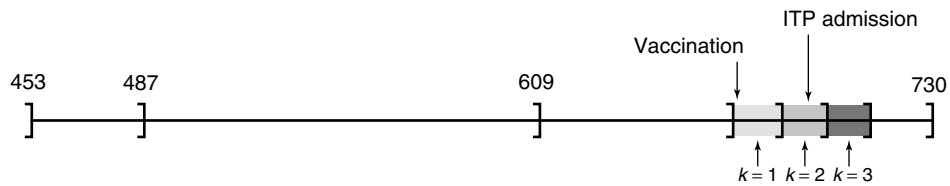


Figure 1 Graphical representation of the observation period for individual $i = 1$ in the ITP and MMR vaccine data

reference to its asymptotic χ^2 (3) distribution yields a p value of less than 0.01.

This analysis was restricted to first ITP. In fact, five children experienced two ITP episodes and one child experienced five. Under the hypothesis that distinct episodes in the same child occur independently, the case series method can be used to analyze all 44 episodes in these 35 children. The results based on all episodes are shown in Table 3. The RIs are slightly lower than for the analysis of first ITPs. Finally, it might be objected that the small numbers of events in the risk periods call into question the validity of the asymptotic theory underpinning the calculation of **confidence intervals** and p values. To allay such concerns, **bootstrap** methods could be used. In a nonparametric bootstrap, cases are resampled with replacement. Alternatively, and more simply, the three postvaccination risk periods could be combined into one 0–42-day risk period. The corresponding RI is then 3.89 (1.79, 8.48) for first ITPs, and 3.23 (1.54, 6.75) for all episodes.

Properties and Implementation

Main Assumptions and Limitations

- The most important and limiting assumption of the case series method is that occurrence of an event does not alter the probability of subsequent exposure. Thus, exposures must be independent of prior events. In some circumstances, such as when an event alters the probability of exposure for only a short period, sensitivity to this assumption can be tested by including a pre-exposure “risk” period. See Farrington and Whitaker [10] for details.
- The observation period must be independent of event times. This can also be restrictive, for

example, if the event increases the short-term mortality rate.

- Events are assumed to arise in a nonhomogeneous **Poisson process**. In practice this is not very restrictive as generally events of interest are rare. Also, as noted above, the method can be used to analyze rare nonrecurrent events.
- It is only possible to estimate the RI: as in case–control studies, the absolute incidence cannot be estimated.
- Variability in the time or age of exposure and event between individuals is required.

Main Advantages

- The need to include only cases reduces **data collection effort**.
- An attractive feature of the case series method is that it is self-controlled; it implicitly adjusts on all non-time-varying confounders, such as genetic and socioeconomic factors.
- Time-dependent confounders such as age can be allowed for in the baseline incidence.
- The method may have high efficiency relative to the underlying cohort method.

Relative Efficiency

By conditioning on the number of events, the case series method loses some efficiency compared with the retrospective cohort method from which it was derived. However, under certain circumstances this loss of efficiency is small. Specifically, efficiency relative to the cohort method is high if the proportion of individuals exposed is high, the length of the risk periods are short in comparison with the observation period and the exposure effect is not large. It is not uncommon for such a situation to arise in practice. On the other hand, the method can be inefficient relative to the cohort method (for a fixed number of cases) if the risk periods are long in comparison with observation periods. Further details are given in Farrington and Whitaker [10].

Designing a Case Series Study

Designing a case series study requires prior specification of the risk periods, and of the observation period of interest. Risk periods are usually defined in relation to a prior hypothesis, which might be based on previous studies or other external knowledge on the

Table 3 All ITP episodes and MMR vaccine: relative incidence by risk period

Risk period (days after MMR)	Number of events	Relative incidence (RI)	95% confidence interval
Control	31	1	–
0–14	2	1.33	(0.31, 5.81)
15–28	8	5.99	(2.55, 14.07)
29–42	3	2.54	(0.73, 8.81)

effect of the exposure on the occurrence of the event. Observation periods are often defined implicitly by a combination of calendar time and age limits: for example, “all cases of acute asthma occurring in children aged 12–23 months inclusive between January 1, 2004 and December 31, 2006”. The observation period of such a case would then stretch from the later of 366 days of age and age (in days) on January 1, 2004 to the earlier of day 730 of age and age on December 31, 2006. Note that it is important to be precise in defining observation periods.

Exposure histories are then documented, using a data source that is independent of event occurrence. Note that all exposures with time at risk within each case’s observation period must be obtained. It is possible only to include cases with a history of exposure, provided that the risk periods of interest are short compared to observation periods (otherwise, there may be confounding between age and exposure effects, which nonexposed cases help to control). Thus, the method is particularly convenient for use with record linkage data based on imperfect identifiers, in which case it is desirable to include only definite matches.

Full details of sample-size formulas for the case series method are given in Musonda *et al.* [11]. Larger sample sizes are required to detect true exposure effects at the same **power** and **significance level** when the risk period is short relative to the observation period, the true RI is close to one, and in the presence of strong age effects, especially if they are monotonic.

Model Fitting

Suitable data, such as hospital admissions data, are usually assembled by listing one record per event. These data must be expanded to list each interval by individual i , age-group j , and risk-group k , along with the number of events that occurred during that interval n_{ijk} and the interval length e_{ijk} . The multinomial likelihood can then be fitted using a Poisson regression model with log link function. The number of events n_{ijk} is the response variable, the log time spent in an interval $\log e_{ijk}$ is offset and the model includes an individual effect γ_i to ensure the fitted total equals the observed values, the age-group effect α_j , and the risk-group effect β_k . The model is thus:

$$n_{ijk} \sim \text{Poisson}(e_{ijk}\lambda_{ijk}) \quad (4)$$

$$\log(\lambda_{ijk}) = \gamma_i + \alpha_j + \beta_k \quad (5)$$

Although this model can be fitted using most statistical software packages, it can be very slow if a study includes a large number of individuals. Fast methods for model fitting are available using the packages R, SAS, STATA, GENSTAT, and GLIM. Examples and instructions of how to expand the data and fit the models using these statistical packages are available from the self-controlled case series website: <http://statistics.open.ac.uk/scs/> [12].

Comparative Studies

The examples below compare the case series method with cohort studies. Example 1 shows the benefit of using the case series method compared to a full cohort study in reducing the data collection effort. Example 2 shows the advantage of using a case series analysis when information on confounders is incomplete.

Example 1: Febrile Convulsions and MMR Vaccine Barlow *et al.* [13] investigated the association between MMR vaccine and febrile seizures using a large cohort of 679 942 children under 7 years, living in the United States. The relative risk of first febrile seizure in the second week after MMR vaccine (8–14 days) was 2.83, with 95% CI (1.44, 5.55). A similar investigation undertaken by Farrington *et al.* [2], which used a case series study of 952 cases aged 12–23 months living in England, showed that the MMR vaccine was associated with febrile convulsions with a relative incidence of 3.04, 95% CI (2.27, 4.07). The case series study gives an estimate of RI close to the value obtained from the full cohort study, which included only 716 cases, but with a narrower confidence interval. When the exposure risk period is short in comparison to the observation period as in this example, the case series analysis has good efficiency relative to a cohort study. Although in this instance, the case series method has greater precision because the cohort study, though very large, included fewer events than the case series study.

Example 2: Asthma and Influenza Vaccine Kramarz *et al.* [14] used data from a cohort of 70 753 asthmatic children aged 1–6 years in 1996 living on the West Coast of the United States to determine whether influenza vaccination precipitates asthma exacerbations. The crude RI of asthma exacerbation

within 2 weeks after influenza vaccination was 3.29, 95% CI (2.55, 4.15), which decreased after adjustment for sex, age, calendar time, indicators of asthma severity, and preventive care practices to 1.39, 95% CI (1.08, 1.77). A case series analysis was undertaken using the 2075 children within the cohort who had at least one asthma exacerbation during the influenza season. The risk of asthma exacerbation within 2 weeks after influenza vaccination was found to be 0.98, 95% CI (0.76, 1.27). One interpretation is that the case series method controls fully for indication **bias** associated with the confounding effect of modeling asthma severity, children with more severe conditions being the most likely to receive the vaccine. In the cohort analysis, only partial control of such confounding may have been achieved.

Some Extensions of the Method

Semiparametric Model

There is also a semiparametric version of the case series model, which avoids misspecification of the age-specific baseline incidence [10]. In this model, the cumulative baseline RI is modeled by a step function with jumps $\exp(\alpha_r)$ at the distinct event ages for all N individuals combined s_r . The semiparametric likelihood for N cases, with case i experiencing n_i events at ages $t_{i1}, \dots, t_{in_i}$, is:

$$L(\beta) = \prod_{i=1}^N \prod_{j=1}^{n_i} \frac{\exp[\alpha_{ij} + \beta x_i(t_{ij})]}{\sum_r \exp[\alpha_r + \beta x_i(s_r)]} \quad (6)$$

where $x_i(t)$ is the exposure at age t for case i , and α_{ij} is the value α_r corresponding to t_{ij} ($j = 1, 2, \dots, n_i$ denotes an event for a specific individual and r denotes an event for any individual). An additional benefit of the semiparametric model is that it implicitly controls for exponential time trends: these are confounded with the age effects but do not influence the exposure effects.

For example, applied to the ITP data described earlier with a single post-MMR risk period of 0–42 days, the semiparametric model yields RI of 3.73 (1.60, 8.71) for first ITPs and 3.01 (1.38, 6.54) for all episodes.

Bivariate Counts

If two types of events, R and S , occur, the effect of an intermittent exposure on the occurrence of R rather

than S , may be evaluated by a conditional RI denoted RRc [8]. The RRc is defined as the ratio between RR^R and RR^S , the relative risks of R and S , respectively, and is estimated by the ratio of their estimates obtained separately by maximizing the appropriate log likelihood (1). The variance of the RRc estimate is simply the sum of the variances of the estimates of RR^R and RR^S under the assumption of independence of the counts of R and S . A test for checking such an assumption was proposed recently by Hocine *et al.* [7]. The test statistic was derived from a conditional likelihood using a bivariate Poisson-generated multinomial model. Robust **estimation** is described in Hocine *et al.* [15].

Other Methods Using Cases Only

Other methods have been proposed to investigate the relationship between an adverse event and a transient exposure using data only on cases. A review of case-only methods used in vaccine safety assessment may be found in Farrington [16]. Prentice *et al.* [17] derived a test for no association from a Cox proportional hazards model, but this test does not yield a readily interpretable association measure. Maclure [18] proposed the case–crossover method, a self-controlled case–control method in which a unique event time is matched to previous time intervals selected from the case’s history. This method yields consistent estimates using conditional logistic regression, only if the joint distribution of exposures is exchangeable as underlined by Vines and Farrington [19] and Whitaker *et al.* [20]. Feldmann [21] used a regression model for rare events with a constant baseline incidence. This model coincides with the case series method when the event of interest is rare and the baseline incidence is constant. A further close connection is to the work of Cox [22] on modulated Poisson processes. Cox’s conditional likelihood for a modulated process is identical to a case series likelihood for a single individual.

References

- [1] Farrington, C.P. (1995). Relative incidence estimation from case series for vaccine safety evaluation, *Biometrics* **51**, 228–235.
- [2] Farrington, P., Pugh, S., Colville, A., Flower, A., Nash, J., Morgan-Capner, P., Rush, M. & Miller, E. (1995). A new method for active surveillance of

- adverse events from diphtheria/tetanus/pertussis and measles/mumps/rubella vaccines, *Lancet* **345**, 567–569.
- [3] Miller, E., Waight, P., Farrington, P., Stowe, J. & Taylor, B. (2001). Idiopathic thrombocytopenic purpura and MMR vaccine, *Archives of Disease in Childhood* **84**, 227–229.
- [4] Andrews, N., Miller, E., Waight, P., Farrington, C.P., Crowcroft, N., Stowe, J. & Taylor, B. (2001). Does oral polio vaccine cause intussusception in infants? Evidence from a sequence of three self-controlled case series studies in the United Kingdom, *European Journal of Epidemiology* **17**, 701–706.
- [5] Mullooly, J.P., Pearson, J., Drew, L., Schuler, R., Maher, J., Gargiullo, P., De Stefano, F. & Chen, R., The Vaccine Safety Datalink Working Group. (2002). Wheezing, lower respiratory disease and vaccination of full-term infants, *Pharmacoepidemiology and Drug Safety* **11**, 21–30.
- [6] Hubbard, R., Farrington, P., Smith, C., Smeeth, L. & Tattersfield, A. (2003). Exposure to tricyclic and selective serotonin reuptake inhibitor antidepressants and the risk of hip fracture, *American Journal of Epidemiology* **158**, 77–84.
- [7] Hocine, M., Guillemot, D., Tubert-Bitter, P. & Moreau, T. (2005). Testing independence between two Poisson-generated multinomial variables in case-series and cohort studies, *Statistics in Medicine* **24**, 4035–4044.
- [8] Hocine, M., Tubert-Bitter, P., Moreau, T., Chavance, M., Varon, E. & Guillemot, D. (2007). A relative risks ratio between antibiotic use and recurrent bacteria colonization in cohort and case-series studies, *Journal of Clinical Epidemiology* **60**(4), 361–365.
- [9] Becker, N.G., Li, Z. & Kelman, C.W. (2004). The effect of transient exposures on the risk of an acute illness with low hazard rate, *Biostatistics* **5**, 239–248.
- [10] Farrington, C.P. & Whitaker, H.J. (2006). Semiparametric analysis of case series data (with discussion), *Applied Statistics* **55**(5), 553–594.
- [11] Musonda, P., Farrington, C.P. & Whitaker, H.J. (2006). Sample sizes for self-controlled case series studies, *Statistics in Medicine* **25**, 2618–2631.
- [12] Whitaker, H.J., Farrington, C.P., Spiessens, B. & Musonda, P. (2006). Tutorial in biostatistics: the self-controlled case series method, *Statistics in Medicine* **25**, 1768–1797.
- [13] Barlow, W.E., Davis, R.L., Glasser, J.W., Rhodes, P.H., Thompson, R.S., Mullooly, J.P., Black, S.B., Shinefield, H.R., Ward, J.I., Marcy, S.M., DeStefano, F. & Chen, R.T., for the Centers for Disease Control and Prevention Vaccine Safety Datalink Working Group. (2001). The risk of seizures after receipt of whole-cell pertussis or measles, mumps, and rubella vaccine, *The New England Journal of Medicine* **345**, 656–661.
- [14] Kramarz, P., DeStefano, F., Gargiullo, P.M., Davis, R.L., Chen, R.T., Mullooly, J.P., Black, S.B., Shinefield, H.R., Bohlke, K., Ward, J.I. & Marcy, S.M., for the Vaccine Safety Datalink Team. (2000). Does influenza vaccination exacerbate asthma? *Archives of Family Medicine* **9**, 617–623.
- [15] Hocine, M.N., Moreau, T. & Chavance, M. (2006). Semiparametric analysis of case series data (Contribution to the discussion), *Applied Statistics* **55**(5), 556–558.
- [16] Farrington, C.P. (2004). Control without separate controls: Evaluation of vaccine safety using case-only methods, *Vaccine* **32**, 1064–1070.
- [17] Prentice, R.L. (1984). On the use of case series to identify disease risk factors, *Biometrics* **40**, 445–458.
- [18] Maclure, M. (1991). The case-crossover design: a method for studying transient effects on the risk of acute events, *American Journal of Epidemiology* **133**, 144–153.
- [19] Vines, S.K. & Farrington, C.P. (2001). Within-subject exposure dependency in case-crossover studies, *Statistics in Medicine* **20**, 3039–3049.
- [20] Whitaker, H.J., Hocine, M.N. & Farrington, C.P. (2007). On case-crossover methods for environmental time series data, *Environmetrics* **18**, 157–171.
- [21] Feldmann, U. (1993). Epidemiologic assessment of risks of adverse reactions associated with intermittent exposure, *Biometrics* **49**, 419–428.
- [22] Cox, D.R. (2000). The statistical analysis of dependencies in point processes, in *Stochastic Point Processes: Statistical Analysis, Theory and Applications*, P.A. Lewis, ed, John Wiley & Sons, New York, pp. 55–66.

MOUNIA N. HOCINE, HEATHER J. WHITAKER
AND C. PADDY FARRINGTON

Early Detection Programs in Medical Care

Introduction

In recent years there is growing attention and interest in preventive medical care [1] as well as early detection programs (EDPs) or screening programs. A screening test is a special examination, which has the potential to diagnose disease while it is asymptomatic and the individual has no signs or symptoms. An EDP is a series of such exams. EDPs for many chronic diseases are now in place. This is especially true for a variety of cancer sites, most notably mammography to detect breast cancer, prostate-specific antigen (PSA) to detect prostate cancer, papanicolaou smears (Pap test) to detect cervical cancer, and fecal occult blood test for colorectal cancer, among others. The motivation for conducting EDPs is the expectation that early diagnosis will result in more cures and longer survival time. Consequently, accurately quantifying the benefit associated with an EDP is of major importance. Mathematical modeling of the disease-screening process, e.g., [2–12], allows us to investigate the performance of various EDPs and study their operating characteristics before we actually carry out any studies.

Different types of study designs, both experimental and observational, have been used to evaluate EDPs [13–15]. Ecological studies correlate mortality and the intensity of screening [16, 17]. Descriptive, uncontrolled studies, based on the experience of individual physicians, hospitals, and nonpopulation-based registries are common [18, 19]. The performance characteristics of various detection tests, i.e., their sensitivity, specificity, and positive predictive value, are generally first reported in such studies. The first evidence that screening may be successful is a stage shift, i.e., an increase in the incidence of early stage cancers coupled with a decrease in late-stage metastatic cancers; later, a reduction in deaths may occur. However, descriptive studies, which lack an appropriate control group and a proper sampling framework, cannot establish efficacy. More rigorous studies include randomized **clinical trials** (RCTs) and case-control studies, described in the sections titled “Case-Control Studies” and “Randomized Clinical Trials”, and a variety of prospective designs.

Prospective cohort studies are similar in many aspects to RCTs, but because the comparison groups are not established by **randomization**, they are generally more difficult to analyze and interpret. Examples of such trials are the UK trial of early detection of breast cancer [20–22] and the Nordic cervical cancer screening trials [23, 24], which showed an impressive reduction in incidence and mortality rate.

Notation and Modeling Assumptions

We consider an idealized model for the natural history of disease in which an individual can be in one of four states denoted by S_h , S_p , S_c , and S_d [2, 9]. An individual in the healthy state, S_h , is either disease free, or has a disease which cannot be detected. An individual in the preclinical state, S_p , has detectable disease but is unaware of it. An individual in the clinical state, S_c , has been diagnosed with disease. The state S_d corresponds to death due to disease. The natural history of disease is assumed to be progressive and is represented by the sequence $S_h \rightarrow S_p \rightarrow S_c \rightarrow S_d$. More generally the model can be embedded in a semi-Markov framework [25–27].

In our numerical calculations we assume that the sojourn time distributions in the various states are exponential with different parameters. Based on reported results [28–30] the mean sojourn times (in years) are given by 55 for S_h and 3.5 for S_p . The median sojourn in S_c is 10 years if the disease was detected by regular medical care and 17 years if the disease was detected by an early exam. For prevalent cases (cases detected at the first exam) the figures are 11 and 20 years respectively. The prevalence of pre-clinical disease is assumed to be 7 per 1000 per year. Examination for early detection are given at age(s) or time(s) τ and their sensitivity is denoted by β . These numerical inputs will be used below.

Case-Control Studies

General

The most common method, or study design, used to assess the benefit of screening is the case-control study [31]. Case-control studies for EDPs are conducted by sampling a group of cases and a group of matched controls and comparing them with respect to their screening history. In most studies cases are

2 Early Detection Programs in Medical Care

individuals that have died from the disease whereas the controls have not. The screening history is considered to be the exposure. Note that the definition of an exposure does not depend on whether the disease was actually diagnosed at a screening examination. The key is that both cases and controls are at risk for the outcome and the exposure [32]. For a review of current methodological approaches specialized to case–control studies for evaluating EDPs, see [33], which discusses many aspects of the design focusing on (a) the selection of cases and controls; (b) the definition of exposure; and (c) common biases. Typical biases associated exclusively with EDPs are length **bias** and lead time bias [2, 34–36]. The theoretical basis for the use of case–control studies for evaluating the benefit of screening has been investigated by [37], whose derivations we follow. For a large-scale simulation study to examine biases resulting from exposure definition, matching, and other factors, see [38].

Analysis

We model a case–control study assuming that individuals may have only one screening examination at a fixed age τ . The calculation involving multiple exams yields a qualitatively similar result. Let $E(\tau) = 1$ if the individual had an early detection exam at age τ and $E(\tau) = 0$ otherwise. Let T denote the follow-up time from the exam. The binary outcome variable is denoted by $D_\tau(T)$, i.e., $D_\tau(T) = 1$ if the individual dies of disease in the interval (τ, T) , and $D_\tau(T) = 0$ otherwise.

The odds ratio when the controls are defined as living individuals, with or without disease, is given by

$$\theta_u(\tau, T) = \frac{\Pr[E(\tau) = 1 | D_\tau(T) = 1] \Pr[E(\tau) = 0 | D_\tau(T) = 0]}{\Pr[E(\tau) = 1 | D_\tau(T) = 0] \Pr[E(\tau) = 0 | D_\tau(T) = 1]} \quad (1)$$

This quantity is observed in practice. The equality of the retrospective and prospective odds means that we can calculate the equation (1) theoretically using the expression

$$\frac{\Pr[D_\tau(T) = 1 | E(\tau) = 1] \Pr[D_\tau(T) = 0 | E(\tau) = 0]}{\Pr[D_\tau(T) = 1 | E(\tau) = 0] \Pr[D_\tau(T) = 0 | E(\tau) = 1]} \quad (2)$$

referred to as the *unrestricted odds ratio*. If the controls are defined as living individuals with clinical disease, the calculations are slightly modified. More formally, let $C_\tau(T) = 1$ if an individual is in the clinical state at T , and zero otherwise. In this population

$$D_\tau(T) + C_\tau(T) = 1 \quad (3)$$

for all individuals. The definition of the odds ratio is modified accordingly and conditions on the event equation (3). With some algebra we find that the odds ratio reduces to

$$\theta_r(\tau, T) = \frac{\Pr[D_\tau(T) = 1 | E(\tau) = 1] \Pr[C_\tau(T) = 0 | E(\tau) = 0]}{\Pr[D_\tau(T) = 1 | E(\tau) = 0] \Pr[C_\tau(T) = 0 | E(\tau) = 1]} \quad (4)$$

which we refer to as the *restricted odds ratio*. The derivation of the above probabilities is somewhat involved [37] and builds on the model for the natural history of the disease and the screening process.

Numerical Examples

We illustrate our calculations of the odds ratio using a hypothetical breast cancer screening trial, using inputs described in the section titled “Notation and Modeling Assumptions”. In Figure 1 we plot the inverse of the unrestricted odds ratio (equation (1)) as a function of the follow-up time for several screening ages and sensitivities. The plots for the restricted odds ratio are qualitatively similar.

Note that the odds ratio decays monotonically to unity as the follow-up time T increases. Cases diagnosed at large follow-up times T are likely to have developed the disease after age τ and an early detection exam has no effect for such individuals. Moreover it can be shown that

$$\sup_{T \geq \tau} \theta_u(\tau, T) = 1 + \beta \left(\frac{\lambda_c^s}{\lambda_c} - 1 \right) \quad (5)$$

where β is the sensitivity of the exam and λ_c^s and λ_c are the hazard functions (evaluated at τ) of the sojourn time distributions in the clinical state for a screen-detected and an undetected individual, respectively. Thus if early detection results in a fivefold increase in the mean survival in the clinical state, e.g., from 10 to 50 years for breast

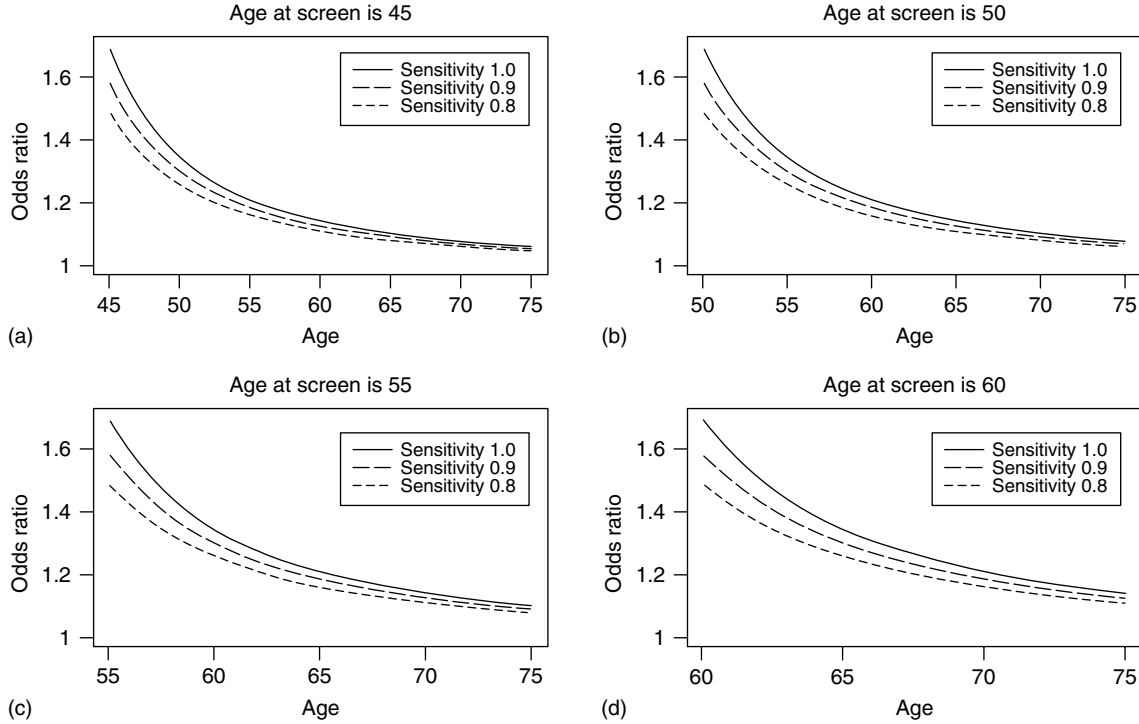


Figure 1 Plot of the unrestricted odds ratio as a function of age for various screening ages and screening exam sensitivities

cancer, which in practice is tantamount to a complete cure, then there will be, *only*, a fivefold increase in the odds ratio. We draw two conclusions from this observation. The first is that relatively modest odd ratios in case-control studies for screening may actually indicate relatively large mortality differences. The second conclusion is that case-control studies result in a significant loss of **power** because even strong clinical effects are associated with odds ratios having a small departure from unity. For example, if a single control is matched to every case, then over 400 pairs are necessary for the power to exceed 50% and almost 1000 pairs are necessary for the power to exceed 80%. Similar results are obtained when more than one control is matched for each case.

The small departure from unity of the odds ratio is explained by the fact that screening is not a regular exposure in the context of case-control studies. Screening has no effect on the survival of an unaffected individual or of an affected but undetected individual. The probability of dying, given a history of screening, is the sum of three terms.

The first is the probability of dying conditional on being in S_h at time τ . The second and third terms are the death probabilities conditional on being in S_p at time τ and being either undetected or detected respectively. The first two terms provide no information about the benefit of screening. Only the third term does. However, the first two terms are larger and dominate the third. A similar effect is observed for the probability of living beyond time T with screening. Consequently the odds ratio does not compare the quantities which directly reflect the effect of early diagnosis. A more meaningful comparison would be for individuals who were undetected in S_p at time τ with matched detected individuals. Unfortunately these are unobservables.

Randomized Clinical Trials

General

The RCT is the method of choice for evaluating the benefit of EDPs and several large-scale trials are currently under way [39–41]. RCTs are designed to

4 Early Detection Programs in Medical Care

eliminate the biases associated with **observational studies**. A typical early detection randomized trial will have a control and a study group. Individuals in the control group receive usual medical care, whereas individuals in the study group will be subject to one or more screening tests. Early detection clinical trials differ from therapeutic trials in that the power is affected by: (a) the number of exams; (b) the time between exams; and (c) the ages at which exams will be given. These design options do not exist in therapeutic trials. Furthermore, long-term follow-up may result in a reduction of power. This feature arises because cases diagnosed in the study group, after the last scheduled examination, may either be those that were evaluated as false negatives or cases that did not have disease at the last examination. It is not possible to distinguish between them. As a result, all deaths are included in the mortality comparisons in order to have an unbiased comparison. A relatively short follow-up time may not have enough deaths to have adequate power. A long follow-up time leads to an increase in the number of diagnosed cases not present at the last scheduled examination. This dilution will result in reduced power. This dilution and the resulting reduction in power are not true for therapeutic trials having a long follow-up time. Another consideration in planning early detection trials is the choice of the target population. In therapeutic trials, all eligible subjects have disease, whereas in early detection trials, subjects do not have disease. However subjects may be incident with disease at a later time. Thus the choice of the target population for the early detection trial is important [42]. A low-risk population may require a very large number of subjects to achieve adequate power.

Analysis

The theoretical foundation for evaluating the benefit of EDPs using RCTs is given in [43]. Their model can be used to investigate various design options such as the number of exams, the optimal spacing between exams, and the length of the follow-up time. Assume that at time τ_0 individuals are registered and randomized into a control and a study group. The study group is offered examinations for the early detection of disease at times $\tau_1 < \tau_2 < \dots < \tau_n$. The i th screening subinterval is denoted by $[\tau_{i-1}, \tau_i)$ for $i = 1, \dots, n$. An individual entering S_p during the i th interval is referred to as an i th-generation individual.

The argument T refers to the chronological time of follow-up; i.e. the total follow-up time is $(T - \tau_0)$.

The probability model calculates the cumulative probability of death for both the control and the study groups under a variety of experimental designs. The cumulative probabilities of death are denoted by $\theta_c(T)$ and $\theta_s(T)$ for the control and the study groups respectively. These probabilities are estimated by $\hat{\theta}_i(T)$, $i = c, s$. The estimators are assumed to follow, in large samples, a normal distribution with mean $\theta_i(T)$ and variance σ_i^2/n_i where n_i is the number of subjects assigned to each group and

$$\sigma_i^2 = \theta_i(T)(1 - \theta_i(T)) \quad (6)$$

The test statistic for testing the equality of these proportions is $Z(T)$ defined by

$$Z(T) = \frac{\sqrt{n} [\hat{\theta}_c(T) - \hat{\theta}_s(T)]}{\sqrt{\hat{\sigma}_c^2 + \hat{\sigma}_s^2}} \quad (7)$$

the large sample power of a one-side α -level test is

$$1 - \Phi [z_{1-\alpha} - \sqrt{n}\delta(T)] \quad (8)$$

where

$$\delta(T) = \frac{\theta_c(T) - \theta_s(T)}{\sqrt{\sigma_c^2 + \sigma_s^2}} \quad (9)$$

and Φ is the standard **normal distribution** function with $\Phi(z_{1-\alpha}) = 1 - \alpha$. The power for detecting a benefit, given by equation (8) is a function of the design parameter of the trial and the sensitivity of the examinations as well as the probability model describing the natural history of disease.

Numerical Examples

We start by investigating the power of a trial comparing: (a) two different schedules (two and four exams), (b) varying the time between exams (1–5 years), and (c) varying the time of the follow-up period (5, 10, 15, optimal in years) for sample sizes of 25 000 in each group. Figure 2(a,b) depict the power calculations for a one-sided level 0.05 test.

The numbers shown on the top curves are the required follow-up years in order to reach the maximum power. The fractions in the figures indicate the proportion of participants receiving all scheduled exams. The calculations show that the greater

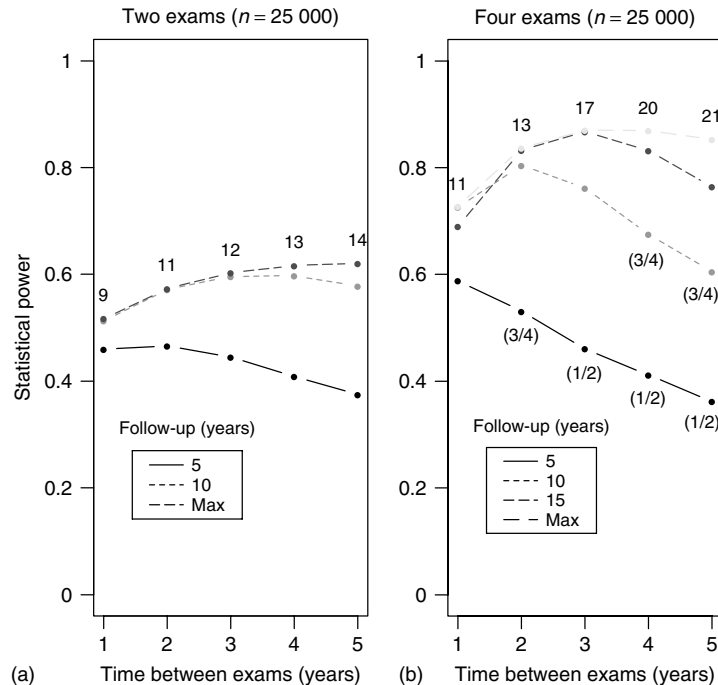


Figure 2 Power calculations for a one-sided level of 0.05 test with $n = 25\,000$ in each group showing relationship with time between exams and follow-up times; (a) and (b) depict results for two and four exams

the interval between examinations, the greater the power, provided the follow-up time is long enough to achieve optimal power. However, larger intervals between examinations require a longer follow-up to reach maximal power. Inadequate follow-up time will result in reduced power with longer intervals between examinations.

There is a trade-off between the number of examinations, time between examinations and optimal follow-up times. For example 80% power can be achieved with a **sample size** of $n = 25\,000$ with four exams spaced 2 years apart, provided the follow-up time is 10 years. Increasing the sample size to $n = 50\,000$ will achieve the same power with two exams spaced 2 years apart or three exams spaced 1 year apart with follow-up times of 7–8 or 5–6 years respectively. Larger number of patients results in a reduction of required follow-up time. Many other features of trial design may be similarly evaluated for more details, see [43].

Experimental Design Issues

The classical clinical design, in which there is a

control and a study group, is an ideal way of planning early detection trials. However, since entry into a clinical trial requires subject consent, there may be resistance in consenting to enter a study in which there is the possibility of being assigned to a control group which has no benefit to the individual. For this reason, other experimental designs have been used to evaluate the benefit of the early detection of breast cancer [44–46]. They have some advantages over a classical design in that subjects in a control group may benefit by participating in the study. We refer to the two experimental designs as the *up-front designs* (UFDs) and *close-out designs* (CODs), see [47] and the references therein. The UFD consists of offering an initial exam to all participants. Then they can be randomized to a usual care group or a screening group. If the outcome of the initial examination is included in the analysis, then the study can answer the question of the benefit of an additional screening program after an initial examination. If the analysis excludes all the cases diagnosed at the initial examination, then the analysis evaluates the benefit of a screening program after elimination of the prevalent

cases. These prevalent cases are most likely to be affected by length bias sampling and consequently will tend to have less aggressive disease and live longer. The COD consists of randomizing subjects to usual care and a screened group. However, the usual care group receives an examination which coincides with the last examination in the control group. The COD suffers from a possible bias. This bias arises because the probability of being in the preclinical state at the close-out exam is higher for the control group relative to the study group; i.e. cases in the study group, which potentially could be present at the last exam, may have been diagnosed on earlier exams and hence removed from the population. As a result there is a higher expected number of false negative cases in the control group when given the close-out exam. Both the UFD and the COD have lower power compared to the classical design.

References

- [1] Davidov, O. & Zelen, M. (2000). Designing cancer prevention trials: a stochastic model approach, *Statistics in Medicine* **19**, 1983–1995.
- [2] Zelen, M. & Feinleib, M. (1969). On the theory of screening for chronic diseases, *Biometrika* **56**, 601–614.
- [3] Prorok, P.C. (1976). The theory of periodic screening I. Lead time and proportion detected, *Advances in Applied Probability* **8**, 127–143.
- [4] Prorok, P.C. (1976). The theory of periodic screening II. Doubly bounded recurrence times and mean lead time and detection probability estimation, *Advances in Applied Probability* **8**, 460–476.
- [5] Albert, A., Gertman, P.M. & Louis, T.A. (1978). Screening for the early detection of disease I. The temporal natural history of a progressive disease, *Mathematical Biosciences* **40**, 1–59.
- [6] Albert, A., Gertman, P.M., Louis, T.A. & Liu, S. (1978). Screening for the early detection of disease II. The impact of screening on the natural history of disease. *Mathematical Biosciences* **40**, 61–109.
- [7] Eddy, D.M. (1980). *Screening for Cancer: Theory, Analysis and Design*, Prentice-Hall.
- [8] Eddy, D.M. (1984). A mathematical model for timing medical tests, *Medical Decision Making* **3**, 34–62.
- [9] Zelen, M. (1993). Optimal scheduling of examinations for the early detection of disease, *Biometrika* **80**, 279–293.
- [10] Yakovlev, A.Y. & Tsodikov, A.D. (1996). *Stochastic Models for Tumor Latency and their Biostatistical Application*, World Scientific.
- [11] Parmigiani, G. (1993). On optimal screening age, *Journal of the American Statistical Association* **88**, 622–688.
- [12] Parmigiani, G. (1997). Timing medical examinations via intensity functions, *Biometrika* **84**, 803–816.
- [13] Cole, P. & Morrison, A.S. (1980). Basic issues in population screening for cancer, *Journal of the National Cancer Institute* **64**(5), 1263–1272.
- [14] Morrison, A.S. (1992). *Screening in Chronic Disease*, Oxford University Press, New York.
- [15] Prorok, P.C. (1984). Evaluation of screening programs for the early detection of cancer, in *Statistical Methods for Cancer Studies*, R.G. Cornell, ed, Marcel Dekker, New York, pp. 225–238.
- [16] Gohagan, J.K., Prorok, P.C., Kramer, B.S. & Cornett, J.E. (1994). Prostate cancer screening in the prostate, lung, colorectal and ovarian cancer screening trial of the national cancer institute, *Journal of Urology* **152**(5 Pt 2), 1905–1909.
- [17] Catalona, W.J., Smith, D.S. & Ornstein, D.K. (1997). Prostate cancer detection in men with serum PSA concentrations of 2.6 to 4.0 ng/mL and benign prostate examination. Enhancement of specificity with free PSA measurements, *Journal of the American Medical Association* **277**, 1452–1455.
- [18] Hisamichi, S., Fukao, A., Fujii, Y., Tsuji, I., Komatsu, S., Inawashiro, H. & Tsubono, Y. (1991). Mass screening for colorectal cancer in Japan. *Cancer Detection and Prevention* **15**(5), 351–356.
- [19] Oshima, A., Hanai, A. & Fujimoto, I. (1979). Evaluation of a mass screening program for stomach cancer, *National Cancer Institute Monographs* **53**, 181–186.
- [20] UK Trial of Early Detection of Breast Cancer Group. (1981). Trial of early detection of breast cancer: description of method, *British Journal of Cancer* **44**, 618–627.
- [21] UK Trial of Early Detection of Breast Cancer Group. (1993). Breast cancer mortality after 10 years in the UK trial of early detection of breast cancer, *The Breast* **2**(1), 13–20.
- [22] Ellman, R., Moss, S.M., Coleman, D. & Chamberlain, J. (1993). Breast self-examination programmes in the trial of early detection of breast cancer: ten year findings, *British Journal of Cancer* **68**, 208–212.
- [23] Hakama, M. (1982). Trends in the incidence of cervical cancer in the Nordic countries, in *Trends in Cancer Incidence Causes and Practical Implications*, K. Magnus, ed, The International Union Against Cancer and the Norwegian Cancer Society, Oslo (Norway), pp. 279–292.
- [24] Laara, E., Day, N.E. & Hakama, M. (1987). Trends in mortality from cervical cancer in the Nordic countries: association with organised screening programmes, *Lancet* **1**(8544), 1247–1249.
- [25] Davidov, O. (1999). The steady state probabilities for regenerative semi Markov processes with application to prevention and screening, *Applied Stochastic Models* **15**, 55–63.
- [26] Limnios, N. & Oprisan, G. (2001). *Semi-Markov Processes and Reliability*, Birkhauser, Boston.

- [27] Parmigiani, G., Skates, S. & Zelen, M. (2002). Modelling and optimization in early detection programs with a single exam, *Biometrics* **58**, 30–36.
- [28] Walter, S.D. & Day, N.E. (1984). Estimation of the duration of a pre-clinical disease state using screening data, *American Journal of Epidemiology* **118**, 865–886.
- [29] Shapiro, S., Venet, W., Strax, P.H., Venet, L. & Rosser, R. (1985). Ten to fourteen year effect of screening on breast cancer mortality, *Journal of the National Cancer Institute* **67**, 65–74.
- [30] Stomper, P.C. & Gelman, R. (1989). Mammography in symptomatic and asymptomatic patients, *Hematology Oncology Clinics of North America* **3**, 611–640.
- [31] Weiss, N.S. (1994). Application of the case control method in the evaluation of screening, *Epidemiologic Reviews* **16**, 102–108.
- [32] Davidov, O. & Herman, A. (2007). *Case Control Studies. In the Encyclopedia of Quality and Reliability*, John Wiley & Sons.
- [33] Cronin, K.A., Weed, D.L., Connor, R.J. & Prorok, P.C. (1999). Case control studies of cancer screening: theory and practice, *Journal of the National Cancer Institute* **90**, 498–504.
- [34] Habbema, J.D.F., Van Oortmarssen, G.J. & Van Putten, D.J. (1983). An analysis of survival differences between clinically and screen detected cancer patients, *Statistics in Medicine* **2**, 279–285.
- [35] Sasco, A.J. (1988). Lead time and length bias in case control studies for the evaluation of screening, *Journal of Clinical Epidemiology* **41**, 103–104.
- [36] Church, T.R. (1999). A novel form of ascertainment bias in case control studies for screening, *Journal of Clinical Epidemiology* **52**, 837–847.
- [37] Davidov, O. & Zelen, M. (2003). The theory of case control studies for early detection programs, *Biostatistics* **4**, 411–421.
- [38] Connor, R.J., Boer, R., Prorok, P.C. & Weed, D.L. (2000). Investigation of design issues in case control studies of cancer screening using micro simulation, *American Journal of Epidemiology* **151**, 991–998.
- [39] Standaert, B. & Denis, L. (1997). The European randomized study of screening for prostate cancer: an update, *Cancer* **80**, 1830–1834.
- [40] Gohagan, J.K., Prorok, P.C., Hayes, R.B. & Kramer, B.S., Prostate, Lung, Colorectal and Ovarian Cancer Screening Trial Project Team. (2000). The Prostate, Lung, Colorectal and Ovarian (PLCO) cancer screening trial of the national cancer institute: history, organization, and status, *Controlled Clinical Trials* **21**(6 Suppl), 251S–272S.
- [41] National Lung Screening Trial (NLST) – National Cancer Institute <http://www.cancer.gov/nlst>.
- [42] Prorok, P.C., Connor, R. & Baker, S. (1990). Statistical consideration in cancer screening programs, *Urologic Clinics of North America* **17**, 699–708.
- [43] Hu, P. & Zelen, M. (1997). Planning clinical trials to evaluate early detection programs, *Biometrika* **84**, 817–829.
- [44] Miller, A.B., Baines, C.J., To, T. & Wall, C. (1992). Canadian national breast screening study: 1. Breast cancer detection and death rates among women aged 40 to 49 years, *Canadian Medical Association Journal* **147**, 1459–1488.
- [45] Frisell, J., Eklund, G., Hellstrom, L., Lidbrink, E., Rutqviste, L.E. & Somell, A. (1991). Randomized study of mammography screening: preliminary report on mortality in the Stockholm trial, *Breast Cancer Research and Treatment* **18**, 49–56.
- [46] Frisell, J., Lidbrink, E., Hellstrom, L. & Rutqviste, L.E. (1997). Followup after 11 years-update of mortality results in the Stockholm mammographic screening trial, *Breast Cancer Research and Treatment* **45**, 263–270.
- [47] Hu, P. & Zelen, M. (2002). Experimental design issues for the early detection of disease: novel designs, *Biostatistics* **3**, 299–313.

ORI DAVIDOV AND PING HU

Case–Control Studies

Introduction

Case–control studies are one of the most powerful and widely used study designs in epidemiology and medical research. They are used to investigate the influence of various risk factors, or exposures, on the development of disease. Case–control studies are conducted by sampling a group of individuals with disease called *cases* and a group of disease-free individuals called *controls* and comparing them with respect to the exposure studied. Note that the comparison is retrospective and conditional on the disease status. Any risk factor or disease may be studied with this methodology. For example a study to investigate the effect of smoking on bladder cancer would sample a group of individuals diagnosed with bladder cancer and a disease-free control group and compare smoking histories between the two groups.

Case–control studies are widely used because they are logistically simpler to carry out and are often more cost effective and time efficient than the competing study designs, namely, cohort studies and cross-sectional studies. In a cohort study a large number of individuals are followed over time and their evolving disease status conditional on their exposure is recorded and compared. Cross-sectional studies, or surveys, record the joint occurrence of disease and risk factors at a specific point in time. Moreover, case–control studies are the only feasible option for studying rare diseases because other designs would require very large sample sizes. It is important to note that major breakthroughs in medical research were achieved using case–control methodology. A notable example is the discovery of the relationship between smoking and lung cancer published by the Surgeon General in 1968 [1]. An extensive literature review [2, 3] showed that similar conclusions were reached by randomized controlled **clinical trials**, the gold standard in medical research, cohort studies and most importantly case–control studies. These investigations further strengthened the validity of case–control studies.

Sampling: The Selection of Cases and Controls

Theoretically, the cases and the controls are a random sample drawn from a hypothetical population of diseased and healthy individuals, respectively. Clearly, valid inferences from case–control studies require that the populations of cases and controls be carefully defined and appropriately sampled.

Case selection (see [4] and [5]), starts by defining the disease under study. Too broad or too narrow a definition should be avoided. Furthermore specific inclusion and exclusion criteria should be adopted to further specify the study population. It is also important to decide if prevalent or incident cases are to be studied. The importance of the proper selection of controls cannot be underestimated. The control group should represent the healthy population and be comparable with the case sample [6]. Similar inclusion and exclusion criteria should be used. Care should be taken to avoid selecting a control who has a disease related to the exposure. For example, choosing individuals after myocardial infarction (MI) as a control group for studying the relationship between smoking and lung cancer will tend to underestimate the risk of smoking because smoking is also a risk factor for MI. Therefore it is more common in MI patients relative to the general population. This type of **bias** can be avoided by selecting a heterogeneous control group, and by comparing several control groups. Controls may also be chosen from among family members, coworkers, or neighbors of the cases. For an extensive discussion on the selection of controls see [7–9]. It is common in practice to match cases and controls [10]. The matching may be at the individual level or the group level where it is called *frequency matching*. The goal of matching is to automatically adjust for confounders, i.e., covariates associated with the disease, most commonly age and sex. However, since the matching covariates may be associated with the exposure **covariate**, matching does not preclude the need for confounding adjustments.

Data Analysis: Basic Models and Methods

The data from a generic case–control study can be displayed in 2×2 table as illustrated in Table 1. The

2 Case-Control Studies

Table 1 Case-control basic data representation^(a)

(a)		
	Case	Control
Exposure	A	B
No exposure	C	D
(b)		
	Esophageal cancer (+)	Esophageal cancer (−)
Smoke (+)	182	533
Smoke (−)	9	255

^(a)In this table we present the basic data representation of the case-control study. The table is a 2×2 table with columns representing the cases/control status, while the rows represent the exposure/no exposure status. Table 1(b) gives a numerical example from a case-control study designed to study the association of esophageal cancer with smoking. For more details see [11]

column (or row) variable represents the disease status and the row (or column) variable is the exposure status.

Let Y be the random variable denoting the disease status, i.e., $Y = 1$ for cases and $Y = 0$ for controls. Let X be the random variable denoting the exposure, with $X = 1$ if exposed and $X = 0$ otherwise. Historically [12], the observation that the prospective odds ratio, which is the parameter of interest, is identical to the retrospective odds ratio, which is the quantity we actually observe, was a key factor in the development of case-control methodology. Formally we have

$$\frac{\Pr(Y = 1|X = 1) \Pr(Y = 0|X = 0)}{\Pr(Y = 0|X = 1) \Pr(Y = 1|X = 0)} = \frac{\Pr(X = 1|Y = 1) \Pr(X = 0|Y = 0)}{\Pr(X = 0|Y = 1) \Pr(X = 1|Y = 0)} \quad (1)$$

where the left-hand side of equation (1) is the prospective odds ratio and the right-hand side is the retrospective odds ratio. Additionally when the disease is rare the odds ratio is close to the relative risk $\Pr(Y = 1|X = 1)/\Pr(Y = 0|X = 1)$. The odds ratio is also called the *cross-product ratio* because it can be calculated directly from the 2×2 table as

$$OR = \frac{AD}{BC} \quad (2)$$

In the example given in Table 1(b) the $OR = (182 \times 255)/(533 \times 9) = 9.6$. Note that an odds ratio of unity indicates that there is no relationship between the exposure and the disease. An odds ratio of 9.6 indicates a very strong relationship. Formally we are interested in testing the hypothesis $H_0 : OR = 1$ versus $H_1 : OR \neq 1$. A test statistic based on the log-odds ratio is often used and its large sample, **normal distribution** is obtained using the δ method. Thus

$$\sqrt{n} [\log(\widehat{OR}) - \log(OR)] \longrightarrow^d N(0, V)$$

where $\widehat{OR}$ is the estimated odds ratio, OR is the true odds ratio, $n = A + B + C + D$ is the **sample size**, and V is the variance of the standardized difference. It can be estimated by $\widehat{V} = 1/A + 1/B + 1/C + 1/D$. For a rigorous development see [13]. In our example $\widehat{V} = 0.24$, and $n = 80$, so a simple standard calculation leads us to a significant p value < 0.001 . A formula for the **confidence interval** for the log-odds ratio is given by

$$\log(\widehat{OR}) \pm Z_{1-\frac{\alpha}{2}} \sqrt{\frac{\widehat{V}}{n}}$$

exponentiating both sides we recover the confidence interval for the odds ratio.

In a widely cited paper Mantel and Haenszel [14] showed how to analyze a sequence of K , 2×2 tables with a common odds ratio. Such data is common after stratification over a confounder variable. Letting A_i , B_i , C_i , and D_i represent the observed cell frequencies in the i th table, the Mantel-Haenszel statistic is

$$\widehat{OR}_{MH} = \frac{\sum_{i=1}^K A_i D_i}{\sum_{i=1}^K B_i C_i} \quad (3)$$

The Mantel-Haenszel statistic is unbiased and consistent even when the observations within strata are dependent and the binomial model fails [15, 16]. Its asymptotic distribution was obtained in two different settings; (a) many small 2×2 tables [17]; and (b) a few large 2×2 tables [18].

A more popular and general approach to the analysis of case-control studies is the logistic regression model which enables researchers to model and

Table 2 Matched case-control data representation^(a)

(a)			
		Cases exposed to the risk factor?	
Control exposed to the risk factor?		Yes	No
	Yes	A'	B'
	No	C'	D'
(b)			
		Cases exposed to the risk factor?	
Control exposed to the risk factor?		Yes	No
	Yes	4	1
	No	6	5

^(a) This table presents data from a matched case-control study for investigating the effect of consumption of specific type (sliced Saxony) ham on a *Salmonella* outbreak in Manhattan (unpublished report <http://staff.pubhealth.ku.dk/~bxc/Epi.2006/slides-cc.pdf>)

analyze the occurrence of disease as a function of categorical and continuous covariates. A major appeal of this model is that a prospective logistic regression model, with continuous or discrete covariates, may be fitted using retrospective data [19]. This discovery further popularized the use of logistic regression models which are currently the main tool in the analysis of a wide variety of case-control designs. The logistic regression model assumes that the probability of disease as a function of the risk factors and other covariates is given by the equation

$$\Pr(Y = 1|X = x, Z = z) = \frac{\exp(\alpha + \beta^T x + \gamma^T z)}{1 + \exp(\alpha + \beta^T x + \gamma^T z)} \quad (4)$$

The parameter α quantifies the baseline risk, β quantifies the direction and strength of the relationship between the exposures and the disease, and γ models the effects of confounders z which are simply covariates in the model. Note that the exposure x may be continuous or discrete, univariate or multivariate. The logistic model is easy to fit and to interpret. In fact the regression parameters β and γ are nothing but the log-odds ratios associated with effects of the exposure(s) and confounder(s). Inferential procedures, that is, **hypothesis testing** and confidence intervals are well developed and statistical software easily accessible [20]. Logistic regression models also

allow for easy incorporation of matching, a common feature in case-control studies. The idea is to view each case and its matching control(s) as consisting of a single strata characterized by a unique intercept α_i in equation (4) as in [21] or [22].

If no confounders are present, matched case-control studies may be analyzed using a 2×2 table such as Table 2(a,b).

Matched pairs with similar exposure are called *concordant*, otherwise they are called *discordant*. The discordant pairs are of principal interest. If there is no relationship between the exposure and the disease, then the probability of observing a discordant matched pair would be $1/2$. More formally we would test the hypothesis $H_0 : \phi = 1/2$ versus $H_1 : \phi \neq 1/2$ where ϕ is the probability of observing a discordant pair estimated by $\hat{\phi} = B'/(B' + C')$. A test based on the statistic

$$\frac{B' - C'}{\sqrt{B' + C'}} \longrightarrow^d N(0, 1)$$

is a common choice. This testing procedure is referred to as *McNemar's test*. The odds ratio is given by $OR = \phi/(1 - \phi)$. Note that for small samples, as in Table 2(b), an exact analysis can be conducted. Under the null hypothesis, B' given $B' + C'$ follows a binomial distribution $Bin(1/2, B' + C')$ and testing is easily conducted; the p value here is 0.0625. Alternatively conditional logistic regression may be used [13, 22]. When more than one control is

matched, the Mantel–Haenszel statistic may be used earlier [23, 24].

Finally we note that more flexible, semiparametric models [25–28] have been fit to case–control data. The advantage of these models is that they make few assumptions regarding the relationship between covariates and outcome. Unfortunately the semiparametric approach does not easily lend itself to statistical tests for evaluating the effect of covariates on the risk of disease.

Variations on the Classical Case–Control Design

Many variations on the basic case–control study design are available. In this section we provide a brief overview and description of some of the options. We start with two variations on the classical design. Hospital-based case–control studies are studies in which both the cases and the controls are sampled from among patients of one or several hospitals. This design has many benefits; it is logistically simple to carry out, the data is usually of high quality, with less recall and nonresponse bias. Its major drawback is that the control population may be different, in important aspects which are difficult to account for, from the at risk population. On the other hand, in population-based case–control studies, the cases and the controls represent the risk population well. However the source population at risk for the disease, also called the *study base*, must be well defined, and randomly sampled. In addition most of the benefits associated with hospital-based studies do not exist [7, 29].

We now mention several designs that extend the classical case–control study. The case–cohort design is a design that combines the strengths of a cohort study with that of a case–control study. At the start of a cohort study a subsample of it is randomly selected. Risk factors for this subsample are carefully measured and monitored. At the end of the follow-up time this subsample is compared to all the cases occurring during the study [30, 31]. Another closely related design is the nested case–control design, which is also cohort based. In this design, the cases are the events in the cohort and the controls are randomly sampled individuals from the cohort who are at risk at the time of each event [32]. Both designs

are useful for studying a rare disease and evaluating the influence of exposures that are either hard or expensive to measure. The sequential case–control study involves repeated hypothesis testing performed as the data accumulates. This is usually done when a new risk factor is discovered and examined and the data accumulates serially [33, 34]. The two-phase case–control study, also referred to as the *double sampling case–control study*, is another design worth mentioning [35, 36]. In this design one first samples a large group of cases and controls from the general population. In a second phase of sampling one samples these groups. The second phase of sampling is differential. For example, one may sample exposed individuals with higher probability than nonexposed. Recently, family-based case–control studies have gained popularity. In these studies, the families of cases are sampled and compared with families of controls. These studies are used to assess the role of genetic factors and their interaction with environmental factors in causing disease [37].

Limitations of Case–Control Studies

Case–control studies are versatile and an increasingly important tool in biomedical research. Nevertheless, despite their usefulness, case–control studies are not foolproof and should always be conducted with care to avoid common biases. Bias is defined as a systemic distortion that leads to an estimated odds ratio that is consistently different from the true odds ratio. Case–control studies are subject to a variety of biases; the most important ones are selection bias, nonresponse bias, and information bias. Measurement error in the exposure variables and unaccounted for confounders may also lead to bias. It is advisable to identify all possible sources of bias in the design phase [38]. Selection bias arises when the cases and controls are not representative of their parent populations. In such situations, the distribution of the exposure in the study population is different from its distribution in the reference population. This bias can be avoided by carefully defining the study base from which cases and controls are sampled [39]. Selection bias may also arise owing to high nonresponse rates among either cases but more commonly in the control population. Care should be taken for the response rate to be as high as possible. In situations where that goal is not

achieved, data analysis should be made to make sure that the responders and nonresponders share the same demographic characteristics. Information bias exists when the exposure is reported or measured differently among cases and controls. Recall bias, the bias associated with reporting on health history, in population selected controls is an example of this type of bias [40]. These biases are forms of the general measurement error problem.

In addition to the aforementioned difficulties, which can be addressed with careful planning, there are some inherent problems associated with case-control studies. Most importantly case-control studies are observational and not randomized studies. Therefore a positive relationship between an exposure and disease cannot be considered to be a causal relationship. This problem is further complicated by the fact that it is never known whether all confounders have been accounted for. Another important drawback of case control studies is that they do not allow for the **estimation** of the overall risk or prevalence of the disease. There are also many subtle statistical issues, for some interesting examples from family studies and from screening studies see [41, 42].

References

- [1] Surgeon General. (1968). *Smoking and Health*, Report of the advisory committee to the surgeon general of the public health service, U.S. Department of Health, Education and Welfare, Washington, DC.
- [2] Benson, K. & Hartz, A.J. (2000). A comparison of observational studies and randomized, controlled trials, *New England Journal of Medicine* **342**, 1878–1886.
- [3] Concato, J., Shah, N. & Horwitz, R.I. (2000). Randomized, controlled trials, observational studies, and the hierarchy of research designs, *New England Journal of Medicine* **342**, 1887–1892.
- [4] Wacholder, S. (1995). Design issues in case-control studies, *Statistical Methods in Medical Research* **4**, 293–309.
- [5] Woodward, M. (1999). *Epidemiology, Study Design and Data Analysis*, Chapman & Hall/CRC, pp. 254–257.
- [6] Woodward, M. (1999). *Epidemiology, Study Design and Data Analysis*, Chapman & Hall/CRC, pp. 257–266.
- [7] Wacholder, S., McLaughlin, J.K., Silverman, D.T. & Mandel, J.S. (1992). Selection of controls in case-control studies. I. Principles, *American Journal of Epidemiology* **135**, 1019–1028.
- [8] Wacholder, S., Silverman, D.T., McLaughlin, J.K. & Mandel, J.S. (1992). Selection of controls in case-control studies. II. Types of controls, *American Journal of Epidemiology* **135**, 1029–1041.
- [9] Wacholder, S., Silverman, D.T., McLaughlin, J.K. & Mandel, J.S. (1992). Selection of controls in case-control studies. III. Design options, *American Journal of Epidemiology* **135**, 1042–1050.
- [10] Rothman, K.J. & Greenland, S. (1998). *Modern Epidemiology*, Lippincott-Raven Publisher, pp. 151–152.
- [11] Breslow, N.E. & Day, N.E. *Statistical Methods in Cancer Research*, IARC Scientific Publications No 32, IARC Lyon.
- [12] Cornfield, J. (1951). A method of estimating comparative rates from clinical data. Applications to cancer of the lung, breast and cervix, *Journal of the National Cancer Institute* **11**, 1269–1275.
- [13] Agresti, A. (1990). *Categorical Data Analysis*, John Wiley & Sons, New York, p. 425.
- [14] Mantel, N. & Haenszel, W. (1959). Statistical aspects of the analysis of data from retrospective studies of disease, *Journal of the National Cancer Institute* **22**, 719–748.
- [15] Liang, K.Y. (1985). Odds-ratio inference with dependent data, *Biometrika* **72**, 678–682.
- [16] Liang, K.Y. (1987). Extended Mantel-Haenszel estimating procedure for multivariate logistic regression models, *Biometrics* **43**, 289–299.
- [17] Robins, J., Greenland, S. & Breslow, N. (1986). Estimators of the variance of the Mantel-Haenszel odds-ratio, *American Journal of Epidemiology* **124**, 719–723.
- [18] Phillips, A. & Holland, P.W. (1987). Estimators of the variance of the Mantel-Haenszel log odds-ratio estimate, *Biometrics* **43**, 425–431.
- [19] Pyke, R. & Prentice, R.L. (1979). Logistic incidence models and case control studies, *Biometrika* **66**, 403–411.
- [20] Kleinbaum, D.G., Klein, D. & Pryor, E.R. (2005). *Logistic Regression*, Springer-Verlag.
- [21] Breslow, N.E. (1995). Statistics in epidemiology: the case control study, *Journal of the American Statistical Association* **91**, 14–28.
- [22] Davidov, O., Murad, H., Lusky, A., Shinwell, E.S., Reichman, B. & Freedman, L. (2003). The analysis of the influence of birth order and other factors in multiple birth data, *Statistics in Medicine* **22**, 3739–3753.
- [23] Miettinen, O.S. (1970). Estimation of relative risk from individually matched series, *Biometrics* **26**, 75–86.
- [24] Woodward, M. (1999). *Epidemiology, Study Design and Data Analysis*, Chapman & Hall/CRC, pp. 273–284.
- [25] Zhao, L.P., Kristal, A.R. & White, E. (1996). Estimating relative risk functions in case control studies using a nonparametric logistic regression, *American Journal of Epidemiology* **144**, 598–609.
- [26] Diggle, P., Morris, S. & Morton-Jones, T. (1999). Case control isotonic regression for investigation of elevation in risk around a point source, *Statistics in Medicine* **18**, 1605–1613.
- [27] Rosenberg, P.S., Katki, H., Swanson, C.A., Brown, L.M., Wacholder, S. & Hoover, R.N. (2003). Quantifying risk factors using non-parametric regression: model

- selection remains the greatest challenge, *Statistics in Medicine* **22**, 3369–3381.
- [28] Ruppert, D., Wand, M.P. & Carroll, R.J. (2003). *Semi-parametric Regression*, Cambridge University Press.
 - [29] West, D.W., Sehaman, K.L., Lyon, J.L., Robinson, L.M. & Allerd, R. (1984). Differences in risk estimation from a hospital and a population-based case-control study, *International Journal of Epidemiology* **13**, 235–239.
 - [30] Prentice, R.L. (1986). A case-cohort design for epidemiologic cohort studies and disease prevalent trials, *Biometrika* **73**, 1–11.
 - [31] Self, S.G. & Prentice, R.L. (1988). Asymptotic distribution theory and efficiency results for case-cohort studies, *Annals of Statistics* **16**, 64–81.
 - [32] Goldstein, L. & Langholz, B. (1992). Asymptotic theory for nested case-control sampling in Cox regression model, *Annals of Statistics* **20**, 1903–1928.
 - [33] O'Neill, R.T. & Anello, C. (1978). Case-control studies: a sequential approach, *American Journal of Epidemiology* **108**, 415–424.
 - [34] Paternack, B.S. & Shore, R.E. (1980). Sample sizes for group sequential cohort and case-control study designs, *American Journal of Epidemiology* **113**, 778–784.
 - [35] Breslow, N.E. & Cain, K.C. (1988). Logistic regression for two-stage case-control data, *Biometrika* **75**, 11–20.
 - [36] Breslow, N.E. & Holubkov, R. (1997). Maximum likelihood estimation of logistic regression parameters under two-phase, outcome-dependent sampling, *Journal of the Royal Statistical Society. Series B* **59**, 447–461.
 - [37] Zhao, L.P., Hsu, L., Davidov, O., Potter, J., Elston, R.C. & Prentice, R.L. (1997a). Population-based family study designs: an interdisciplinary research framework for genetic epidemiology, *Genetic Epidemiology* **14**, 365–388.
 - [38] Skegg, D.C.G. (1988). Potential for bias in case-control studies of oral contraceptives and breast cancer, *American Journal of Epidemiology* **127**, 205–212.
 - [39] Steineck, G. & Ahlbom, A. (1992). A definition of bias founded on the concept of the study base, *Epidemiology* **3**, 477–482.
 - [40] Swan, S.H., Shaw, G.M. & Schulman, J. (1992). Reporting and selection bias in case-control studies of congenital malformations, *Epidemiology* **3**, 356–363.
 - [41] Davidov, O. & Zelen, M. (2001). Referent sampling, family history and relative risk: the role of length biased sampling, *Biostatistics* **2**, 173–181.
 - [42] Davidov, O. & Zelen, M. (2003). The theory of case control studies for early detection programs, *Biostatistics* **4**, 411–421.

ORI DAVIDOV AND AMIR HERMAN

Patient Opinion Measures

Introduction

In healthcare a wide range of patient opinion measures has been developed. Most of these measures are questionnaires, which may be provided to patients by (e-)mail, in face-to-face or telephone contacts, or through the Internet. Many measures have focused on either patient reported symptoms and health-related quality of life, or on patient opinions on healthcare – the latter is the focus of this contribution. Patient opinion measures have been used in scientific research for a long time, but only in recent decades they are increasingly used for quality improvement, accreditation of healthcare providers, and public reporting on quality of care providers. In these applications the users are not necessarily researchers, but citizens, patients, healthcare providers, health authorities, and policy makers. In many countries a market of patient survey research has developed in recent years and users of patient opinion measures have to choose from a range of competing measures.

From a methodological and statistical perspective, there is nothing unique in the validation or application of patient opinion measures compared to other measures. More recent developments in questionnaire methodology, such as computer-assisted interviewing and item–response theory, have not yet reached the field of patient opinion measures. In fact, many patient opinion measures, which are currently used in applications, have hardly been validated. A few measures have been validated reasonably, but it can be observed that methodological issues related to their application of the instrument tend to be ignored [1]. This can seriously **bias** the results, also if the measure itself was carefully developed and validated. This contribution aims to provide an overview of some statistical issues in the application of patient opinion measures. It uses the Europep instrument to illustrate the points made [2] (Box 1).

Methods

Sampling Participants

An appropriate sampling procedure is crucial for the validity of any population-based study. This requires,

among others, a clear definition of the sampling frame (from which cases are sampled) and a sampling procedure that avoids bias (ideally, random sampling). Bias caused by the sampling procedure cannot be compensated for by larger sample sizes. A small but appropriately sampled study population may be preferable to a large but biased sample. New statistical methods for handling biased samples (e.g., sampling participants through the Internet) may widen the opportunities in the future. The Europep instrument was not made for one particular sampling procedure. Many users have consecutively recruited patients among visitors of a general practice, but others have randomly sampled patients from practice registers or population registers.

Response rate can be defined as the number of completed questionnaires divided by the total of questionnaires that reached the target population. A low response rate could induce selection bias. Many studies with the Europep instrument have achieved response rates of 70% or higher. It is recommended to strive for a response rate of at least 60% in using the instrument. Sending reminders increases the response rate by about 10% and repeated mailing of the questionnaire was only marginally more effective than a simple postcard reminder [3]. In many situations it is necessary to use reminders to achieve acceptable response rates. It is recommended to interpret results carefully if the response rate is low (e.g., lower than 60%). While the results may still be valuable for educational feedback and scientific research, they are probably less useful for accreditation and public reporting.

It is difficult to provide recommendations on absolute figures for the **sample size** in studies, which use patient opinion measures. Inclusion of more patients increases the **power** of the study and results in more accurate estimations (smaller **confidence intervals**). If figures are presented for specific practices or practitioners, it should be taken into account that the data are clustered. Statistical advice on appropriate power calculations may be needed. Some studies have suggested that a minimum of 60 respondents are needed per practice/practitioner to allow for reliable figures at the level of practice/practitioner [4]. Lower numbers may be acceptable in educational feedback or scientific research, depending on the research question.

2 Patient Opinion Measures

Box 1 Construction and revision of the Europep questionnaire

The Europep instrument was developed by an international consortium of researchers and family physicians in the years 1995–1998. It was developed from the beginning as an international instrument for patient evaluations of general practice, using rigorous translation and validation procedures. We aimed at use for educational purposes in general practices and regions as well as nationwide surveys and international comparisons. A series of studies were performed for its development, including an international study on patient priorities and studies to examine protoversions of the questionnaire. The questionnaire is focused on evaluations of specific aspects of general practice (not: priorities, wishes, reports, experiences, satisfaction, utilities, etc.). It comprises 23 items, which use a five-point scale (poor–excellent). Explicit criteria were used for the final selection of items, which focused on coverage of domains of general practice, importance to patients, item response, language problems, and discrimination (specific quantitative criteria were formulated). Providing effective feedback on Europep data was not explicitly addressed, but some countries have developed elaborated feedback procedures. More information can be found on www.swisspep.ch/pages/EUROPEP.html.

In 2006, the users of the Europep instrument felt a need to revise and possibly extend the Europep questionnaire. The questionnaire had been used in more than 20 European countries, in some cases on a very large scale (e.g., in Denmark and Switzerland). We sought a balance between improving the questionnaire and maintaining a core set of items – needed for comparisons over time. We planned to exclude and revise items on the basis of a systematic process with a group of users of the Europep instrument. The criteria were specified *a priori*: (a) An item would be omitted or revised if there was a serious ambiguity or translation problem in any country, leading to misinterpretations (item interpretation). (b) An item would be omitted or revised if there would be an unexpected response of lower than 30% in more than half of the countries (item response). (c) An item would be omitted or revised if more than 80% of the respondents used only one of the five answering categories in most (>80%) countries (item discrimination). The EPA project data on Europep data (collected in the period 2000–2005) provided data on item response and discrimination (www.ru.nl/topas-europe). Eight rounds of proposals and feedback were done in the group of users over the email, using a cumulative list of interpretation problems and revision proposals. The criteria regarding item response and discrimination did not lead to omission of items, but a range of interpretation problems was identified which led to a number of small revisions. The result was the prototype of the Europep 2006 questionnaire:

What is your assessment of the general practitioner over the last 12 months with respect to:

1. Making you feel you have time during consultation*
2. Showing interest in your personal situation*
3. Making it easy for you to tell him or her about your problem
4. Involving you in decisions about your medical care
5. Listening to you
6. Keeping your records and data confidential
7. Providing quick relief of your symptoms*
8. Helping you to feel well so that you can perform your normal daily activities
9. Thoroughness of the approach to your problems*
10. Physical examination of you
11. Offering you services for preventing diseases (e.g., screening, health checks, immunizations)
12. Explaining the purpose of examinations, tests, and treatments*
13. Telling you enough about your symptoms and/or illness*
14. Helping you deal with emotions related to your health status*
15. Helping understand why it is important to follow the GP's advice*
16. Knowing what has been done or told during previous contacts in the practice*
17. Preparing you for what to expect from specialists, hospital care, or other care providers*

Box 1 continued

What is your assessment of the general practice over the last 12 months with respect to:

18. The helpfulness of the practice staff (other than the doctor) to you*
19. Getting an appointment to suit you?
20. Getting through to the practice on telephone?
21. Being able to talk to the general practitioner on the telephone*
22. Waiting time in the waiting room?
23. Providing quick services for urgent health problems?

*Revised items

Handling of Items

Health services research is a multidisciplinary research domain, which borrows from clinical epidemiology and social science methodology. The two traditions do not always take the same approach, which is certainly the case in the handling of questionnaire items. The psychometrics tradition regards items as repeated measurements for an underlying concept, so the scores of correlated items should be aggregated to compose a measure of the underlying concept. It is this aggregated measure, which is most relevant for the study. The clinimetrics tradition, on the other hand, tends to regard items either as independent measures (each item reflects one concept) or as aspects of a known biological process (e.g., symptoms of a disease). In the latter case, items are simply added up without a need for an observed high **correlation** between the item scores. In health services research, often a balance has to be found between the two approaches.

In the Europep instrument each item has been developed to reflect a specific concept, such as “availability of the GP on the telephone” or “time in the consultation”. In other words, different items were not primarily designed as repeated measures for underlying concepts. Therefore, items were presented and interpreted separately. Scales are aggregations of different items, which reflect an underlying concept. While the precise meaning of the individual items gets lost, the advantage of scales is that these are usually more accurate than individual items (smaller confidence intervals). Explorative psychometric analyses of Europep data (principal factor analysis, orthogonal rotation, accepting factors with eigenvalue higher than 1) showed that the Europep instrument contains

two scales: an evaluation of clinical management and an evaluation of practice management.

The following procedure was used in the analysis of Europep data. Reliability analysis was used to verify the internal consistency (Cronbach’s α) on each of the two scales in the Europep instrument. Raw scores were used for analysis; cases with missing values were excluded by list-wise deletion. It was felt that Cronbach’s α should be 0.70 or higher (an arbitrary value); in reality, α ’s were much higher. While different procedures for calculation of scale scores may be acceptable, the following procedure was recommended as a standard for Europep users. Items are dichotomized according to 5 *versus* 1–4 (“excellent *vs* less than excellent”). The category “excellent” is coded 1, while the other valid values (excluding the missing values) are coded 0. Then the mean value of the items in the scale is determined. For this purpose a new variable is defined, which is the mean value of the selected items. The result is a score for each individual, which can be interpreted as the proportion of items, which were assessed as “excellent”.

Missing values are a problem for any study, which cannot be fully solved satisfactorily. Many users of the Europep instrument have regarded the category “do not know/not applicable” as a missing value. This implies that cases with missing values were not included in the presentation of scores on the items. It is recommended to Europep users to use this approach as the standard procedure. In the calculation of scales of the Europep instrument, missing values were handled as follows. Patients are excluded if more than one-third of the items in the scale were missing values (no answers or not applicable/not relevant). **Imputation** of missing values in scales

4 Patient Opinion Measures

(e.g., by mean on the other items in the scale) is not recommended, because this artificially increases the consistency of the scale scores.

Presentation of Findings

For educational purposes little more than presentation of absolute numbers and percentages in each answering category for each item may be needed. For comparison with other providers or patient groups potential confounders should be considered. The most obvious confounder in studies of patient views is patient age, although the absolute size of effect of age is small. It has been found fairly consistently that older patients have more positive views of the care received (for which no evidence-based explanation is available). The impact of other patient characteristics is less consistent and probably small, after adjustment for patient age. In the application of the Europep instrument it is recommended to control for age in comparisons, particularly if they are used for accreditation or public reporting.

Furthermore, it is recommended to present all figures with confidence intervals, certainly if they are used for accreditation or public reporting. In many, if not all, applications data are clustered within practices or practitioners (different patients from the same practice/practitioner are included). This implies that random effects models should be used, for which expert statistical advice is often required. Inappropriate statistics are most likely to overestimate differences and changes, because confidence intervals tend to be wider after correction for clustering. This implies that differences between care providers or changes over time are observed, while these are in fact not significant.

In practical applications, the ultimate goal is to improve quality of care (which may be operationalized in different ways). Despite the enormous number of patient opinion studies, the evidence on its impact on quality of care or patient outcomes is very scarce. A randomized trial showed that providing educational feedback on patient opinions of

general practice was positively received by physicians, but did not lead to observable changes in professional behavior or practice management [5]. The use of patient opinion measures in accreditation and public reporting may be more effective, but research evidence to support this claim is currently not available.

Conclusions

Patient opinion measures need to be validated and used sensibly, like any measure. Some statistical aspects of sampling participants, handling of items, and presentation of findings were discussed in this contribution. The Europep instrument was used throughout the paper to illustrate the points made. Statistical aspects of applying patient opinion measures should be addressed to avoid bias, also if the measure itself was carefully developed and validated.

References

- [1] Wensing, M. & Elwyn, G. (2003). Methods for incorporating patients' views in health care: validity, effectiveness and implementation, *British Medical Journal* **326**, 877–879.
- [2] Grol, R., Wensing, M., Mainz, J., Jung, H.P., Ferreira, P., Hearnshaw, H., Hjortdahl, P., Olesen, F., Reis, S., Ribacke, M. & Szecsenyi, J. (2000). Patients in Europe evaluate general practice care: an international comparison, *British Journal of General Practice* **50**, 882–887.
- [3] Wensing, M. & Schattenberg, G. (2005). Response rates after three follow-up procedures in a survey among elderly adults: a randomised trial, *Journal of Clinical Epidemiology* **58**, 959–961.
- [4] Wensing, M., Van der Vleuten, C., Grol, R. & Felling, A. (1997). The reliability of patients' judgements of general practice care: how many questions and patients are needed? *Quality in Health Care* **6**, 80–85.
- [5] Vingerhoets, E., Wensing, M. & Grol, R. (2001). Educational feedback on patients' evaluations of care: a randomised trial, *Quality in Health Care* **10**, 224–228.

MICHEL WENSING

Clinical Behavior Change

Introduction

Clinical Behavior Change: A Systematic Model for Change

Many patients do not receive optimal care. They do not receive, or receive inappropriately, the care they should have been given according to current guidelines for good clinical practice. Or they receive superfluous care. They are exposed to medical interventions that are no longer indicated because of undesirable side effects or proven lack of effect. Clinical behavior change is directed at improvements in patient care, including adoption of evidence-based medicine guidelines, procedures, technologies, or care programmes. Clinical behavior change may also be directed at removal of undesirable (i.e. unnecessary, harmful, or inefficient) routines and variations in the care provided [1].

The number of quality improvement projects in healthcare is increasing throughout the world. Different approaches to clinical behavior change can be observed, each based on different assumptions and theories of human and organizational behavior. Some focus on internal and others on external influences. Despite the varying orientation, most approaches suggest that implementing change effectively requires a systematic, stepwise process. The following steps can be distinguished [2]:

1. Formulation of a concrete, well-developed, and attainable proposal for change in practice, with clear targets.
2. Analysis of the target group and the setting: what are the problems in care provision, and what factors are stimulating or hampering the process of change?
3. Development or selection of a set of strategies for change: strategies for both the effective dissemination and the effective implementation and maintenance of change.
4. Developing and executing an implementation plan that contains activities, tasks, and a timetable.
5. Evaluation and, if necessary, revision of the plan: continuous monitoring on the basis of indicators for the quality of healthcare that is under study.

Interventions for Clinical Behavior Change

Numerous interventions can be used to improve behavior. The best-evaluated interventions are the professional-oriented interventions such as distribution of educational materials, conferences, local consensus processes, audit and feedback, or reminders. Behavior change in professionals can also indirectly be strived for by patient-oriented strategies, e.g., by direct mailings to middle-aged women alerting them on regular screening for cervix cancer and stimulating them to remind their doctor. Another important category of clinical behavior change interventions is on the level of organizational measures, such as, e.g., changes in medical record systems, the introduction of telemedicine, the introduction of multidisciplinary teams, or revision of professional roles. Less well studied are financial measures aimed at quality improvement, such as, e.g., a performance-based payment, or legal regulation and/or rules such as, e.g., change in medical liability or accreditation measures. A taxonomy for clinical behavior change strategies is available on the website of the Cochrane Collaboration Effective Practice and Organisation of Care Group: www.epoc.uottawa.ca.

Statistical Issues in Research on Clinical Behavior Change

Research on clinical behavior change is the scientific study of methods to promote the systematic uptake of clinical research findings into routine practice and hence, to improve the quality and effectiveness of healthcare. It identifies the behavior of healthcare professionals and healthcare organizations as key sources of variance requiring improved empirical and theoretical understanding before effective interventions can be reliably achieved [3].

Both quantitative and qualitative approaches in **data collection** may play their role in clinical behavior change research. In this chapter we focus on statistical issues in the collection of quantitative data for steps 2 and 5 of the implementation model. In step 2 the gap between current care and the proposal for change in practice should be analyzed, and in step 5 the effect of an implementation strategy on indicators for the desired behavior. The validity of measurements of clinical behavior, clustering effects, and the magnitude of interdoctor variation are important issues. In experimental studies on the effectiveness of implementation strategies, the same statistical issues

2 Clinical Behavior Change

play a role as in each (randomized) controlled trial. Clustering is a specific statistical issue in trials on clinical behavior change, as many studies of clinical behavior change are cluster randomized trials, effecting **power** calculation, and level of analysis. We will briefly discuss these issues, illustrated with examples from a study on the test-ordering behavior of general practitioners (GPs) in the Netherlands.

Methods

The Measurement of Clinical Behavior

An all-inclusive measurement of clinical behavior is usually too complex to be possible, so a limited number of crucial aspects of care should be selected for measurement, such as key recommendations in a clinical guideline or those aspects of care that cause most problems to patients. The key recommendations should be translated into quality indicators. A quality indicator is “a measurable element of practice performance for which there is evidence or consensus that it can be used to assess the quality, and hence change in the quality of care provided” [4]. Quality indicators can refer to structures, processes (interpersonal or clinical), or outcomes of care. Examples are given in Table 1 [4].

There is a long debate on the value of outcome indicators *versus* process and structure indicators. Ideally, indicators refer to the effectiveness and safety of healthcare in terms of health status and other patient outcomes. A problem is that many patient outcomes are influenced by many other factors in addition to the process and structure of healthcare, such as patients’ compliance with treatment, hygiene, and nutrition status. In

many studies, indicators refer to processes of care: clinical procedures, communication, and decision-making processes. These indicators tend to be closely linked to professional performance and organization of care, so they can be interpreted more easily in terms of quality of care. Process indicators may have a low or high degree of specificity. This can be illustrated by different examples of prescription rates, ranging from low to high specificity:

- Number of prescriptions per 1000 patients per professional per year.
- Number of prescriptions of an obsolete drug per 1000 patients per professional per year
- Number of a specific first-choice antibiotic drug per 1000 patients per professional per year.
- Number of a specific first-choice antibiotic drug in the correct doses per 1000 patients per professional per year.

General volume data might be easily accessible in existing computerized data sets and can give important information on variation between professionals. The outcome may become more meaningful by integrating patient data in the denominator of the indicator. This is usually expressed as a proportion; for example: the number of first-choice antibiotic drugs/100 patients diagnosed with the disease for which the antibiotic is indicated/professional/year.

The development and validation of indicators for performance or quality of care is a developing science. Existing methodology could be used, although examples of methodological studies of indicators are as yet scarce. For instance, for aggregation of indicators into mean scores, scale construction methods can be used, such as principal factor analysis. Methods from epidemiological diagnostic studies could be used for assessing indicators, if a gold standard is available. Generalizability analysis can be used for assessing the consistency of indicator scores per practice, institution, or geographic region. This chapter focuses on one aspect of the analysis of indicators, which is the statistical clustering of data on clinical behavior.

Clustering of Clinical Behavior

A fundamental assumption of the standard statistics used to analyze outcome indicators on patient level

Table 1 Different types of quality indicators

Indicators	Measures
Structure indicators	Number of professionals (full-time equivalents) per 1000 patients Special clinic rates
Process indicators	Referral rates Prescription rates Vaccination rates
Outcome indicators	Hospital readmission rates Postoperative wound infections rates Pressure ulcers rates

is that the outcome for an individual patient is completely unrelated to that for any other patient – they are said to be “independent”. This assumption is violated if two patients within any one cluster are more likely to respond in a similar manner on the outcome indicator than two patients from different clusters [5]. For example, the management of patients in a single general practice is more likely to be consistent than management of patients across different general practices. The clustering effect may also play its role on the level of professionals and their scoring on process and structure indicators. Two nurses within a hospital might be more likely to respond in a similar manner to an investigation of adherence to a guideline on pressure ulcer treatment than two nurses from different hospitals. Two GPs working closely together in one practice might be more likely to respond in a similar manner to an investigation of their degree to delegate tasks to practice assistants than two GPs from different practices.

A statistical measure of the extent of clustering is the “intracluster **correlation** coefficient” (ICC) which is based on the relationship of the between-cluster to within-cluster variance [5, 6]. According to Eccles *et al.*, ICCs for measures of the process of care tend to be higher than those for measures of patient based outcomes, and ICCs for measures of the process of care in secondary care settings tend to be higher than those for similar measures in primary care settings. A database of ICCs across a range of activities and settings is available from <http://www.abdn.ac.uk/hsru/epp/cluster.shtml>.

Analysis of Clinical Behavior to Describe the Gap between Optimal and Current Care

A wide range of simple or complex data collection methods may be used and all have their advantages and disadvantages. Examples include the manual or electronic abstraction of data from medical records or from existing databases on certain patient groups, e.g., insurance company data, the (prospective) documentation of activities by care providers or patients, and documentation of activities by an observer. In Box 1 an example is given of an analysis of data derived from regional laboratories describing test-ordering rates of GPs in the Netherlands [7]. The aim of the study was to explore the determinants of interdoctor variation. In a study with 229 GPs in 40 local GP groups^a from five regions, the

test-ordering behavior of the GPs was described, to establish professional-oriented and contextual determinants of inclination of the GPs to order tests. The data had a clear hierarchical structure, with GP groups operating under single regional laboratory facilities (regions) and GPs collaborating within GP groups. The test-ordering behavior of two GPs within the same GP group may be more similar than that of two GPs from different GP groups. The same holds for GP groups within a region being perhaps more similar – due to differences in quality assurance activities of the laboratory facility – than two GP groups randomly chosen from different regions. Therefore, the data were modeled in a

Box 1. An example of analysis of current care practice. Variation in test-ordering behavior of general practitioners (GPs)

Data on 19 laboratory and eight imaging tests ordered by GPs, were collected from existing databases in 5 regional laboratory facilities, and cross sectionally analyzed. In a multivariable multilevel regression analysis, these data were linked with **survey data** on GP characteristics such as knowledge about and attitude towards test ordering, and with data on contextual factors such as practice type or distance to the laboratory facility.

Large interdoctor variations were found with skewed distributions. The total median number of tests per GP per year was 998 (interquartile range 663–1500), with large differences between the regions ($p < 0.001$). The response to the survey was 97%. Three variables appeared to be strongly associated with the numbers of tests ordered. “Individual involvement in developing guidelines” (GP level), and “group practice” and “more than 1 year of experience working with a problem-oriented laboratory order form disseminated by the diagnostic center” (contextual level) were associated with reduced numbers of tests ordered of 27, 18, and 41%, respectively, compared with the reference categories.

three-level multilevel analysis model with GP group and region as the random coefficients. The random

variation due to local GP group level proved to be small and insignificant after three GP group level related variables had been omitted. The intraclass correlation coefficient at the region level was 0.304, meaning that the variation between regions was large compared with the variation within regions, which supports the assumption that variability in test ordering is strongly correlated with a region factor.

Contamination in Experiments of Strategies for Clinical Behavior Change

While it is possible to conduct trials of clinical behavior change strategies that allocate individual patients, this may not always be ideal because of contamination in the control group. An example of contamination is an obstetrician who is given an electronic reminder for folic acid advice for women allocated to the intervention arm and not for women in the control group, which after a short while results in the desired behavior in both groups. The ultimate protection against contamination in a trial is to allocate on the level of clusters instead of patients. This results in a group of obstetricians who are exposed to the reminder in all their patient contacts *versus* a group of obstetricians in the control arm who are not exposed at all to the reminder. The disadvantage is that allocation at the level of the professional has lower statistical power than a patient-randomized trial of equivalent size [5].

How to Deal with Clustering in Experiments of Strategies for Clinical Behavior Change

Cluster randomized trials can help to diminish the disadvantage of lower statistical power. In such trials the **randomization** is at one level (professional or organization) and data collection is at a different level (patient) [8]. There are three general approaches to the analysis of cluster randomized trials: analysis at the cluster level; the adjustment of standard tests; and advanced statistical techniques using data recorded at both the individual and cluster levels [8]. Cluster level analyses use the cluster as the unit of randomization and analysis. A summary statistic (e.g., mean, proportion) is computed for each cluster and as each cluster provides only one data point, the data can be considered to be independent, allowing standard statistical tests to be

used. Patient-level analyses can be undertaken using adjustments to simple statistical tests to account for the clustering effect. However this approach does not allow adjustment for patient or practice characteristics. Advances in the development and use of new modeling techniques to incorporate patient-level data allow the inherent correlation within clusters to be modeled explicitly, and thus a “correct” model can be obtained. These methods can incorporate the hierarchical nature of the data into the analysis. For example, in a primary care setting, we may have patients (level 1) treated by a GP (level 2) nested within practices (level 3) and may have covariates measured at the patient level (e.g., patient age or gender), the GP level (e.g., gender, time in practice), and at the practice level (e.g., practice size) [6]. In Box 2 an example is given of an experimental evaluation of a strategy to improve test-ordering behavior of the GPs [9]. In this study, a multifaceted strategy using guidelines, feedback, and social interaction resulted in modest improvements in test ordering by GPs. The strategy consisted of individualized feedback, education in the GP group on guidelines and peer group discussion based on mutual comparison of test-ordering behavior. To account for clustering within local groups, a three-level model was used with the local group as level 3, individual GPs as level 2, and the total number of tests per 6 months per GP as level 1. In multilevel analysis the point estimate and standard deviation were about the same as in the analysis of **covariance** at the individual level and therefore no correction for local groups was needed even though the ICC for the local GP groups was about 0.10.

Sample sizes for cluster randomized trials need to be inflated to adjust for clustering. Both the ICC and the cluster size influence the inflation required; the sample-size inflation can be considerable, especially if the average cluster size is large. The extra numbers of patients required can be achieved by increasing either the number of clusters in the study – the more efficient method [10] – or the number of patients per cluster. In general, little additional power is gained from increasing the number of patients per cluster above 50. Researchers often have to trade-off the logistic difficulties and costs associated with recruitment of extra clusters against those associated with increasing the number of patients per cluster [6]. A power-calculation

Box 2. An example of evaluation of a strategy for clinical behavior change; the effects of a multifaceted strategy aimed at improving the performance of test-ordering behavior of the GP

The strategy was executed among 26 local GP groups including 174 GPs in five regions in the Netherlands. It was evaluated in a cluster randomized controlled trial with a balanced, **incomplete block** design and randomization at group level. Thirteen GP groups underwent the strategy for three clinical problems (arm A; tests for cardiovascular topics, upper and lower abdominal complaints), while 13 other groups underwent the strategy for three other clinical problems (arm B; tests for chronic obstructive pulmonary disease and asthma, general complaints like fatigue, degenerative joint complaints). Each arm acted as a control for the other. During the 6 months of intervention (Jan–June 1999), physicians discussed three consecutive, personal feedback reports in three small group meetings, related them to three evidence-based clinical guidelines, and made plans for change. According to existing national, evidence-based guidelines, a decrease in the total numbers of tests ordered per clinical problem was considered a quality improvement. Baseline data were from July to December 1998, follow-up data were from July to December 1999. All effects were analyzed with analysis of covariance using the number of tests during the follow-up period as the dependent variable and the number of tests at baseline and the region, which appeared to be an important variable, as independent variable.

For clinical problems allocated to arm A, the mean total number of requested tests per 6 months per physician was reduced from baseline to follow-up by 12% among physicians in the arm A intervention, but was unchanged in the arm B control ($P = 0.01$). For clinical problems allocated to arm B, the mean total number of requested tests per 6 months per physician was reduced from baseline to follow-up by 8% among physicians in the arm B intervention, and by 3% in the arm A control ($P = 0.22$).

programme for cluster randomized trials is available at <http://www.abdn.ac.uk/hsru/epp/cluster.shtml>.

Discussion

In this chapter we have stressed the importance of a systematic approach in clinical behavior change research and highlighted the issue of statistical clustering of data on clinical behavior. Considerations are given for the choice of the outcome measure. Careful analysis of clinical behavior to describe the gap between optimal and current care or to find determinants of interdoctor variation may result in useful insights leading to keys for quality improvement. Statistical issues in trials of clinical behavior change research do not differ widely from controlled trials on other topics, but attention has to be paid to the issue of clustering of patients or clinical behavior.

New modeling techniques can incorporate the hierarchical nature of the data into the analysis and at the same time covariates measured at the patient level can be incorporated. Which of the methods

is better to use is still a topic of debate [6]. The most appropriate analysis option will depend on a number of factors including the research question; the unit of inference; the study design; whether the researchers wish to adjust for other relevant variables at the individual or cluster level (covariates); the type and distribution of outcome measure; the number of clusters randomized; the size of cluster and variability of cluster size; and statistical resources available in the research team [6].

End Notes

^a GPs in a local GP group share activities aimed at continuing medical education.

References

- [1] Grol, R. (2005). Implementation of changes in practice, in *Improving Patient Care. The Implementation of Change in Clinical Practice*, R. Grol, M. Wensing & M. Eccles, eds, Elsevier, Butterworth-Heinemann, pp. 6–14.

- [2] Grol, R. & Wensing, M. (2005). Effective implementation: a model, in *Improving Patient Care. The Implementation of Change in Clinical Practise*, R. Grol, M. Wensing & M. Eccles, eds, Elsevier, Butterworth-Heinemann, pp. 41–60.
- [3] Eccles, M. & Mittman, B. *BMC Implementation Science*, <http://www.implementationscience.com/info/about/>.
- [4] Braspenning, J., Campbell, S. & Grol, R. (2005). Measuring changes in patient care: development and use of indicators, in *Improving Patient Care. The Implementation of Change in Clinical Practise*, R. Grol, M. Wensing & M. Eccles, eds, Elsevier, Butterworth-Heinemann, pp. 41–60.
- [5] Eccles, M., Grimshaw, J., Campbell, M. & Ramsay, C. (2003). Research designs for studies evaluating the effectiveness of change and improvement strategies, *Quality and Safety in Health Care* **12**, 47–52.
- [6] Eccles, M., Grimshaw, J., Campbell, M. & Ramsay, C. (2005). Experimental evaluations of change and improvement strategies, in *Improving Patient Care. The Implementation of Change in Clinical Practise*, R. Grol, M. Wensing & M. Eccles, eds, Elsevier, Butterworth-Heinemann, pp. 235–247.
- [7] Verstappen, W.H.J.M., Ter Riet, G., Dubois, W.I., Winkens, R., Grol, R.P.T.M. & van der Weijden, T. (2004). Variation in test ordering behaviour of GPs: professional or context-related factors? *Family Practice* **21**, 387–395.
- [8] Donner, A. & Klar, N. (2000). *Design and Analysis of Cluster Randomization Trials in Health Research*, Arnold, London.
- [9] Verstappen, W.H.J.M., van der Weijden, T., Sijbrandij, J., Smeele, I., Hermsen, J., Grimshaw, J. & Grol, R.P.T.M. (2003). Effect of a practice-based strategy on test ordering performance of primary care physicians. A randomized trial, *Journal of the American Medical Association* **289**, 2407–2412.
- [10] Diwan, V.K., Eriksson, B., Sterky, G. & Tomson, G. (1992). Randomization by group in studying the effect of drug information in primary care, *International Journal of Epidemiology* **21**, 124–130.

TRUDY VAN DER WEIJDEN
AND MICHEL WENSING

Quality in Critical Care Medicine

Introduction

Critical care medicine is a growing element of modern hospital medicine. Critically ill patients consume about of 20% of hospital budgets and approximately 0.8% of the GNP in the United States and 0.2% in Canada [1]. It is therefore a very expensive part of hospital medicine.

Intensive care units (ICUs) can be structured in different ways: according to patient population (such as surgical *vs* medical, neuro, and burns) and according to organization of patient care (open *vs* closed units).

The recent interest in quality improvement in healthcare has not skipped the ICU environment and there are many opportunities to improve on the delivery of critical care provided in ICUs. To describe these, a brief look at the current issues that are problematic in caring for critically ill patients is needed.

The patients in the ICU are frequently the most complex patients in the hospital. These are the patients who require intensive monitoring and often complex and multiple therapeutic interventions. Many drug prescriptions are provided to these patients on a daily basis and this has been shown to be a ground for drug errors, which can lead to adverse patient outcomes [2].

Many factors impact the processes and therefore the outcomes of patients in the ICU. The physical design and equipment available, the organizational workflow of various medical services, and the interactions between the different members of the ICU team: physicians, nurses, pharmacists, physiotherapists, respiratory therapists, and social workers are all critical to achieve optimal outcome, and their availability as well as the communication between them all have a critical impact on the function of the ICU.

Determinants of Quality

A look at quality in ICUs requires observation of a number of components. Design, both physical and organizational setup, of the ICU can impact the

process and clinical outcome. Process, includes the approach to patients with regard to using protocols, workflow of organization, and guidelines, and outcome primarily deals with the clinical outcome of patients admitted to the ICU, but also pertains to other markers of significance such as cost, length of stay, and resource use. An important component of outcome is patient and family satisfaction with the process in the ICU. An admission to intensive care is frequently a very dramatic episode in a person's life and the impact it has on the patient and family can and should be addressed, measured, and improved upon.

Physical design of the ICU can have significant effect on the workflow, on personal and patient satisfaction, and therefore on clinical outcome. Poorly designed work stations can lead to increased rate of medical errors, may impair communication, and increase the risk for nosocomial infections.

ICU Structure: Open *versus* Closed

Many ICUs are organized in an "open" manner. This means that the patients are admitted to the ICU by a primary physician who continues to carry the responsibility for all aspects of the patients' stay in the ICU. This implies that although consultants are used, the treatment of different patients may be quite diverse with some difficulty in implementing routines and protocols for the different clinical issues that may arise. "Closed" units have a primary ICU team that directs care of the patients when they are admitted to the ICU. These clinicians are "dedicated intensivists" and spend a major part of their time in caring for critically ill patients. Studies have shown differences in outcomes between closed and open units with an advantage to a closed unit [3].

The clinical process in the ICU spans the encounter with the patient from the decision on admission to the ICU and the therapy afforded to patients even before they are admitted to the ICU. Rivers and colleagues have shown that early goal directed therapy in patients with severe infections can dramatically reduce mortality. They applied this approach to critically ill patients who required intensive care while they were still in the emergency room waiting for admission to the ICU [4, 5].

Staffing

Staffing of the ICU by nurses affects the quality of healthcare delivered to the patients and to families. This has been shown in the effect of nursing ratio on nosocomial infections (infections acquired in the ICU). A low nursing ratio led to a 30% increase in nosocomial infections. These infections led to increased morbidity and mortality [6]. Other studies have also shown that staffing by nurses can impact various outcome parameters. The availability of other paramedical personnel in the ICU, such as pharmacists, can also lead to improved quality and cost reduction.

Medical Error and Error Prevention

Medical errors are responsible for a great number of adverse events in critically ill patients. A model of medical error was described by Leape [7]. Most errors are due to a system failure rather than a failure by a particular individual. Errors are difficult to reduce because of the complexity of medical systems, and also because healthcare professionals are often reluctant to report and disclose the details of medical errors because of concern of punishment. A comprehensive approach to reducing medical errors requires an open and blame-free attitude. This will allow reporting of incidents that can be collected and studied, and with the use of information from large data-bases allow conclusions as to root causes of errors.

Errors are often divided into those due to faulty planning and those due to faulty execution, and they can be errors of commission or errors of omission. The usual model for medical errors describes the “Swiss cheese model”. The theory behind this model is that all processes have inherent faults in them, the so-called holes, which do not lead to error because each process is composed of many components and the faults in each component need to be aligned to lead to an actual error. To minimize or eliminate errors, efforts must be made to investigate possible errors and near misses, strive to reduce the holes in the system, and allow for backup mechanisms that will decrease the likelihood of a number of faults occurring in the same process in a manner that can lead to error.

Among the various mechanisms designed to reduce errors are the use of computerized systems and

the implementation of appropriate communications protocols between staff. Other important issues are the empowerment of all clinical personnel as well as patients and their families to question a process or a decision and to alert and even abort a process if potential error is suspected.

Evidence-Based Medicine

Many activities and clinical procedures in the ICU are based on personal experience and anecdotal information. There is now a growing body of evidence regarding the best policies and procedures that should be used to achieve the best outcome in critically ill patients. This information can direct the appropriate therapy with regard to ventilatory management in patients with acute respiratory distress syndrome (ARDS) [8], weaning from mechanical ventilation [9], the initial approach to hemodynamic management of patients in septic shock [4], and metabolic management of patients with regard to the control of glucose levels and insulin administration [10].

Despite the fact that these issues have significant data to support their implementation, studies have shown that clinicians are not always ready to adopt evidence-based therapies and their implementation is not universal by any means. For example, following the publication of the National Institutes of Health (NIH) study that showed that reduced tidal volumes can decrease mortality in mechanically ventilated patients with ARDS, there was a consistent but moderate application of this information in the treatment of patients with ARDS as seen in a study in a number of ICUs in the New England area [11].

Even established and standard therapies that are uniformly considered to be part of routine care are not universally complied with. Pronovost *et al.* measured compliance with routine procedures such as adequate sedation, elevation of the head of the bed, deep venous thrombosis (DVT) prophylaxis, and appropriate use of blood transfusion in the ICUs of 13 teaching and community hospitals. They found the compliance for these approaches to be between 33 and 89% [12].

It should be noted that although these and many other interventions have substantial levels of evidence that indicate the appropriate therapy to be followed when a number of choices are available, the treatment of critically ill patients is very complex and chaotic.

The physiologic response to illness and to therapy is very diverse and general approaches may not always be able to indicate the correct or optimal strategy in each and every case. Probably the most significant determinant of outcome in the treatment of complex critically ill patients is the management by an experienced, dedicated clinician.

Guidelines, Protocols, and Bundles

Although the importance of an experienced intensivist cannot be overstated, it is not possible to provide this service to every critically ill patient. This amplifies the importance of treatment provided using the best available clinical knowledge by following protocols and guidelines. Protocols can provide a framework and suggest the optimal diagnostic and therapeutic interventions required to handle various situations.

A new concept regarding the application of best practice techniques is the bundle concept: the combination of a number of proven therapies to achieve a better outcome. There are a number of clinical issues for which the bundle concept has been applied. The concept includes the combination of a number of practices for which there is substantial evidence regarding their effect on improving outcome. At the same time, not all interventions that have a scientific basis to them are included in the bundle since consideration is given to ease of implementation and cost, in deciding on the components of a bundle. Therefore maximal effort to reduce complications should include the bundle as a minimum, but by no means as the sole, component of the intervention.

The Institute of Healthcare Improvement (IHI) has promoted the ventilator-associated pneumonia (VAP) bundle and the central catheter-related infection bundle among others.

Ventilator-Associated Pneumonia

Pneumonia in ventilated patients is a major clinical problem. Critically ill patients requiring ventilatory support may develop pneumonia owing to the treatment, and this nosocomial infection leads to increases in length of stay in the ICU, to increased time of mechanical ventilation and to increased mortality. It is therefore very important to decrease the incidence of VAP to improve outcome. The VAP bundle

includes a number of strategies that have been shown to reduce VAP, and are “bundled” together. There are four components of the ventilator bundle and they are elevation of the head of the bed, daily assessment of patient ability to be extubated by stopping sedation, prophylaxis for upper GI bleeding, and DVT prophylaxis.

This means that when assessing compliance with the approach to reduce VAP, complete application of the bundle is considered as compliance with the bundle. There is no partial application and measurement of complete compliance can be shown to reduce the rate of VAP. It has been shown that compliance with the VAP bundle can be achieved when staff are informed about the importance and scientific validity of the bundle, and are given feedback regarding measured compliance with the bundle. Applying these principles, Cocanour and colleagues have shown a large reduction in the VAP in a trauma ICU [13]. Others have also shown significant reductions in the incidence of VAP in critically ill patients by applying this approach [14].

Central Venous Catheter Infection

Central lines, which are an essential component in patients treated in the ICU, can lead to blood stream infections that can have a significant impact on morbidity and mortality. Almost 50% of patients in the ICU have central lines, which are used for monitoring and delivery of critical drugs. Infections associated with these central lines lead to significant burden in terms of cost, an average cost of a central line associated infection being between \$3700 and \$29 000. There is also a significant attributable increase in length of stay and mortality. The central venous catheter bundle includes the following components: hand hygiene, maximal barrier precautions upon insertion of the central line, chlorhexidine skin antisepsis, optimal catheter site selection, with subclavian vein as the preferred site for nontunneled catheters, and daily review of line necessity with prompt removal of unnecessary lines.

Applying a complex of interventions designed to reduce infections, catheter-related infections could be eliminated in a single institution [15], and in a study of a large number of ICUs it was seen that the application of a group of interventions lead to a significant and sustained reduction in catheter-related infections [16].

Who Regulates? Who Measures?

The increasing concerns over both quality and cost have led many regulatory agencies to try to enforce various process changes into critical care. Not only governmental agencies but privately organized bodies with a large consumer power such as the Leap Frog group have also determined process components that they regard as mandatory to improve outcome. Among these are the implementation of computerized physician order entry, and the staffing of dedicated intensive care physicians 24 h a day all week. The involvement of regulatory agencies is perceived by some clinicians as adding complexity to reporting and regulations without a quality or outcome benefit. However, it is a factor that will probably only increase in magnitude and lead to a growing demand on clinicians and administrators to apply best practice techniques, document, and measure them. The effect of these regulations on quality is as yet unknown.

Summary

Quality improvement in the ICU is an ongoing endeavor that begins with design of the physical aspects of the ICU and workflow structure, entails measurement of different parameters, and impacts outcome of patients, both in terms of requirement of ICU care, length of stay, and mortality.

References

- [1] Jacobs, P. & Noseworthy, T.W. (1990). National estimates of intensive care utilization and costs: Canada and the United States, *Critical Care Medicine* **18**(11), 1282–1286.
- [2] Donchin, Y., Gopher, D., Olin, M., Badihi, Y., Biesky, M., Sprung, C.L., Pizov, R. & Cotev, S. (1995). A look into the nature and causes of human errors in the intensive care unit, *Critical Care Medicine* **23**(2), 294–300.
- [3] Carson, S.S., Stocking, C., Podsadecki, T., Christenson, J., Pohlman, A., MacRae, S., Jordan, J., Humphrey, H., Siegler, M. & Hall, J. (1996). Effects of organizational change in the medical intensive care unit of a teaching hospital: a comparison of 'open' and 'closed' formats, *The Journal of the American Medical Association* **276**(4), 322–328.
- [4] Rivers, E. (2006). The outcome of patients presenting to the emergency department with severe sepsis or septic shock, *Critical Care* **10**(4), 154.
- [5] Rivers, E., Nguyen, B., Havstad, S., Ressler, J., Muzzin, A., Knoblich, B., Peterson, E. & Tomlanovich, M. (2001). Early goal-directed therapy in the treatment of severe sepsis and septic shock, *The New England Journal of Medicine* **345**(19), 1368–1377.
- [6] Hugonnet, S., Chevrolet, J.C. & Pittet, D. (2007). The effect of workload on infection risk in critically ill patients, *Critical Care Medicine* **35**(1), 76–81.
- [7] Leape, L.L. (1997). A systems analysis approach to medical error, *Journal of Evaluation in Clinical Practice* **3**(3), 213–222.
- [8] The Acute Respiratory Distress Syndrome Network. (2000). Ventilation with lower tidal volumes as compared with traditional tidal volumes for acute lung injury and the acute respiratory distress syndrome. *The New England Journal of Medicine* **342**, 1301–1308.
- [9] Esteban, A., Frutos, F., Tobin, M.J., Alia, I., Solsona, J.F., Valverdu, I., Fernandez, R., de la Cal, M.A., Benito, S. & Tomas, R. Spanish Lung Failure Collaborative Group. (1995). A comparison of four methods of weaning patients from mechanical ventilation, *The New England Journal of Medicine* **332**(6), 345–350.
- [10] van den Berghe, G., Wouters, P., Weekers, F., Verwaest, C., Bruyninckx, F., Schetz, M., Vlasselaers, D., Ferdinande, P., Lauwers, P. & Bouillon, R. (2001). Intensive insulin therapy in the critically ill patients, *The New England Journal of Medicine* **345**(19), 1359–1367.
- [11] Young, M.P., Manning, H.L., Wilson, D.L., Mette, S.A., Riker, R.R., Leiter, J.C., Liu, S.K., Bates, J.T. & Parsons, P.E. (2004). Ventilation of patients with acute lung injury and acute respiratory distress syndrome: has new evidence changed clinical practice? *Critical Care Medicine* **32**(6), 1260–1265.
- [12] Pronovost, P.J., Berenholtz, S.M., Ngo, K., McDowell, M., Holzmüller, C., Haraden, C., Resar, R., Rainey, T., Nolan, T. & Dorman T. (2003). Developing and pilot testing quality indicators in the intensive care unit, *Journal of Critical Care* **18**(3), 145–155.
- [13] Cocanour, C.S., Peninger, M., Domonoske, B.D., Li, T., Wright, B., Valdivia, A. & Luther, K.M. (2006). Decreasing ventilator-associated pneumonia in a trauma ICU, *The Journal of Trauma* **61**(1), 122–129; discussion 129–130.
- [14] Youngquist, P., Carroll, M., Farber, M., Macy, D., Madrid, P., Ronning, J. & Susag, A. (2007). Implementing a ventilator bundle in a community hospital, *Joint Commission Journal on Quality and Patient Safety* **33**(4), 219–225.
- [15] Berenholtz, S.M., Pronovost, P.J., Lipsett, P.A., Hobson, D., Earsing, K., Farley, J.E., Milanovich, S., Garrett-Mayer, E., Winters, B.D., Rubin, H.R., Dorman, T. & Perl, T.M. (2004). Eliminating catheter-related bloodstream infections in the intensive care unit, *Critical Care Medicine* **32**(10), 2014–2020.

- [16] Pronovost, P., Needham, D., Berenholtz, S., Sinopoli, D., Chu, H., Cosgrove, S., Sexton, B., Hyzy, R., Welsh, R., Roth, G., Bander, J., Kepros, J. & Goeschel, C. (2006). An intervention to decrease catheter-related bloodstream infections in the ICU, *The New England Journal of Medicine* **355**(26), 2725–2732.

ERAN SEGAL

Population Characteristics of Proficiency Test Data

Introduction

Silicate in seawater, acenaphthylene in sediment, and DDT in fish muscle, are some of the many chemicals found in the sea. Interpretation of information on the concentrations of such chemical contaminants in the environment is highly dependent on the reliability of the analytical measurement. The quality of these measurements is frequently assessed through a laboratory's participation in an appropriate proficiency test (PT) scheme [1, 2]. Homogeneous environmental test materials are provided for analysis and the resultant data from the chemical measurements are compared by the scheme organizer to determine the level of agreement between participants.

A statistical evaluation of the PT should provide a robust and reliable estimate of the true value and the variance of the data, without being influenced by extreme values, many of which are random in origin. The assessment of the individual laboratory performance can then be based on these estimates.

The true value of the concentration of a determinant in a natural PT is usually unknown. Therefore we need to estimate a consensus value from the data provided by the laboratories. When data are generated under statistical control they will generally be normally distributed. We can then use the arithmetic mean and standard deviation (SD) as summary statistics. In reality, however, we find that overlying any **normal distribution** are positively or negatively skewed values that can create bi- or multimodal distributions due to outliers, biased measurements, and left-censored values (LCVs). **Outlier** tests have been used with some caution to eliminate extreme values, but we need alternative methods to evaluate the nonnormal distributions generated in PTs. Robust methods using the Huber [3–5] or Hampel [6, 7] estimators effectively weigh down the effect of skewed and outlying values on the mean where they comprise <10% of the data. These methods, however, cannot cope with highly skewed data, bimodality, or LCVs.

The Cofino model [8, 9] overcomes many of these limitations by taking an alternative approach. Our model identifies clusters of values within a dataset that exhibit a high level of agreement and, for each

cluster, calculates the mean, variance, and the percentage of the whole dataset. This model, therefore, minimizes the contribution of the skewed or outlying data to the mean of the main mode of the data.

Theoretical Background

The Cofino model [8, 9] uses a **probability density function (pdf)** for each observation as a starting point. We obtain a pdf representing all data in a PT by summing each pdf associated with the individual measurement. The overall pdf is the starting point for our model. Instead of calculating the mean of the data, our model establishes the most probable value, given the overall pdf. As an analog to the wave functions in quantum mechanics, observation measurement functions (OMFs, φ_i) are defined as the square root of the pdf attributed to the individual observation, which is the key to calculating the population characteristics. The set of OMFs forms a space, or a basis set, in which we construct the so-called *population measurement functions (PMFs)*. The PMF ψ_i is constructed as a linear combination of OMFs: $\Psi_i = \sum c_{ij}\varphi_j$.

We obtain the coefficients by locating the, unnormalized, PMF₁, with the highest probability in the basis set. This probability is given by the integral $\int \psi_1^2 dc$. We use Lagrange multipliers and impose the additional constraint, that the sum of the squared coefficients is equal to one. The mathematical elaboration requires a solution to the eigenvector–eigenvalue equation $Sc = \lambda c$. In this equation, S represents the matrix of overlap integrals, with S_{ij} given by $\int \phi_i \phi_j dc$. The overlap integral S_{ij} can range between 0 (no overlap) and 1 (100% overlap) when the observations have identical pdfs.

From a set of n basis vectors, OMFs, the model creates a total of n eigenvectors, c , with eigenvalues, λ . The eigenvalue λ_i gives the probability, in the basis set, of the corresponding eigenfunction i . The highest probability and thus maximum value for λ is equal to the number of data n , which is obtained when all data have exactly the same pdf, and each OMF has a coefficient equal to $1/\sqrt{n}$. The eigenvector with the highest eigenvalue λ is the PMF₁. The remaining $n-1$ linear combinations are ranked according to probability (i.e., eigenvalue) and are denoted as PMF₂ . . . PMF _{n} . PMF₂ and higher PMFs may sometimes be additional modes, but are frequently only

2 Population Characteristics of Proficiency Test Data

clusters of data ordered according to their degree of overlap. When the squared OMFs of all PMFs are summed together over the entire concentration range, the pdf of the entire dataset is reconstructed.

For each PMF ψ_i the expectation value and variance can be calculated as follows:

$$\bar{x}_i = \frac{\int c \times \psi_i^2 dc}{\int \psi_i^2 dc} \quad (1)$$

$$s_i^2 = \frac{\int c^2 \times \psi_i^2 dc}{\int \psi_i^2 dc} - \bar{x}_i^2 \quad (2)$$

Also, the eigenvalues λ enable a quantitative assessment of the degree of comparability and the character (unimodal, bimodal) of the dataset. So we can convert the eigenvalue of the mode, or cluster proper into a percentage of the overall pdf. This percentage, therefore, quantitatively describes the fraction of the dataset that is accounted for by the PMF in question.

We have extended the model to include LCVs by defining the appropriate pdfs for these observations. The simplest approach is to assume that any concentration between zero and the limit of quantization (LOQ) has an equal probability. We can then use the square root of a rectangular pdf as basis function. Explicitly, when an LCV is reported, the basis function is equal to $\sqrt{1/LOQ}$ in the interval between zero and LOQ, or otherwise zero. These basis functions have an expectation value $\int \phi_i^2 x dx = LOQ/2$ and a variance $\int \phi_i^2 x^2 dx - \bar{x}^2 = LOQ^2/12$. Other pdf profiles, such as triangular, could be used, but we have concluded that in our studies we have no justification to use pdfs that would require more assumptions.

Inputs for the Model

In our studies we use normal distributions to define the basis functions. Therefore, in addition to the measured values, an estimate of the *within-laboratory* SD is required to define the pdfs and thus the OMFs. The *within-laboratory* SD may be calculated, for instance, when laboratories submit replicate data. The main constraint here is the ability to assess the robustness of the methods used by the laboratories to replicate the measurements.

We have specifically developed the Cofino model using a normal distribution approximation (NDA)

to apply to situations where the reliability of the *within-laboratory* SD cannot be assessed, or when *within-laboratory* SDs are not available at all [8, 9]. The NDA approach is based on the premise that the population characteristics of a normal distribution can be reproduced when each data point is given a normal distribution with the reported concentration as mean and a *within-laboratory* SD constructed as

$$SD_{NDA} = 1.168 \times \text{median}(\text{abs}(x_i - \text{median}(x))) \quad (3)$$

The NDA assumption has been fully tested and verified [8]. We have developed the model for normally distributed populations, but it will also handle nonnormally distributed populations as equally well as the robust techniques developed by Huber and Hampel [8].

Output of the Model

In this application, we have used the Cofino NDA model. We obtain both numerical **descriptive statistics** and a graphical output to aid interpretation of PT data. We have also compared similar information using the Huber and Hampel estimators. The Cofino mean, SD of each mode in the data, and the percentage of data associated with each mode were calculated. The main mode (PMF₁) accounts for >75% of data when the distribution is close to normal. In such cases, the mean and SD we obtained by the different methods of calculation are very comparable. The measurement of silicate in seawater is given as an example (Table 1 and Figure 1).

The three main graphical outputs from the Cofino statistical calculations provide an insight into the structure of the distribution (Figure 1). The pdf conforming for all data (light area) and for data in the main mode (PMF₁) (dark area), allows us to compare the measurements associated with the consensus value with that of the whole distribution. The overlap matrix of each value with all other values is obtained during the calculation which provides us with a map of the homogeneity of the data. The white area (1) denotes overlap (agreement) and the black area (0) denotes no overlap (no agreement). Modes (square areas of similar color) that can be seen within the overlap matrix are not always immediately obvious in the PMF plots (see **Multiple Modes in Proficiency Test Data**). The ranked overview provides us with the individual data (or mean of replicates) with

Table 1 A comparison of the summary statistics from the proficiency test for silicate in seawater and acenaphthylene in sediment

Matrix	Determinant	Total NObs	LCV NObs	Percentage of data in first mode	Units	Cofino		Huber		Cofino SD		Huber SD		Hampel		Hampel SD		Arithmetic		Arithmetic SD	
						mean		mean		mean		mean		mean		mean		mean		mean	
Seawater Sediment	Silicate	47	0	79	$\mu\text{mol l}^{-1}$	9.79	0.70	9.91	0.73	9.88	0.72	9.95	1.72	20.94	17.35	24.76	24.29	9.95	1.72	24.76	24.29
	Acenaphthylene	14	1	69	mg kg^{-1}	12.16	14.79	23.72	25.10	20.94	17.35	24.76	24.29	20.94	17.35	24.76	24.29	24.76	24.29	24.76	24.29

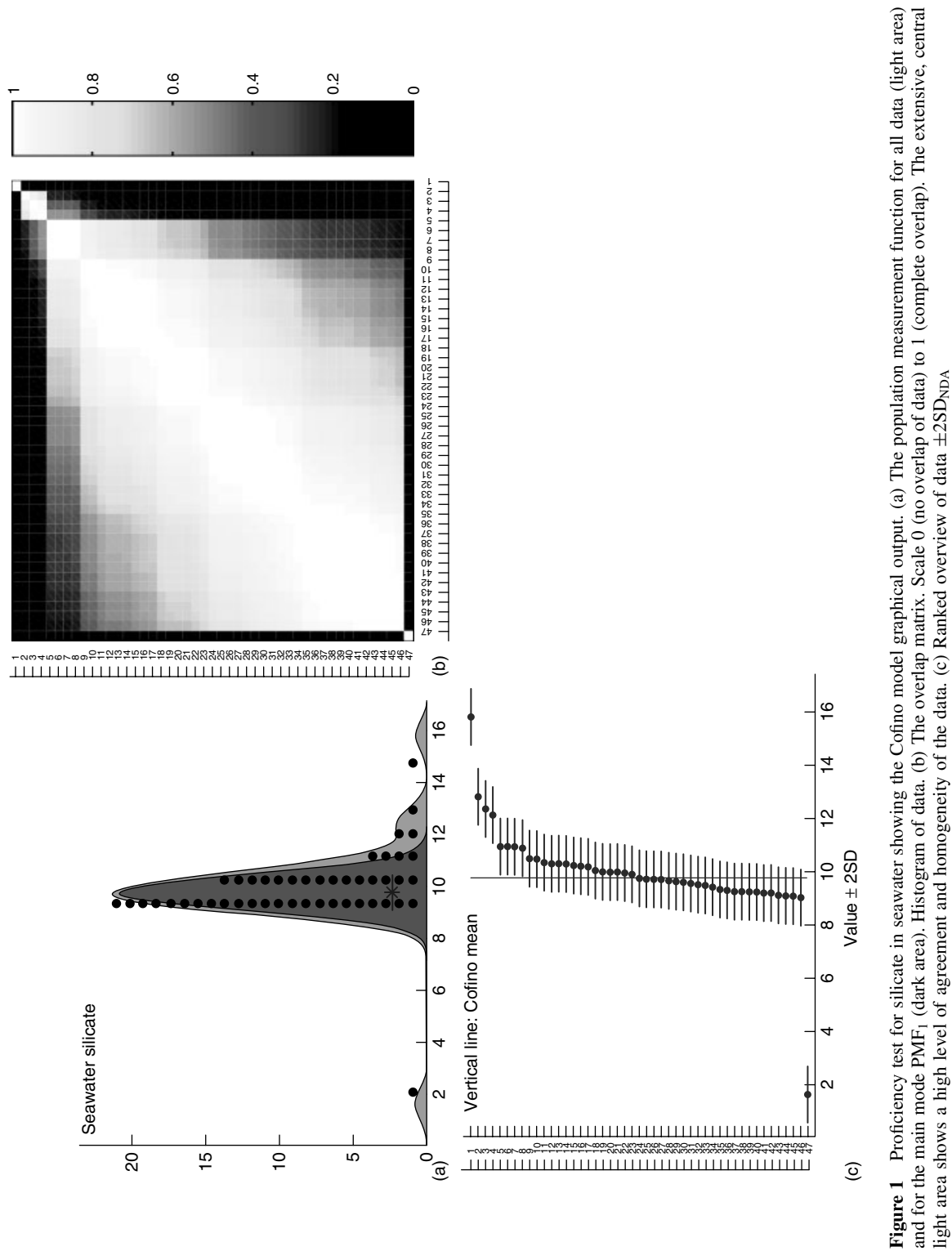


Figure 1 Proficiency test for silicate in seawater showing the Cofino model graphical output. (a) The population measurement function for all data (light area) and for the main mode PMF₁ (dark area). Histogram of data. (b) The overlap matrix. Scale 0 (no overlap of data) to 1 (complete overlap). The extensive, central light area shows a high level of agreement and homogeneity of the data. (c) Ranked overview of data $\pm 2SD_{NDA}$

the SD reconstructed from the NDA (equation 3) and also gives the location of the consensus value in relation to the distribution. The range of the LCVs is also included.

We have provided two examples for comparison. The mean and SD for silicate in seawater determined by the Huber and Hampel robust estimators differ from our model by <2%. So where data closely reflect a normal distribution, different methods of calculation can produce similar results. Larger differences between these methods of calculation arise with skewed and tailing data. The example of data from the analysis of acenaphthylene in sediment (Table 1 and Figure 2) shows almost 100% difference in the mean using the Huber estimator (23 mg kg⁻¹) and the Cofino model (12.2 mg kg⁻¹). The positive tailing values are *ca* 50% of all numerical values (Figure 2) and have a clear influence on the robust estimators compared with the Cofino model. Why is this so? In Figure 2(b) we can see that there is no overlap of the tailing data with the values associated with the first, main mode, PMF₁ and so effectively do not contribute to the construction of the mean of this mode. The Cofino model, therefore, provides a more meaningful consensus value for the purposes of this PT.

The evaluation of the LCVs by the Cofino model allows us to include all data and provide a more complete assessment. In some cases, where the concentration of the determinant is very low, the LCVs may constitute most of the data. Ignoring these LCVs can lead to a biased interpretation. We have, as an example, measurements of *op'*-DDT in fish muscle (Figure 3) where there are 6 numerical values and 11 LCVs. The LCVs and the two lowest numerical values are in good agreement while the remaining four numerical values are at considerably higher concentration, possibly due to contamination during chemical analysis. The Cofino model *without* LCVs gives a consensus value of 13.9 µg kg⁻¹ while the robust mean (Huber estimator) is 18.9 µg kg⁻¹. Using the full Cofino model *with* the LCVs, the consensus value reduces to *ca* 1.5 µg kg⁻¹, which is more in line with the range of the majority of the data. We should emphasize that in this case the consensus value would most likely be indicative, due to the high uncertainty of the consensus value.

Benefits of the Model

The Cofino model has been developed and tested for use in the determination of population characteristics specifically, but not exclusively, for PTs. The current algorithm used in the NDA implementation of the model assumes an underlying normal distribution, but has a same level of robustness as the Huber and Hampel approaches. The model can be applied to a wider range of data types.

Currently, this model provides comprehensive summary statistics for each mode in a dataset combined with an extensive graphical description of the structure of the data. These outputs are obtained in a single computation without any manipulation of the data, trimming or subjective elimination, including LCVs and missing values.

The model can be applied directly to tailing and skewed data, including extreme values (outliers) and bimodal (multimodal) distributions. In the case of bimodal data, the model may identify each mode and provide an estimate of the mean, SD, and percentage of data associated with each mode (*see Multiple Modes in Proficiency Test Data*). The model can include LCVs and has been tested on datasets with up to 60% LCVs. The model works well without the need to have the *within-laboratory* SD by using the NDA developed for this purpose. However, where these values are both available and reliable then they may provide a more sensitive replacement for the NDA.

The detailed, graphical information from the model provide a plot of the PMFs and/or of the population density functions. The color density plot of the matrix overlap is very sensitive to identifying structure of data, especially modality, where it is not so clear in other plots. Bimodality can be detected with the overlap matrix plot, even when both modes have equal density (*see Multiple Modes in Proficiency Test Data*). A ranked overview shows the means, the SD of each laboratory, and the consensus value that includes both numeric and LCVs.

The model has been fully tested with randomly generated normal distributions with the number of observations (NObs) 10–200, relative SD 2–30%, to compare performance with conventional statistics. We have also tested normal distributions overlaid to construct bimodal data. The sensitivity to bimodality, with respect to the distance, and the relative amplitude of the modes has also been tested. An evaluation

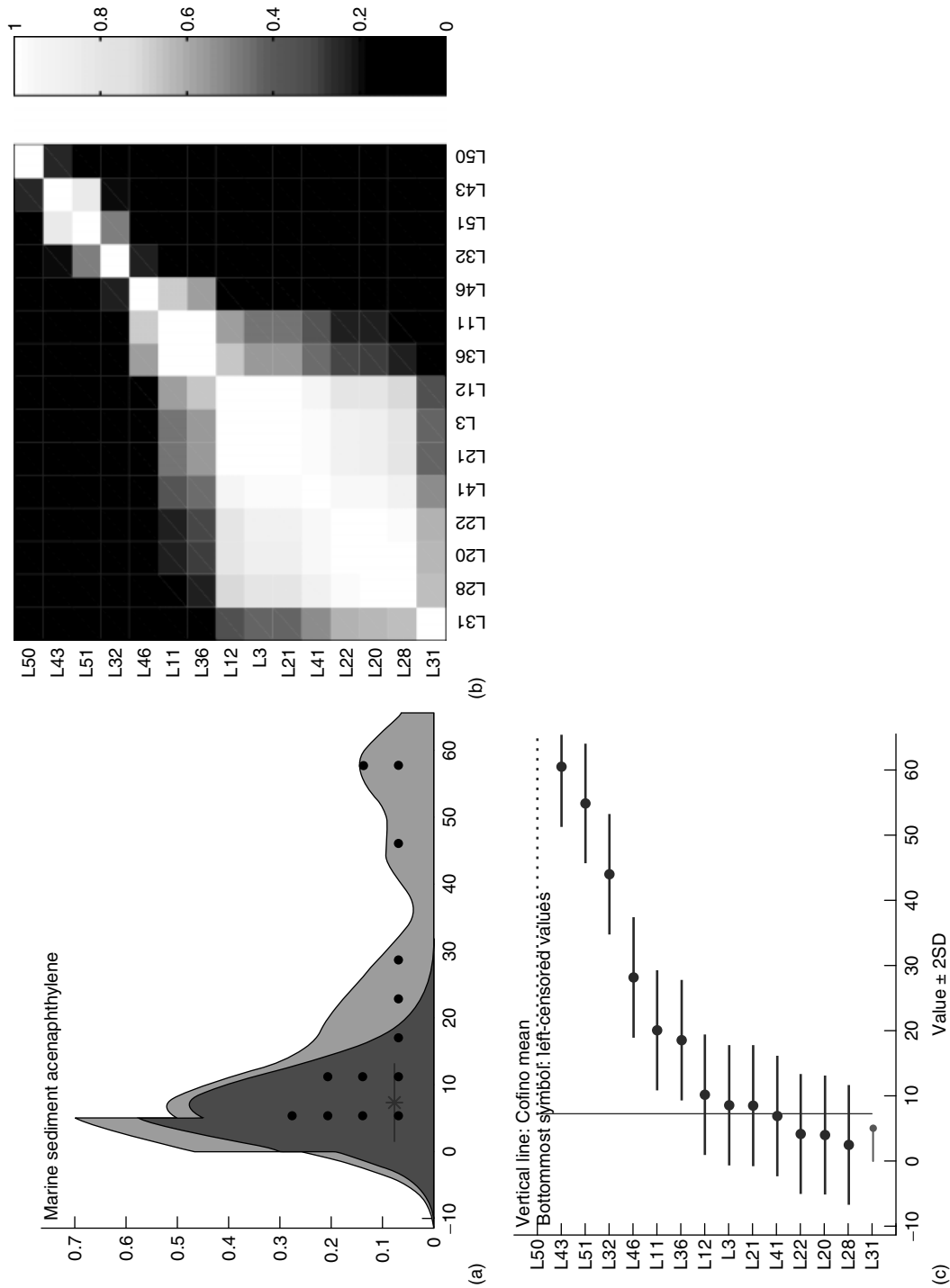


Figure 2 Proficiency test for acenaphthylene in sediment showing the Cofino model output. (a) The population measurement function for all data (light area) showing extensive tailing, and for the main mode PMF₁ (dark area). (b) Matrix overlap. Scale 0 (no overlap of data) to 1 (complete overlap). The light square shows data below L36 in agreement with the positive tailing values. (c) Ranked overview of data $\pm 2\text{SD}_{\text{NDA}}$ dark values denote LVCs

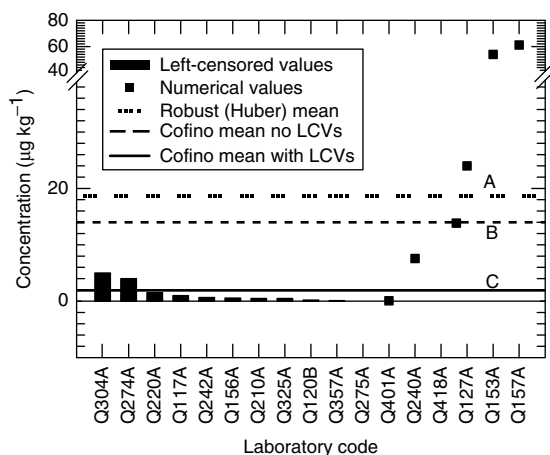


Figure 3 The ranked distribution of data for *op'*-DDT in fish muscle. Of the 17 data, 11 are LCVs. Only the output from the Cofino model with the LCVs included give an estimate that reflects the values found by most laboratories. i.e., $<2 \mu\text{g kg}^{-1}$. A: Robust (Huber) mean $18.9 \mu\text{g kg}^{-1}$ also only includes numerical data. B: Cofino mean $13.9 \mu\text{g kg}^{-1}$ without LCVs. C: Cofino mean $1.5 \mu\text{g kg}^{-1}$ with LCVs included

of over 2000 datasets taken from the QUASIMEME Laboratory Performance Studies has been made using this model.^a

End Note

^a The Cofino model has been developed using the MATLAB programming language. Copies of the code and associate information are available from the authors. A compiled version is available for use without the need to obtain the MATLAB programs. The compiled versions cannot be amended.

Acknowledgments

The authors acknowledge the members of the QUASIMEME team, in particular Judith Scurfield, in testing this model, and the participants of the QUASIMEME scheme

for providing invaluable data for evaluation. The authors also acknowledge Ian Davies for his helpful critical review of the manuscript.

References

- [1] Cofino, W.P. & Wells, D.E. (1994). Design and evaluation of the QUASIMEME inter-laboratory performance studies: a test case for robust statistics, *Marine Pollution Bulletin* **29**, 149–158.
- [2] Thompson, M., Ellison, S.L.R. & Wood, R. (2006). The international harmonized protocol for the proficiency testing of analytical chemistry laboratories, *Pure and Applied Chemistry* **78**, 145–196.
- [3] Huber, P.J. (1972). Robust statistics: a review, *Annals of Mathematical Statistics* **43**, 1041–1067.
- [4] Analytical methods committee robust statistics—how not to reject outliers. Part 1. Basic concepts, *The Analyst* **114**(12) 1693–1697 (1989).
- [5] Analytical methods committee robust statistics—how not to reject outliers. Part 2. Inter-laboratory trials, *The Analyst* **114**(12), 1699–1702 (1989).
- [6] ISO/CD 20612 (2005). *Water Quality-Interlaboratory Comparisons for Proficiency Test of Laboratories*, International Standards Organization, Geneva.
- [7] Müller, C.H. & Uhlig, S. (2001). Estimation of variance components with high breakdown point and high efficiency, *Biometrika* **88**(2), 353–366.
- [8] Cofino, W.P., Wells, D.E., Ariese, F., van Stokkum, I.H.M., Wengener, J.W. & Peerboom, R. (2000). A new model for the inference of population characteristics from experimental data using uncertainties, *Journal of Chemometrics and Intelligent Laboratory Systems* **53**, 37–55.
- [9] Cofino, W.P., van Stokkum, I.H.M., van Steenwijk, J. & Wells, D.E. (2005). A new model for the inference of population characteristics from experimental data using uncertainties. Part II. Application to censored datasets, *Analytica Chimica Acta* **533**, 31–39.

DAVID E. WELLS AND WIM P. COFINO

Multiple Modes in Proficiency Test Data

Introduction

Chromium in sediment and biota, chlorobiphenyls in mussels, and aluminum in sediment, are some of the many chemical contaminants in the marine environment. Proficiency Test (PT) schemes are used to provide the external quality assurance for chemical analytical laboratories that measure these types of contaminants [1]. Homogeneous natural test materials are distributed for chemical analysis on a regular basis and the returning data are assessed to determine the extent of the agreement between laboratories. Many scheme providers, such as QUASIMEME^a [2], allow laboratories to use their own optimized analytical methodology to obtain a single value for the measurement. A consensus value is obtained from these data that forms the basis of the assessment. Ideally, these data should be normally distributed, but frequently they are skewed, bi-, or multimodal. As a consequence the consensus value can be highly influenced by the distribution of these data. Many factors contribute to tailing or bimodality and it is important that we identify the data associated with each mode and, where possible, the root causes of these differences. In some instances, it is due to the relative **bias** of some procedure in the chemical analytical process. In such cases, we need to establish whether the bimodality is true or is an artefact.

We also require qualitative and, where possible, quantitative information on the degree of separation needed between modes before detection is possible and a simple screening test to identify bimodality. In addition, we require information on the population characteristics of each mode, preferably without separating the data and a clear, sensitive, graphical representation of the bimodality. We have described the Cofino model and its applications in **Population Characteristics of Proficiency Test Data** and we describe the specific problem of the evaluation of bimodal data here.

By using the Cofino model (*see Population Characteristics of Proficiency Test Data*) [3, 4] to evaluate the data, we can obtain the population characteristics (mean, standard deviation (SD) and the percentage of data attributed to that mean (λ))

for each of the population measurement functions (PMFs). In most cases, over 90% of the data are accounted for by the first two modes, PMF₁ and PMF₂ when the **normal distribution** approximation (NDA) is applied. So, for the bimodality to be significant, the data associated with PMF₂ needs to be distinct from the data associated with PMF₁. Since the data associated with both PMF₁ and PMF₂ have a known mean and SD from the Cofino model, we can use the Student's *t* test as a simple, preliminary, indicator of bimodality.

The identification of bimodality primarily depends on two factors. The first is the degree of separation between the two modes and the second, the percentage of data in each mode.

Resolution of the Modes

To test the resolving **power** of the model, we randomly generated a series of datasets to give $N(100, 2)$ with 50 values and $N(100 + i, 2)$ with 30 values, where $i = 1 : 10$. The SDs of the two populations s_{pop} are thus equal. Data for each degree of separation $N(100, 2) + N(100 + i, 2)$ were combined to provide bimodal distributions with a separation i between the modes. We then applied the Cofino NDA model to obtain the population characteristics of both modes.

In Figure 1 we have plotted the separation of the two modes expressed as i/s_{pop} (thus normalized with respect to s_{pop}) against (a) their means, (b) the ratio of the percentage of data associated with PMF₁ and PMF₂, and (c) Student's *t* value.

We detected very little separation up to a $3 \times s_{\text{pop}}$ difference in the mode means, as both distributions overlapped and influenced the mean of the two modes, PMF₁ and PMF₂. As we increased the separation, the Cofino model mean of the first mode (PMF₁) decreased to approximately the known value (100), while that of the second mode (PMF₂) increased to its known value. Simultaneously, the ratio of percentage of data in the two modes fell to that of the two independent datasets $50/30 = 1.7$.

Percentage of Data in Each Mode

The second factor affecting the identification of bimodality is the number of observations (NObs) associated with each mode.

2 Multiple Modes in Proficiency Test Data

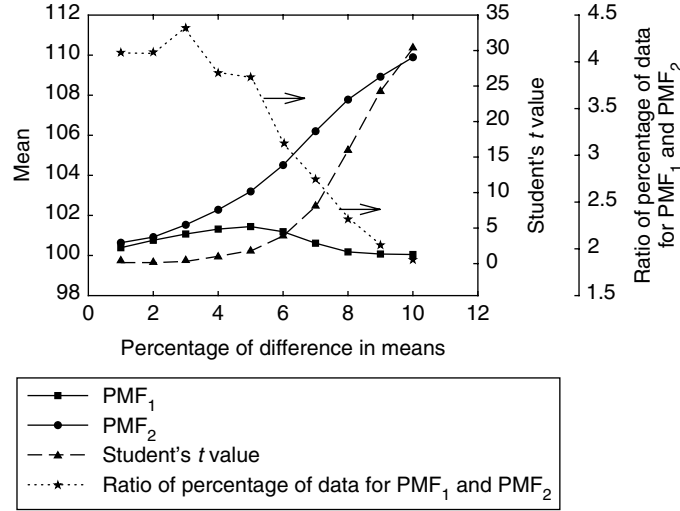


Figure 1 The effect of separation of two overlapping distributions $N(100, 2)$ and $N(100 + i, 2)$ where $i = 1 : 10$ on the mean of data associated with PMF_1 (first mode) and PMF_2 (second mode) obtained using the Cofino model. Both modes are clearly resolved when the difference between the means $> 6\%$

In the example we assigned a ratio of 50:30 for the data in the two modes. However, if we change this ratio to 50:50 it becomes more difficult to identify and isolate the two modes due to the way in which the SD, using the NDA model, is constructed.

Again we generated and combined data with $N(100, 2)$ and with NObs from 30 to 50 in increments of 5 and a second distribution with $N(107, 2)$ and $NObs = 30$. The values for the mean of data associated with PMF_1 , and PMF_2 are plotted against the ratio of NObs in the two modes. Superimposed on this plot is the value of the Student's t value for each of the sets of data (Figure 2).

When the ratio of NObs for the two modes is 1 there is no separation. Here the values for the Cofino mean for PMF_1 and PMF_2 are 103.4 and 103.6, respectively. As we increase the ratio of NObs between the two datasets, we begin to detect separation. When the ratio of NObs exceeds 1.7 the Cofino means of the two modes are close to their independent values, the Student's t value is > 2 , i.e. t_{crit} , and there is a clear separation. Both means of the data associated with PMF_1 and PMF_2 are now close to their independent values 100.2(100) and 106.6(107) respectively.

The breakpoint for separation and obtaining the population characteristics of both modes depends both on the ratio of NObs in the two modes and on

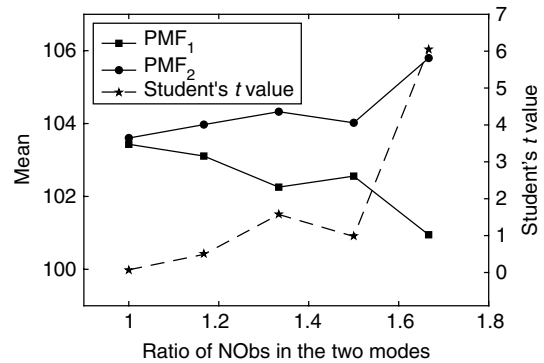


Figure 2 The effect of the ratio of NObs in the two modes $N(100, 2)$ and $N(107, 2)$ on the mean of data associated with PMF_1 and PMF_2 . Both modes are resolved when the ratio of NObs in the two modes $> 1:1.66$

the difference between their means and their SDs. In our example, using one and the same s_{pop} of about 2% of population mean, the ratio of NObs had to be greater than 1.7 and the difference between the means greater than $3 \times s_{pop}$. Below these breakpoints it is still possible to detect the bimodality using the matrix overlap, even though the population characteristics of each mode cannot be obtained from the whole dataset (see **Population Characteristics of Proficiency Test Data**) (Figure 3b). In such a circumstance, the data

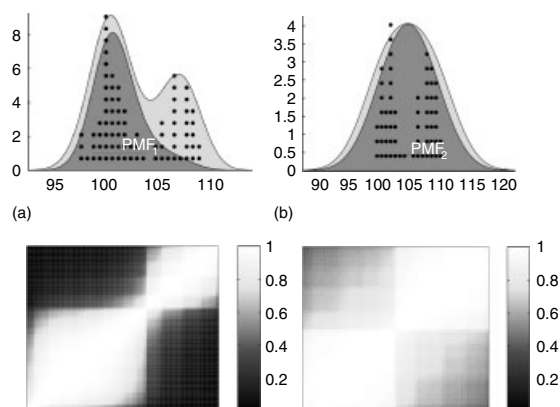


Figure 3 The PMF (upper) for all data (light shade area) and for PMF_1 (dark shades area) and the overlap matrix (lower) for the combined distributions $N(100, 2)$ and $N(107, 2)$ (a) with NObs 50:30 and (b) with NObs 50:50 in each mode. A distinct separation occurs when the means differ by $>6\%$ and the ratio of NObs is $>1:1.66$. Even when there is no evidence of separation of the modes, as in PMF profile (b – upper), the two constituent modes are distinctly evident in the overlap matrix (b – lower)

in the two modes can be identified and each subset of data treated separately. For many other datasets the assessment can be made in a single calculation. Also by inspection of the distribution of the data in the matrix-overlap plots, it is possible to identify any inhomogeneity that might otherwise go undetected.

Evaluation of the Causes of Bimodal Data in PTs

In any PT where the contents of the matrix are not characterized, we need to estimate the population characteristics of each main mode in the data and provide a clear explanation of any bimodal distribution. The Cofino model often allows us to achieve this without separating the data, while the normal arithmetic or robust mean using the Hu [5–7] or Ha [8, 9] estimators will only provide a mean value that may lie somewhere between the two modes.

Chromium in Sediment and Biota

Different laboratories extract the chromium from the sediment using either a nondestructive method or hydrofluoric acid (HF) digest (*total*), or by *partial*

digestion methods with aqua regia. For one PT, the robust mean of all data, using the Hu estimator, was 85 mg kg^{-1} while the Cofino model identifies two modes at 89 and 67 mg kg^{-1} which could be associated with *total* and *partial* digestion methods, respectively. When we subsequently separated the data into two sets according to the digestion method and reapplied the robust statistics, we obtained means of 91.5 and 69.1 mg kg^{-1} .

The level of chromium in biological tissue is generally low and contamination can easily occur during chemical analysis unless an ultraclean environment is used. Data from a PT for chromium in cod liver were bimodal, with a Cofino mean for the data associated with the main mode (PMF_1) around $110 \mu\text{g kg}^{-1}$ (Figure 4). The digestion media included hydrochloric acid, aqua regia, nitric acid, nitric acid with hydrogen peroxide, and five others unspecified. There was no clear preference for either group but most of the laboratories used nitric, with or without hydrogen peroxide. Although there was a range of different digestive media, six laboratories used an open heating (OH) system for digestion and five of these were within the upper 50% of the data, with three being in the second mode. Microwave (M) digestion and the pressure bomb (PB) formed most of the data in the first mode (Figure 4).

Although there may be other differences between methods, e.g., detection systems, there is a clear bias caused by contamination of the samples digested in open vessels.

A comparison of the data obtained by the Cofino model, the Hu and Ha estimators are given in Table 1. The mean values are also indicated in Figure 4. From the summary statistics in Table 1 and from the matrix-overlap plot in Figure 4 we see that the Cofino model provides a consensus value more closely aligned with the main mode of the data while the robust statistics give a mean between the two modes.

Chlorobiphenyl (CB105) in Mussels

CB105 is amongst the more difficult congeners to completely separate chromatographically from other CBs, in particular, from CB132. Therefore, PT data can be a result of measuring just CB105 or a mixture of CB105 and other, unseparated congeners. Higher values can also be caused by contamination due to insufficient sample cleanup. The matrix-overlap

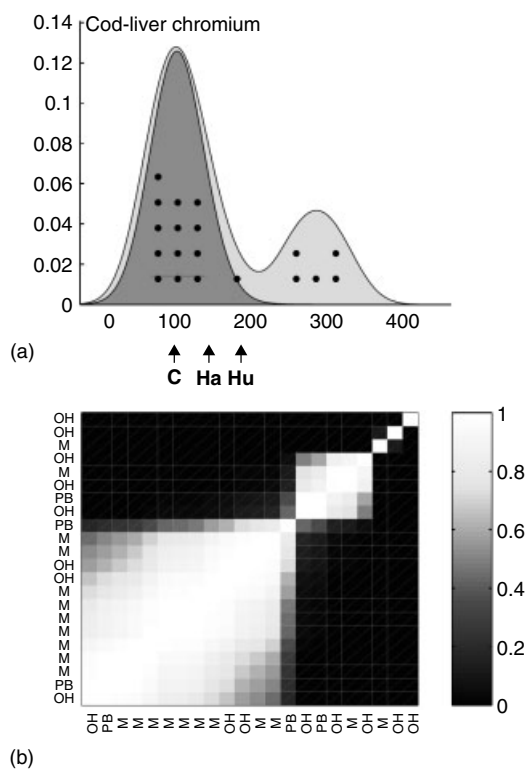


Figure 4 The PMF (a) and the overlap matrix (b) for chromium in cod liver PT. The distribution of returning data is bimodal. The methods of sample preparation used to obtain the data in the main, mode (a) comprise microwave (M) and pressure bomb (PB), both sealed systems, and a few open vessel heating (OH). In the smaller mode at higher concentration, the positively tailing data included five out of eight laboratories that used the OH system, which suggests that this maybe a source of local contamination. C: Cofino mean, Ha: Hampel estimator, Hu: Huber estimator. Our evaluation concluded that the Cofino model provided the best estimates of means for *both* modes

plot (Figure 5a) shows two modes indicated by the lighter colored areas and two very high values (top right) The mode at the lower concentration is more homogeneous with a greater overlap of data. The second mode, at the higher concentration is more diffuse. These higher values in the second mode are most likely to be associated with an insufficiently cleaned up sample, degrading or inappropriate chromatographic column selection, and/or optimization.

The mean of data associated with the first mode at a lower concentration is $0.27 \mu\text{g kg}^{-1}$ while the second mode is about 2.5 times higher at $0.64 \mu\text{g kg}^{-1}$.

Table 1 A comparison of the summary statistics from the proficiency test for chromium and chlorobiphenyl, CB105, in mussel and aluminum in sediment. Each of these PT data shows a bimodal distribution due to different methods of analysis (see Figures 4–6). Both robust methods using the Huber (Hu) and Hampel (Ha) to estimate the mean for chromium are affected by the second mode at higher concentrations for chromium and CB105 and by the second mode at lower concentrations for aluminum. In each case, the Cofino mean estimates the mean where there is the greatest overlap of data. The Cofino model can provide the mean and SD for both modes from all data in one computation

	Determinand	Total NObs	LCV NObs	Percentage of data in first mode		Units	Cofino		Huber		Hampel		Arithmetic	
				NObs	LCV		mean	SD	mean	SD	mean	SD	mean	SD
Biota	Chromium	25	3	70	$\mu\text{g kg}^{-1}$	125.21	80.22	187.39	130.34	150.61	83.50	683.82	2139.42	
	CB105	22	0	75	$\mu\text{g kg}^{-1}$	0.31	0.16	0.42	0.24	0.39	0.18	1.08	3.12	
	Mussel	27	0	75	%	4.62	1.46	4.09	2.04	4.18	1.41	4.10	1.84	
Sediment	Aluminum													

In most cases where there is clear information on likely additive interferences, the best estimate for the consensus value is often the mode with the lower value. However, this does not always need to be the case and careful judgment with supporting chemistry is essential. Again, mean using the Hu estimator is *ca* 30% higher than that when using the Cofino model (Table 1).

Aluminum in Sediment

The selection of digestion methods affects the recovery of aluminum from minerals such as feldspar in the

sediments. A range of strong acids are normally used, with HF being amongst the most effective at breaking the crystal lattice and releasing the aluminum. Other digestive techniques used are aqua regia, nitric acid as well as nondestructive techniques such as X-ray fluorescence. Most laboratories currently use HF as the preferred digestive “wet” technique.

The results for aluminum in the marine sediment show three obvious modes in the data that are clearly visible in the overlap matrix plot (Figure 6a). The highest mode, around 7.5% Al, is probably a result of two laboratories obtaining a particularly high

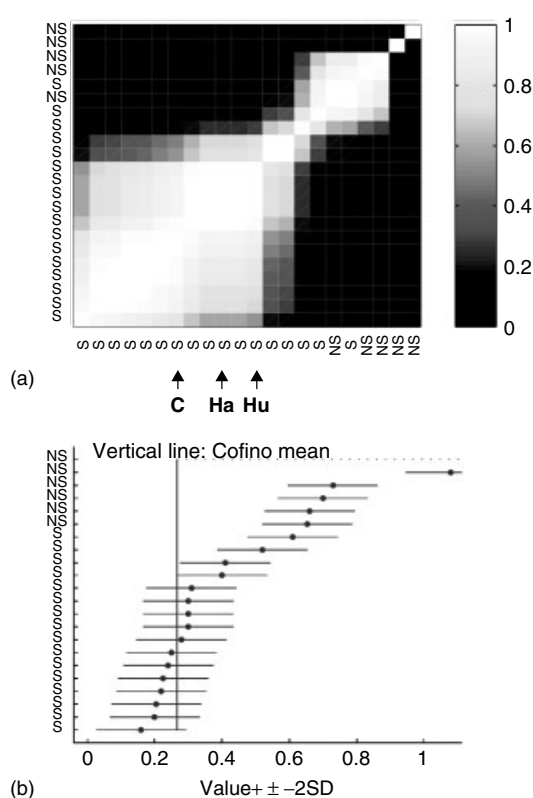


Figure 5 The overlap matrix (a) and the ranked distribution (b) for chlorobiphenyl, CB105 in mussels. The highly tailing data creates a bimodal distribution caused by poor chromatographic resolution, insufficient sample cleanup or contamination (NS) during the chemical analysis which leads to elevated results for these laboratories. The data in the main mode of the distribution (S) have been generated under more optimum conditions. C: Cofino mean, Ha: Hampel estimator, Hu: Huber estimator. The Cofino model provides the best estimates of means for *both* modes

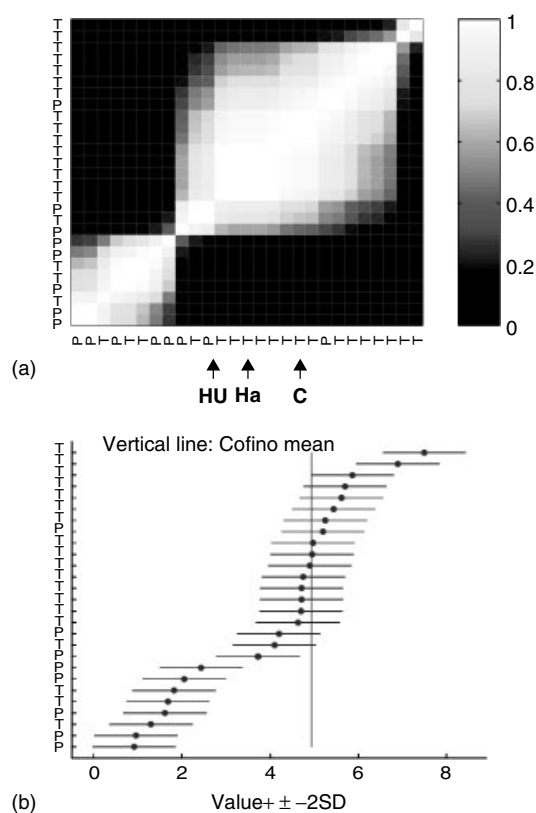


Figure 6 The overlap matrix (a) and the ranked distribution (b) for aluminum in sediment. The bimodal distribution results from different methods of sediment digest. The partial digest using acids other than hydrofluoric acid (HF) form most of the lower mode while the main mode at higher concentrations consists of total digest methods (HF) and nondestructive techniques such as X-ray fluorescence. C: Cofino mean, Ha: Hampel estimator, Hu: Huber estimator. The Cofino model provides the best estimates of means for *both* modes

value. The main mode has a concentration of 4.95% and consists almost exclusively of digestion methods involving HF or nondestructive, *total* techniques. The lower mode at 1.6% Al includes methods using aqua regia (two), nitric acid (two), nitric acid and peroxide (*partial*), as well as three laboratories using HF. The non-HF, wet techniques clearly do not extract all of the aluminum, but the use of HF *per se* does not always guarantee complete recovery of the element. The robust estimates are again affected by the second mode at lower concentration, due mainly to the data obtained by the partial digest of the sediment. In this case, the most homogeneous and reliable consensus value was based on the main mode at the higher concentration.

Benefits of the Cofino Model

Using the Cofino model we are able to estimate the population characteristics of each mode in a bi- or multimodal distribution. In most cases, we can achieve this using the whole dataset, including left censored values (LCVs), without the need to trim the data or remove outliers. The resolution between the modes depends on the ratio of data in the two modes and on the difference between the means in relation to the SDs of the modes. However, even in unfavourable conditions, the matrix overlap plot will still suggest the presence of two modes.

Evaluation using the Hu or Ha robust estimator will provide similar consensus values where the tailing and **outlier** values are *ca* < 10% of all data, but unlike the Cofino model, these techniques are less able to cope effectively with heavily tailing or bimodal data. In some cases, as we have shown here, the robust mean using these estimators can occur between the modes where there are very few data. The combination of the information gained from the Cofino model with scientific expertise will, in many cases as shown above, enhance the reliability of interpretation of PT data.

End Note

^a. QUASIMEME, Quality Assurance of Information for Marine Environmental Monitoring (in Europe),

funded initially by the EU (1992–1996) for European Laboratories, is now a worldwide project funded by the subscribing participating laboratories www.quasimeme.org.

Acknowledgments

The authors wish to acknowledge the members of the QUASIMEME team, in particular Judith Scurfield, in testing this model, and the participants of the QUASIMEME scheme for providing invaluable data for evaluation. The authors would also like to acknowledge Ian Davies for his helpful critical review of the manuscript.

References

- [1] Thompson, M., Ellison, S.L.R. & Wood, R. (2006). The international harmonized protocol for the proficiency testing of analytical chemistry laboratories, *Pure and Applied Chemistry* **78**, 145–196.
- [2] Cofino, W.P. & Wells, D.E. (1994). Design and evaluation of the QUASIMEME inter-laboratory performance studies: a test case for robust statistics, *Marine Pollution Bulletin* **29**, 149–158.
- [3] Cofino, W.P., Wells, D.E., Ariese, F., van Stokkum, I.H.M., Wengener, J.W. & Peerboom, R. (2000). A new model for the inference of population characteristics from experimental data using uncertainties, *Journal of Chemometrics and Intelligent Laboratory Systems* **53**, 37–55.
- [4] Cofino, W.P., van Stokkum, I.H.M., van Steenwijk, J. & Wells, D.E. (2005). A new model for the inference of population characteristics from experimental data using uncertainties. Part II. Application to censored datasets, *Analytica Chimica Acta* **533**, 31–39.
- [5] Huber, P.J. (1972). Robust statistics: a review, *Annals of Mathematical Statistics* **43**, 1041–1067.
- [6] Analytical methods committee robust statistics—how not to reject outliers. Part 1. Basic concepts, (1989). *The Analyst* **114**(12), 1693–1697.
- [7] Analytical methods committee robust statistics—how not to reject outliers. Part 2. Inter-laboratory trials, (1989). *The Analyst* **114**(12), 1699–1702.
- [8] ISO/CD 20612 (2005). *Water Quality-Interlaboratory Comparisons for Proficiency Test of Laboratories*, International Standards Organization, Geneva.
- [9] Müller, C.H. & Uhlig, S. (2001). Estimation of variance components with high breakdown point and high efficiency, *Biometrika* **88**(2), 353–366.

DAVID E. WELLS AND WIM P. COFINO

Disease and Clinical Trial Modeling

Introduction

Clinical trials are considered the gold standard among the experiments aimed at providing scientific evidence on the effects of treatments and other factors that determine the progression of diseases. However, clinical trials are often long and expensive, and their results are uncertain. Moreover, they pose individuals at risk. It is therefore important to ensure that their objectives will be met at the lowest possible cost in terms of health and resources. Poorly designed clinical trials might fail in assessing the hypothesis tested with the required level of statistical significance, which means that resources will have been wasted and patients will have been unnecessarily put at risk. Modeling techniques allow decision makers to explore the consequences of their decisions in planning processes and activities and to identify the best courses of action under uncertainty [1, 2].

Modeling has been used for a long time in the field of clinical research as an aid for ensuring the attainment of positive results of experiments at the stage of experimental design. But disease and clinical trial modeling may have broader goals and take place at different phases of the trial. A model might sometimes be specified in a way that allows solving it by analytical methods, often providing a single solution, but this is often not the case. Simulation might then provide a better methodological approach, especially in dealing with complexity and uncertainty, as it is usually the case in health research. The models most widely used for assessing the effects of pharmaceuticals and other healthcare technologies are discrete health states and stochastic models, where health is represented by variables that change at discrete periods of time [3]. Although simulation has sometimes been used to extrapolate trial results when cost or time restrictions did not permit performance of a new clinical trial, simulation has mainly been used to explore in advance the likely outcomes of future or hypothetical trials; for instance, simulation was employed to illustrate the possible reasons why a randomized trial of an acute stroke treatment might be underpowered [4]. When applied to the design and management of clinical trials, modeling can increase

the likelihood of success and reduce the costs of failures and the inefficient performance of clinical trials.

From an economic perspective, the global objective of a clinical trial can be formulated in terms of a function whose arguments are statistical **power**, total duration, and cost of the clinical trial, because these three variables determine the net profit of putting a pharmaceutical product or health technology in the market from the point of view of a private investor. We assume that the investor tries to maximize the net profit of a clinical trial by maximizing the present value of the stream of costs and revenues related to the product. We further assume that for a limited period the product will benefit from market exclusivity that allows it to be sold at a price that yields extraordinary profits. A shorter duration of the trial would lead to a longer period of high monopoly price of the drug. A higher statistical power, in its turn, will lead to a higher probability of obtaining market approval, which is a condition for obtaining revenues and, hence, profits from the sales of the product under evaluation.

If we define *ENB* as the total expected net benefits, *NBM* as the net average extraordinary profit from annual sales, *n* as the duration of the product life cycle, *P_e* as the success probability of marketing the product under trial, *C_t* as the total cost of clinical trial, and *d* as the total duration of the clinical trial, *ENB* can be expressed as

$$ENB = P_e \cdot NBM \cdot (n - d) - C_t \quad (1)$$

The objective of this paper is to show how simulation modeling can be used to estimate the *P_e*, *d*, and *C_t* that result from a certain trial design or from similar decisions. Together with some estimates of *NBM* and *n* from other sources, decision makers might be able to assess and maximize the expected net profits of a clinical trial as represented by the expression (1). In the third section of the article, the previous analysis is illustrated with the results of a case study with both actual and hypothetical data.

Methods

The general model for optimization of a clinical trial that we propose in this work consists of three basic elements: structure, parameters, and assumptions. The structure is based on two submodels. The first one

2 Disease and Clinical Trial Modeling

is the recruitment submodel, which represents the arrival of volunteers at the center and their posterior inclusion in the trial. Its objectives are to estimate the time and cost required for completing the process of recruitment of patients in the trial. This process will depend on several factors, such as, the number of patients required, the number and capacity of the centers involved in the trial, the arrivals rates of volunteers to the centers at each time period, and the probability of inclusion of the patients in the trial.

The second one is the follow-up submodel, which describes the course of the disease under the various options evaluated. The disease progression is represented by a statistical model considering either continuous and/or discrete health states to evaluate changes in response variables. For example, the evolution of the cholesterol level over time can be modeled as a continuous variable that changes at discrete periods of time or as discrete health states, which are then defined in terms of the cholesterol level. In the second case, the progression of patients between health states depends on certain predefined transition probabilities. These probabilities depend on time and patient characteristics, such as age, morbidity, and on initial and intermediate health states.

Parameters are classified in two categories, the first one includes statistical significance test, endpoints, number of centers, and duration of follow-up, which are considered as trial design variables, and whose characteristics and values vary within certain constraints. Some of these constraints might be imposed by the clinical trial regulation agency or by traditional practice. The second category represents the external variables that may affect the final results of the trial. These are recruitment rate, center capacity, inclusion probability, inclusion cost, cost per center, treatment effects, probability of withdrawal, and non-compliance. Other external variables such as resource use and unit costs of resources are integrated into the model in order to assess the economic outcomes, e.g., the total cost of the trial. Variables might be integrated into the model as fixed parameters, as endogenous variables with either parametric or empiric distribution, or as restrictions.

Figure 1 illustrates construction of a disease and clinical trial model, based on assessing the likely effects of alternative trial designs, which in economic terms means choosing combination of inputs and transforming them into outputs. The analyst can predict which design maximizes the target function.

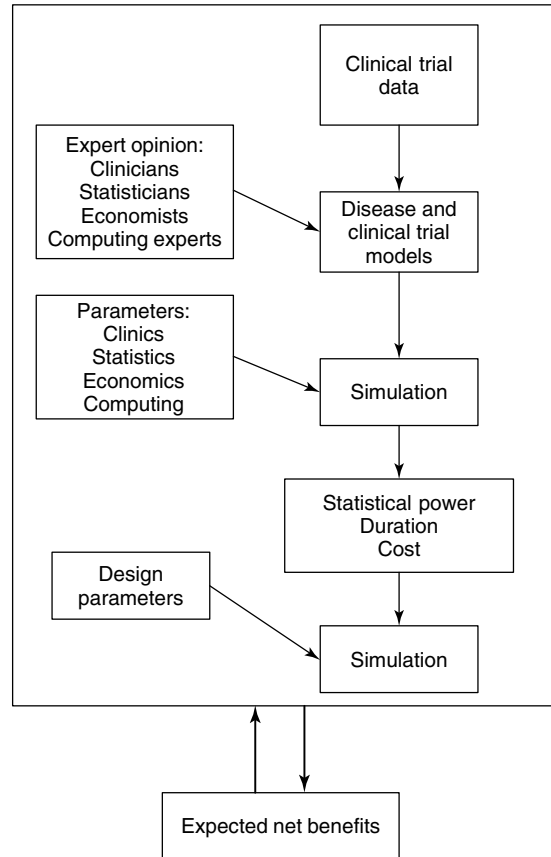


Figure 1 Disease and clinical trial model construction

Data is used to construct the model and adjustments to the structure of the model, later, can be incorporated based on different expert opinions: clinicians, statisticians, economists, and informatics. The objective is to adapt the model to parameters that come from different sources. Simulation is performed in order to calculate the P_e , d , and C_t . Design parameters were integrated assuming a range of values for each of them and simulation is carried out again in order to calculate the expected net benefit.

The optimization of a clinical trial design can be performed by selecting different feasible combinations of the input parameters and executing the clinical trial model. This can be carried out with the help of an automatic strategy that explores the outcomes of different combinations of design variables; finally, the results are evaluated using an optimization

target function such as equation (1). By carrying out simulations with the appropriate number of repetitions of each possible combination, the designer of a clinical trial can explore the value of the optimization function within the assumed restrictions. In order to assess the robustness of the results, sensitivity analysis can easily be carried out using a range of possible parameter values and scenarios.

Results

We applied the approach explained above to a case study of real clinical trial conducted and described by Negredo *et al.* [5] from the Hospital Universitari Germans Trias i Pujol. The main objectives of this trial were to compare the efficacy of three regimes, one protease inhibitor (PI)-containing regime and two PI-sparing regimes, including nevirapine or efavirenz, to maintain the suppression of plasma HIV-1 RNA and to allow the progressive immunologic improvement of patients. Secondary objectives were to assess the impact of PI-sparing regimes on the metabolic profile, on other adverse events and on quality of life. The trial was also aimed at evaluating the variations in body-shape occurring in patients suffering lipodystrophy at baseline. The 100 patients who had been receiving a PI-containing regime and who had long-lasting plasma HIV-1 RNA were recruited after 12 months and were randomized to either continue with the same therapy or to have the PI replaced by nevirapine or efavirenz. These patients were followed for up to 1 year. Only 49 patients completed the trial (23 in the PI-sparing arm and 26 in the PI-based antiretroviral therapy) and only the data of these patients were used to illustrate the applicability of our simulation model, although it can be easily extended to account for the influence of this missing data on the simulation results in a sensitivity analysis.

The structure of the model was adjusted according to the opinions of clinical, statistic, and economic experts. Design parameters including, different methods to quantify the endpoint, number of centers, were identified. Then, unit cost parameters were incorporated in order to estimate the total cost of the trial as function of statistical power, and the duration considering different sample sizes, number of centers, and center capacities. The unit costs were €6500 per center, €60 per patient recruited, and €60 for

patient follow-up, and the monthly cost of delaying the finalization of the trial could be estimated as €6000. The design that minimizes both duration and cost was achieved with eight centers and a duration of 10 months, leading to the evaluation of the effect of the treatment on 300 patients. Thus, the estimated duration and cost were 7.70 months and €120 200, respectively.

We finally substituted the resulting power, duration, and cost by the P_e , d , and C_t in formula (1) to assess the expected net profits. Assuming the remaining time until the expiration of the product to be 10 years, $ENB = 0.83 \times 6000 \times (10 \text{ years} \times 12 \text{ months} - 7.70 \text{ months}) - 120\,200 = €547\,234$.

Conclusions

The approach proposed in this work allows substantial improvements in clinical trial design and management. Moreover, simulation models integrate clinical research with statistics, economics, and computer disciplines, in order to optimize the global results of clinical trials in terms of economic benefits. The same approach can be used to assess the social value of a clinical trial, by reformulating equation (1) in terms of socially relevant effects, i.e. health and resource effects.

For simplification purpose, deterministic calculation was performed to determine the expected net benefits, where uncertainty has not been considered. Nevertheless, when we consider uncertainty, the simulation can be easily used to perform a large number of replications with all possible input combinations. Hence, the expected value is calculated for each combination, and the highest value is chosen to decide the **optimal design** of the trial. The model can be conveniently adjusted to specific conditions in order to represent a planned clinical trial and to prospectively simulate the results for alternative designs and assumptions.

The approach is being applied to phase III clinical trials, although it can also be applied to other phases. Finally, there is the possibility of applying it in real-time practice, such as in the follow-up period of a clinical trial, in order to inform stop-go decisions, possible increases in the numbers of patients or centers, and so on. The model might also be improved by progressively incorporating on-going information and results of the trial.

References

- [1] Jonsson, S.C. (1998). The role of simulation in the management of research: What can the pharmaceutical industry learn from aerospace industry? *Drug Information Journal* **32**, 961–969.
- [2] Backhouse, M.E. (1998). An investment appraisal approach to clinical trial design, *Health Economics* **7**, 605–619.
- [3] Weinstein, M.C., O'Brien, B., Hornberger, J., Jackson, J., Johannesson, M., McCabe, C. & Luce, B.R. (2003). Principles of good practice for decision analytic modelling in health-care evaluation: report of the ISPOR task force on good research practices-modelling studies, *Value in Health* **6**(1), 9–17.
- [4] Samsa, G.P. & Matchar, D.B. (2001). Have randomized controlled trials of neuroprotective drugs been underpowered? An illustration of three statistical principles, *Stroke* **32**(3), 669–674.
- [5] Negredo, E., Cruz, L., Paredes, R., Ruiz, L., Fumaz, C.R., Bonjoch, A., Gel, S., Tuldrà, A., Balaqué, M., Johnston, S., Arnó, A., Jou, A., Tural, C., Sirera, G., Romeu, J. & Clotet, B. (2002). Virological, immunological, and clinical impact of switching from protease inhibitors to nevirapine or to efavirenz in subjects with human immunodeficiency virus infection and long-lasting viral suppression, *Clinical Infectious Diseases* **34**, 504–510.

ISMAIL ABBAS, JOAN ROVIRA, ERIK COBO,
JOAN ROMEU AND JOSEP CASANOVAS

Statistical Process Control in Clinical Medicine

Introduction

The quality of a clinical process may be assessed provided it is in a state of statistical control, i.e., predictable. **Control chart** techniques originally developed for industrial **process control** [1] may be applied to develop and monitor the quality of a clinical process [2]. However, if the quality measure/indicator used is an outcome measure, e.g., mortality following a coronary artery bypass graft operation, its value depends not only on the quality of the clinical process, but also on risk factors, e.g., severity of the patient's disease, comorbidity, age, sex, etc. Various risk-adjusted control charts may then be used [3], although the results obtained should be considered provisional, requiring further investigation.

Outcome results of several healthcare providers (physicians, hospital departments, etc.) may be compared if the patients are assigned at random to the providers. Usually (but not always [4]) this is not possible for practical and/or economical reasons. Then, one has to use observational data. The influence of patient case-mix differences between the providers, therefore, has to be adjusted for prior to the comparison [5]. This is usually (but not always [6, 7]) achieved by regressing the provider category and important risk factors on the outcome. The approach is fraught with problems. In the following we will present the statistical techniques used in this paper, demonstrate the construction and subsequent use of a control chart by way of a straightforward example, and illustrate and discuss the perils of risk adjustment.

Methods

Control Chart

A control chart is used to characterize a clinical process and subsequently to monitor it. The chart depicts the values of a measurable quantity characterizing the process, e.g., the fraction of patients dying per month during a specified surgical procedure, the weekly mean of patients' waiting times until seen by

a physician, the cost per patient treated for a hip fracture, etc. If a clinical process is in statistical control, the causes of variation cannot be determined, but the results follow a probability distribution that characterizes the process. Its mean (μ) and standard deviation (σ) define the state of the process. A control chart is constructed, using 20–30 samples obtained while the process was in statistical control. It is furnished with a centerline (the grand mean of the samples), an upper and a lower control limit (the grand mean plus, respectively minus three times the *sample* standard deviation), and a time axis. The control limits are constructed so that as a rule the sample means will be located within the control limits when the process is in statistical control. If a sample mean falls outside the control limits, the added variation has an **assignable cause** with a high probability, and a search for this is initiated.

Measuring the Quality of a Process

Assume that the results of a clinical process in statistical control follow a Gaussian distribution with mean μ and standard deviation σ , which are known. Let the upper specification limit (USL) be a value separating acceptable results from unacceptable results that are larger than this value. Then, $C_{pu} = (USL - \mu)/(3\sigma)$ is a measure of the quality of the process. If C_{pu} is 1.00, 0.135% of the values of the process are expected to be $\geq USL$, and if it is > 1.00 , less than 0.135% are, provided the true μ and σ are known. According to industrial standards $C_{pu} \geq 1.66$ is an acceptable quality. If the parameters are not known the **sample size** should be 100 or more [8].

Logistic Regression

Let p_{ij} be the probability that patient i , treated by healthcare provider j , dies; $\text{logit}(p_{ij})$ is the logarithm of the corresponding odds ($p_{ij}/(1 - p_{ij})$), X_{vij} is the v th quantity that is measured in this patient and is related to the probability, and $H_j (j = 1, \dots, k - 1)$ is an indicator variable that is equal to 1 if the patient is treated by healthcare provider j and 0 otherwise (provider k of the k providers serves as a reference, and $k > 1$). A logistic regression model is defined by $\text{logit}(p_{ij}) = \alpha + \sum \delta_j H_j + \sum \beta_v X_{vij} + \varepsilon_{ij}$ where α is a constant, β_i and δ_j are regression coefficients and ε_{ij} a random error. If β_v is significantly different from 0, risk factor v is significantly related to $\text{logit}(p_{ij})$

2 Statistical Process Control in Clinical Medicine

and thereby to p_{ij} . If an overall test shows that the intercepts of the providers differ significantly, the significance of the individual δ_j 's may be assessed in the same way as the β_v 's. This fixed-effect model is adequate for the example discussed in this paper because there are only two providers and the values are statistically independent. When this is not the case, a hierarchical model is preferable [9, 10].

Simulation

Assume, e.g., that we have two types of patients. The probability that a type-1 patient dies during a specified and properly administered surgical procedure is 10%, and the probability that a type-2 patient dies is 30%. If these probabilities are constant over time, the quality of the surgical procedure is stable. The behavior of such a process over time may be simulated. Assume we wanted to simulate the outcome of 1000 consecutive type-1 patients subjected to the surgical procedure. Performing 1000 independent experiments each simulating the outcome of one consecutive patient does this. Each experiment is a random draw of one ball from an urn, containing 10% red balls and 90% white balls. Thus, the probability of drawing a red ball is 10%, simulating the 10% probability of death during operation. If we wanted to simulate the outcome when a mixture of type-1 and type-2 patients were operated, we would use two urns, one for type-1 patients containing 10% red balls and another for type-2 patients containing 30% red balls. Instead of drawing balls from urns the experiments may be done using standard statistical packages.

Example of the Construction and Use of a Control Chart

Problem

At an outpatient clinic the management was notified that the USL for a patient's waiting time from arrival until blood sampling and/or electrocardiogram (ECG) recording begins is 30 min.

Design of Experiment

The management decided to take random samples comprising 30 outpatients on each weekday for

four weeks, to study the patient waiting times. The employees at the outpatient clinic were not aware of this investigation. The waiting time from arrival until seen by a technologist was recorded for each patient.

Analysis of Data

The mean and standard deviation of each sample was calculated. The grand mean ($\hat{\mu} = 15.93$ min) estimates the process mean. An unbiased estimate of the process standard deviation ($\hat{\sigma}$) is calculated by dividing the average of the sample standard deviations by a factor removing the **bias** [1]. Dividing this estimate by $\sqrt{30}$ (here the sample size is 30) one obtains an estimate of the standard deviation of the sample mean.

Figure 1 shows the $\bar{X}$ chart with the control limits equal to the mean ± 3 sample standard deviations. The observed sample means are depicted on the chart. Three of the values (samples 5, 10, and 15 all sampled on a Friday) are located above the upper control limit (UCL). Therefore, in all likelihood the process is not in statistical control. When these values were removed and a new chart calculated using the remaining values (chart not shown), the value corresponding to the fourth Friday (sample 20) lay above the UCL.

It turned out that on Fridays the patient mix differed from the other weekdays in that an unusual large number of patients from the cardiology department were scheduled for ECG recordings in addition to blood specimen collection. These patients required more time than patients not scheduled for ECG recordings, and a retrospective analysis revealed

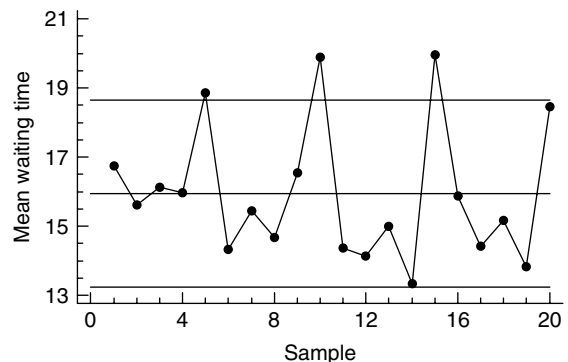


Figure 1 Mean waiting time of patients received at an outpatient clinic on weekdays

Table 1 Procedure for investigating outlying values of daily mean waiting times (start at step 0 and proceed from left to right)

Step	Actions	Questions and routing in table	
0	Control data and data processing	Data error?	If yes go to step 4 If no go to step 1
1	Compute production in number of venipuncture equivalents. Control staffing of ambulatory	Production or staffing changed?	If yes go to step 4 If no go to step 2
2	Define productivity as venipunctures/technologist hour. Compute (a) average productivity (b) average productivity for each 30-min period (c) average productivity for each technologist (d) average productivity for each technologist in each 30-min period. Identify significantly outlying values	Go to step 3	
3	Interview manager of ambulatory and technologists	Go to step 4	
4	Write report and stop		

that overtime was much more usual on Fridays than on any other weekday.

Figure 2 shows the control chart based on the remaining values. The points marked as crosses correspond to the Fridays. Now, the process seems to be in control. The management reorganized the technologist shifts and staffing. New data were collected for a 3-week period and a new control chart made using these data. This chart was similar to that shown in Figure 2. The management purchased a system for automatic recording of waiting times. The ID of a technologist servicing a given patient was already automatically captured by the current clinical data-processing system. They designed a protocol (see

Table 1) to be followed if values outside control limits were found. The estimated process mean and standard deviation were 15.72 and 4.85 min respectively. The C_{pu} calculated to assess the quality of the process was $(30.00 - 15.72)/(3 \times 4.85) = 0.98$, based on a sample size of $15 \times 30 = 630$. This is a questionable quality according to industrial standards. The management, therefore, (a) reorganized and upgraded the staffing, (b) collected new data, and (c) calculated C_{pu} to see if the state of the process had improved sufficiently. They repeated this cycle until the quality was acceptable. Using the resulting control chart, they subsequently monitored the quality of the process by plotting the mean values on the chart and used the procedure in Table 1 to track assignable causes of excess variation.

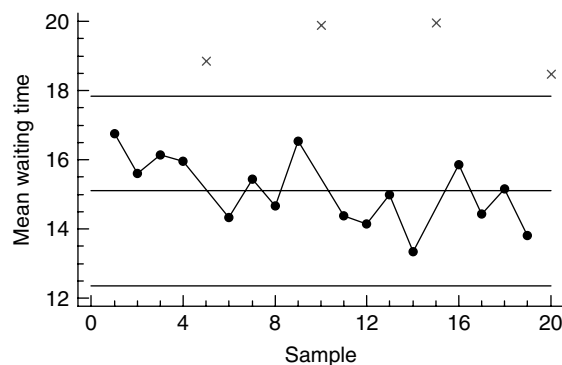


Figure 2 Mean waiting time of patients received at an outpatient clinic on weekdays. The crosses represent mean values recorded on Fridays. They were not used when the control chart was calculated

Perils of Risk-Adjusting Observational Data

A process variable like the above is straightforward to use for monitoring a single clinical process and for comparing several healthcare providers during a specified period. If the variable is an outcome measure in patients, risk adjustment may be necessary to adjust for variation between the patients in terms of severity of disease, comorbidities, etc. [3]. However, statistical risk adjustment of observational data has the potential of generating results that may be very misleading [11]. Thus, great care should be taken when applying the results. The following study simulates a

4 Statistical Process Control in Clinical Medicine

Table 2 The results of a simulation study where the probability of death during a surgical procedure as a function of two risk factors (presence of a genetic trait and presence of infection of the organ operated on) was simulated. The probability of death for given risk factor combination, is the same in the two hospitals, but the case mix differs between the hospitals

Hospital ^(a)	Genetic trait ^(b)	Infected organ ^(c)	Simulated model ^(d)	Result of simulation ^(e)		Genetic trait assumed unknown ^(f)	
				Mortality ^(g)	<i>n</i> (number of dead) ^(h)	Mortality ⁽ⁱ⁾	<i>n</i> (number of dead) ⁽ⁱ⁾
1	No	No	$p = 0.100$	0.060	100 (6)	0.227	300 (68)
	Yes	No	$p = 0.300$	0.310	200 (62)		
	No	Yes	$p = 0.300$	0.300	200 (60)		
	Yes	Yes	$p = 0.600$	0.574	500 (287)		
2	No	No	$p = 0.100$	0.096	700 (67)	0.131	840 (110)
	Yes	No	$p = 0.300$	0.307	140 (43)		
	No	Yes	$p = 0.300$	0.233	120 (28)		
	Yes	Yes	$p = 0.600$	0.825	40 (33)		

^(a) It is assumed that the quality of treatment and care is the same at the two hospitals.

^(b) The presence of the genetic trait influences the probability of death (see parameters of simulation model in column 4).

^(c) The presence of infection of the organ operated on influences the probability of death (see parameters of simulation model in column 4).

^(d) The parameter value of the simulation model generating the data. For instance when the genetic trait is present and the organ is infected the probability of death is 0.600.

^(e) Columns 5 and 6 show the result of the simulation for each combination of risk factors in each hospital.

^(f) Columns 7 and 8 show how the results would have been classified had the information of the genetic trait been unknown to the medical community. For instance, the results of the two first set of experiments would have been pooled and classified as data from patients without organ infection, etc.

^(g) The mortality observed in each of the eight sets of simulation experiments. When compared to the value in column 4 one may appreciate the effect of the random variation.

^(h) *n* is the number of independent simulations done per experiment. *n* has been varied between the hospitals to imitate difference between the case mix of the two hospitals. Number of dead is the number of fatal outcomes occurring during the *n* simulations.

⁽ⁱ⁾ Using the same data that were generated by the simulation. For given value of “infected organ” and given hospital the data have been pooled imitating the situation where the risk factor “genetic trait” is unknown to the medical community.

simple and intuitively obvious mechanism whereby misleading results may be produced.

Imagine two hospitals both receiving patients belonging to the same well-defined clinical entity and suffering from the same grave disease that requires major surgery. Assume that the same surgical procedure is used in both hospitals and that two patient factors influence these patients’ risk of dying during operation: the presence of a specific genetic trait (R_1 : 1 if the factor is present, otherwise 0) and the presence of infection of the organ to be operated on (R_2 : 1 if present, otherwise 0). Patients with neither risk factor present have a 10% probability of dying during the procedure, regardless of the hospital they are admitted to. For patients with one factor present the probability is 30%, and for those with both factors present it is 60%.

We now conduct a simulation experiment based on the above mentioned assumptions simulating the

experience of 2000 patients admitted to the two hospitals. To study the effect of patient case mix we vary the patient case mix between the hospitals. We examine the effect of risk adjustment when both risk factors are known and when only risk factor 2 is known.

Table 2 shows for each of the above two hospitals the case mix of 1000 patients admitted to the hospital and the results of the simulation experiments. These results have been analyzed under the assumption that both risk factors are known to the medical community (columns 5 and 6) and under the assumption that only one is known (columns 7 and 8). For example row 1 shows the simulated probability of death in column 4. It is 0.100 because the patients of this category do not have any of the two risk factors present (see columns 2 and 3). The result of the simulation of 100 patients’ experience is that 6 patients died (column 6). Row 2 shows the corresponding results for patients

with risk factor 1 present and risk factor 2 absent. Here the simulated probability of death is 0.300 and the result of the simulation that 62 of the 200 patients died giving a simulated death rate of 0.310. Columns 7 and 8 show how the results would have been interpreted had risk factor 1 been unknown to the medical community. Now the two types of patients cannot be distinguished and we calculate the mortality for patients with risk factor 1 absent as $68/300 = 0.227$, etc.

The overall mortality rates (hospital 1: 0.415 and hospital 2: 0.171) differ significantly between hospitals. However, we need to adjust for the difference in case mix. To do so, we perform a regression of risk factors and the hospital effect on $\text{logit}(p_{ij})$ where p_{ij} is the probability that patient i , treated at hospital j , dies. The independent variables include R_1 , R_2 , and H_1 (1 if the patient is treated at hospital 1 and 0 if treated at hospital 2). We have the model

$$\text{logit}(p_{ij}) = \alpha + \delta_1 H_1 + \beta_1 R_1 + \beta_2 R_2 + \varepsilon_{ij} \quad (1)$$

Table 3 shows the result of the logistic regression analysis. The coefficients of both risk factors are highly significantly different from 0 while that of the hospital effect is not. This is no surprise to us because we conducted the experiment so that the hospitals only differed in terms of their case mixes.

Table 2 (columns 7 and 8) shows how the data resulting from the simulation would have been classified had the genetic trait been unknown. Patients

with both risk factors present would not be distinguishable from patients with only risk factor 2 present and patients with only risk factor 1 present would be classified as low-risk patients. The logistic regression model therefore becomes

$$\text{logit}(p_{ij}) = \alpha + \delta_1 H_1 + \beta_2 R_2 + \varepsilon_{ij} \quad (2)$$

Table 3 shows that now the hospital coefficient is significantly different from 0 implying that the mortality rate of hospital 1 is higher than that of hospital 2 even if we adjust for case-mix differences. Thus, the presence of an unknown or unmeasured significant risk factor may bias the comparisons between healthcare providers. Needless to say that following a simulation of random assignment of the patients to the two hospitals the death rates did not differ significantly, neither with nor without risk adjustment.

Empirical studies [11] have demonstrated that statistical adjustment of observational data may fail to remove the main part of bias and occasionally dramatically increase systematic bias. This may, e.g., result from omitted or unknown risk factors (as illustrated in the above example), mis-specification of continuous variables (inappropriate conversion to binary variables or failure to recognize nonlinear relationships), misclassification through the use of poor proxies for the proper **covariate**, measurement errors, and within patient instability in covariate (e.g., because of circadian rhythms). As a minimum

Table 3 The estimated coefficients of two logistic regressions of hospital (H_1)^(a) and risk factors^(b) (R_1 and R_2) on $\text{logit}(p)$, where p is the probability of dying during an operation. In the first regression analysis^(c) both risk factors were known. In the second analysis^(d) it was assumed that only one of the risk factors (factor 2) was known, imitating the situation where an important risk factor is unknown to the medical community

Parameter	All risk factors known				One risk factor assumed unknown			
	Estimate	SE ^(e)	p ^(f)	Estimated odds ratio ^(g)	Estimate	SE	p	Estimated odds ratio
β_2	1.30	0.127	<0.00005	3.67	1.29	0.123	<0.00005	3.63
δ_1	-0.139	0.144	0.34	0.871	0.57	0.125	<0.00005	1.77
β_1	1.46	0.127	<0.00005	4.32	Factor assumed unknown			

^(a) H_1 is a binary variable that is 1 if the patient was operated at hospital 1 and 0 if he/she was operated at hospital 2.

^(b) Risk factor 1 (R_1), "presence of genetic trait". $R_1 = 1$ if trait is present, otherwise 0. Risk factor 2 (R_2), "presence of infection" in organ operated on'. $R_2 = 1$ if infection is present, otherwise 0.

^(c) $\text{logit}(p) = \alpha + \beta_1 R_1 + \beta_2 R_2 + \delta_1 H_1 = -2.23 + 1.46R_1 + 1.30R_2 - 0.139H_1$.

^(d) $\text{logit}(p) = \alpha + \beta_2 R_2 + \delta_{11} = -1.86 + 1.29R_2 + 0.57H_1$.

^(e) SE = **standard error** of estimate.

^(f) p-value.

^(g) e^{estimate} . For example for the parameter β_2 we have $e^{1.30} = 3.67$.

the following six points may be considered when statistical risk adjustment of observational data is attempted:

1. Data are first presented without adjustment.
2. Continuous variables such as age, weight, etc., are not dichotomized since this throws away information and increases the likelihood of improper model specification.
3. The assumptions of the model building such as linearity, etc. are checked (see [12]).
4. When appropriate, a hierarchical model is used [9, 10].
5. The variation not accounted for by the risk-adjustment model is measured and reported.
6. As a rule of thumb observational data are considered hypothesis generating material and treated accordingly and very explicitly presented as such.

Acknowledgments

Our thanks are due to Christian Gluud for a thorough and constructive review.

References

- [1] Montgomery, D.C. (2001). *Introduction to Statistical Quality Control*, 4th Edition, John Wiley & Sons, New York.
- [2] Winkel, P. & Zhang, N. F. (2007). *Statistical Development of Quality in Medicine*, John Wiley & Sons, Chichester.
- [3] Grigg, O. & Farewell, V. (2004). An overview of risk-adjusted charts, *Journal of The Royal Statistical Society. Series A* **167**, 523–539.
- [4] Cebul, R.D. (1991). Randomized, controlled trials using the metro firm system, *Medical Care* **29**(Suppl. 7), JS9–JS18.
- [5] Iezzoni, L.I. (1997). *Risk Adjustment for Measuring Healthcare Outcomes*, Health Administration Press, Chicago.
- [6] Rosenbaum, P.R. & Rubin, D.B. (1983). The central role of the propensity score in observational studies for causal effects, *Biometrika* **70**, 41–55.
- [7] D’Agostino Jr, R.B. (1998). Tutorial in biostatistics: propensity score methods for bias reduction in the comparison of a treatment to a nonrandomized control group, *Statistics in Medicine* **17**, 2265–2281.
- [8] Zhang, N.F., Stenback, G.A. & Wardrop, D.M. (1990). Interval estimation of process capability index Cpk, *Communications in Statistics: Theory and Methods* **19**, 4455–4470.
- [9] Greenland, S. (2000). Principles of multilevel modeling, *International Journal of Epidemiology* **29**, 158–167.
- [10] Goldstein, H., Browne, W. & Rasbash, J. (2002). Tutorial in biostatistics multilevel modeling of medical data, *Statistics in Medicine* **21**, 3291–3315.
- [11] Deeks, J.J., Dinnes, J., D’Amico, R., Sowden, A.J., Sakarovitch, C., Song, F., Petticrew, M. & Altman, D.G. (2003). Evaluating non-randomised intervention studies, *Health Technology Assessment* **7**, 1–186.
- [12] Bagley, S.C., White, H. & Golomb, B.A. (2001). Logistic regression in the medical literature: standards for use and reporting, with particular attention to one medical domain, *Journal of Clinical Epidemiology* **54**, 979–985.

PER WINKEL AND NIEN FAN ZHANG

Coherent Systems

The Structure of a Coherent System

Consider a system with n components. Let c_i denote the i th component. For now, we will consider a fixed time t . Define

$$x_i = \begin{cases} 1, & \text{if } c_i \text{ works} \\ 0, & \text{if } c_i \text{ fails} \end{cases} \quad (1)$$

The state of the system depends solely on the state of its components. A system works if certain combinations of components work. Define

$$\phi(x) = \begin{cases} 1, & \text{if system works} \\ 0, & \text{if system fails} \end{cases} \quad (2)$$

where $x = (x_1, \dots, x_n)$.

The function $\phi : \{0, 1\}^n \rightarrow \{0, 1\}$ is called the *structure function* of the system. Let us take a look at some examples. A *series system* works if and only if all of its components work. So,

$$\phi(x) = \min\{x_1, \dots, x_n\} = \prod_{i=1}^n x_i \quad (3)$$

A *parallel system* works if and only if at least one of its components works. So,

$$\phi(x) = \max\{x_1, \dots, x_n\} \quad (4)$$

Define the *cup operator* $\coprod$ through its Boolean addition properties:

$$0 \coprod 0 = 0 \quad (5)$$

$$0 \coprod 1 = 1 \quad (6)$$

$$1 \coprod 0 = 1 \quad (7)$$

$$1 \coprod 1 = 1 \quad (8)$$

We can write the structure function for a parallel system as

$$\phi(x) = \coprod_{i=1}^n x_i \quad (9)$$

Note that for all $x \in \{0, 1\}^n$, $\coprod x_i = 1 - \prod (1 - x_i)$. Clearly, $\prod x_i \leq \coprod x_i$.

A **k-out-of-n system** works if and only if at least k of its components work. The structure function for a 2-out-of-3 system may be written as

$$\phi(x) = (x_1 x_2) \coprod (x_1 x_3) \coprod (x_2 x_3) \quad (10)$$

Define

$$(1_i, x) = (x_1, \dots, x_{i-1}, 1, x_{i+1}, \dots, x_n) \quad (11)$$

$$(0_i, x) = (x_1, \dots, x_{i-1}, 0, x_{i+1}, \dots, x_n) \quad (12)$$

A structure function ϕ may be represented as

$$\phi(x) = x_i \phi(1_i, x) + (1 - x_i) \phi(0_i, x) \quad (13)$$

for any fixed i . This is known as the *pivotal decomposition of a structure function*.

Definition 1 The i th component in a system with structure function ϕ is irrelevant if

$$\phi(1_i, x) = \phi(0_i, x) \quad (14)$$

for all x . Otherwise, we say the i th component is relevant.

Definition 2 A structure function is monotone if

$$x_i \leq y_i, \text{ for all } i, \text{ implies } \phi(x) \leq \phi(y) \quad (15)$$

A component is *relevant* when there exists at least one case where the state of the system depends on whether the component works or has failed. A structure function is *monotone* when the state of the system will never improve after the failure of a component.

Definition 3 A system is coherent if each of its components is relevant and its structure function is monotone.

A coherent system consists of components that, when working, never harm the system and improve the system in at least some instance. Thus, each working component is beneficial to the system.

There are five coherent systems with $n = 3$ components. The different structure functions are

$$\phi_1 = x_1 x_2 x_3 \quad (16)$$

$$\phi_2 = x_1 \coprod x_2 \coprod x_3 \quad (17)$$

2 Coherent Systems

$$\phi_3 = (x_1 x_2) \coprod (x_1 x_3) \coprod (x_2 x_3) \quad (18)$$

$$\phi_4 = x_1 \coprod (x_2 x_3) \quad (19)$$

$$\phi_5 = x_1 \left(x_2 \coprod x_3 \right) \quad (20)$$

There exists 20 coherent systems with $n = 4$ components. The interesting question of how many coherent systems exist in general for n components is open.

It seems reasonable to say that system B with structure function ϕ_B is better than system A with structure function ϕ_A if

$$\phi_A(x) \leq \phi_B(x) \text{ for all } x \in \{0, 1\}^n \quad (21)$$

with a strict inequality at least once.

Theorem 1 For any coherent system in n components with structure function ϕ ,

$$\prod_{i=1}^n x_i \leq \phi(x) \leq \coprod_{i=1}^n x_i \quad (22)$$

Proof Let $\mathbf{0} = (0, \dots, 0)$ and $\mathbf{1} = (1, \dots, 1)$. All coherent systems have the property that $\phi(\mathbf{0}) = 0$ and $\phi(\mathbf{1}) = 1$.

If $\prod x_i = 1$, then $\coprod x_i = 1$ and $x = \mathbf{1}$. Since $\phi(\mathbf{1}) = 1$, the result is established.

If $\prod x_i = 0$, then $\prod x_i = 0$ and $x = \mathbf{0}$. Since $\phi(\mathbf{0}) = 0$, the result is established.

If $\prod x_i = 0$, and $\coprod x_i = 1$, then the result is trivial.

Thus, as one would believe intuitively, the series system is the worst arrangement of components in a coherent system and the parallel system is the best.

The Reliability of a Coherent System

Suppose components $c_1, \dots, c_n$ for a coherent system have probabilities of working at a fixed time t defined as $p_1, \dots, p_n$, respectively. The *reliability* of a system is the probability that the system works. Let $p = (p_1, \dots, p_n)$. The system reliability as a function of p is defined as

$$h(p) = P\{\phi(x) = 1\} \quad (23)$$

Assume that the component states are stochastically independent of each other. That is,

$$X_1, \dots, X_n \sim \text{indep BIN}(1, p_i) \quad (24)$$

So $E\{X_i\} = p_i$. The state of the system is also a binary random variable,

$$\phi(X) \sim \text{BIN}(1, h(p)) \quad (25)$$

so $E_p\{\phi(X)\} = h(p)$.

For a series system, the reliability function becomes

$$\begin{aligned} h(p) &= E\left\{\prod_{i=1}^n X_i\right\} \\ &= \prod_{i=1}^n E\{X_i\} = \prod_{i=1}^n p_i \end{aligned} \quad (26)$$

For a parallel system, the reliability function becomes

$$\begin{aligned} h(p) &= E\left\{\coprod_{i=1}^n X_i\right\} = E\left\{1 - \prod_{i=1}^n (1 - X_i)\right\} \\ &= 1 - \prod_{i=1}^n E\{1 - X_i\} \\ &= 1 - \prod_{i=1}^n (1 - p_i) = \prod_{i=1}^n p_i \end{aligned} \quad (27)$$

It is tempting to believe that $h(p)$ may be obtained from $\phi(x)$ simply by replacing x_i with p_i . This is not so. Consider the following example with

$$\begin{aligned} \phi(x) &= x_1 \left(x_2 \coprod x_3 \right) \\ &= x_1 (1 - (1 - x_2)(1 - x_3)) \\ &= x_1 x_2 + x_1 x_3 - x_1 x_2 x_3 \end{aligned} \quad (28)$$

So

$$\begin{aligned} h(p) &= E\{X_1 X_2 + X_1 X_3 - X_1 X_2 X_3\} \\ &= p_1 p_2 + p_1 p_3 - p_1 p_2 p_3 \\ &= p_1 \left(p_2 \coprod p_3 \right) \end{aligned} \quad (29)$$

Since x_1 is a binary variable, we have $x_1^2 = x_1$. We could alternatively write the structure function as

$$\begin{aligned} \phi(x) &= x_1 x_2 + x_1 x_3 - (x_1 x_2)(x_1 x_3) \\ &= (x_1 x_2) \coprod (x_1 x_3) \end{aligned} \quad (30)$$

But $p_1^2 \neq p_1$, so

$$p_1 (p_2 \coprod p_3) \neq (p_1 p_2) \coprod (p_1 p_3) \quad (31)$$

A pivotal decomposition can be defined for the reliability function h . Let $h(1_i, p) = E\{\phi(1_i, X)\}$ and $h(0_i, p) = E\{\phi(0_i, X)\}$. Then

$$\begin{aligned} h(p) &= E\{\phi(X)\} \\ &= E\{X_i \phi(1_i, X) + (1 - X_i) \phi(0_i, X)\} \\ &= p_i h(1_i, p) + (1 - p_i) h(0_i, p) \end{aligned} \quad (32)$$

where independence applies since $\phi(1_i, X)$ and $\phi(0_i, X)$ are functions of $X_1, \dots, X_{i-1}, X_{i+1}, \dots, X_n$ only.

Theorem 2 *The reliability function $h(p)$ of a coherent system in n components is strictly increasing in each p_i for all $p \in (0, 1)^n$.*

Proof From the pivotal decomposition,

$$\begin{aligned} \frac{\partial h(p)}{\partial p_i} &= h(1_i, p) - h(0_i, p) \\ &= E\{\phi(1_i, X) - \phi(0_i, X)\} \end{aligned} \quad (33)$$

Since ϕ is monotone, $\phi(1_i, x) \geq \phi(0_i, x)$ for all x . So

$$\frac{\partial h(p)}{\partial p_i} \geq 0 \quad (34)$$

Since each component is relevant, there exists a case x^* such that $\phi(1_i, x^*) > \phi(0_i, x^*)$. At each $p \in (0, 1)^n$,

$$P\{X = x^*\} = \prod_{i=1}^n p_i^{x_i^*} (1 - p_i)^{1-x_i^*} > 0 \quad (35)$$

So

$$\frac{\partial h(p)}{\partial p_i} > 0, \text{ for all } p \in (0, 1)^n \quad (36)$$

Each working component is beneficial to the system, so an increased probability of a component working necessarily increases the probability that the system works.

The Signature of a Coherent System

We now consider system reliability as a function of time t . Let $T_1, \dots, T_n$ denote the lifetimes of components $c_1, \dots, c_n$, respectively. Assume that component lifetimes are independent and identically distributed with cumulative probability function $F(t) = P\{T_i \leq t\}$ and survival function

$$\overline{F}(t) = P\{T_i > t\} \quad (37)$$

At time t , component c_i works if and only if its lifetime surpasses t . So

$$[X_i = 1] \iff [T_i > t] \quad (38)$$

and $p_i = \overline{F}(t)$ for $i = 1, \dots, n$.

Let T denote the lifetime of a coherent system with structure function ϕ . At time t , the system works when system's lifetime surpasses t . So

$$[\phi(X) = 1] \iff [T > t] \quad (39)$$

Thus, the survival function for the system is

$$\begin{aligned} \overline{F}_T(t) &= P\{T > t\} \\ &= P\{\phi(X) = 1\} \\ &= h(p_1, \dots, p_n) \\ &= h(\overline{F}(t), \dots, \overline{F}(t)) \end{aligned} \quad (40)$$

Let $T_{(1)} \leq T_{(2)} \leq \dots \leq T_{(n)}$ denote the ordered component lifetimes. From the distribution theory for order statistics,

$$P\{T_{(k)} > t\} = \sum_{j=0}^{k-1} \binom{n}{j} (F(t))^j (\overline{F}(t))^{n-j} \quad (41)$$

Let $\pi_k = \Pr\{T = T_{(k)}\}$ be the probability that the system fails at the k th failure of a component. Define the *signature* of a coherent system as the probability vector $\pi = (\pi_1, \dots, \pi_n)$. The survival function for a coherent system can be written in terms of its signature. Write

$$\begin{aligned} \overline{F}_T(t) &= \sum_{k=1}^n P\{T = T_{(k)}\} \cdot P\{T > t \mid T = T_{(k)}\} \\ &= \sum_{k=1}^n \pi_k P\{T_{(k)} > t\} \\ &= \sum_{k=1}^n \sum_{j=0}^{k-1} \pi_k \binom{n}{j} (F(t))^j (\overline{F}(t))^{n-j} \end{aligned} \quad (42)$$

4 Coherent Systems

Consider the five coherent systems with structure functions given in equations (16–20). For the series system, $\pi(\phi_1) = (1, 0, 0)$ since a series system always fails at the first failure of a component. For the parallel system, $\pi(\phi_2) = (0, 0, 1)$ since a parallel system does not fail until the last failure of a component. A 2-out-of-3 system always fails upon the second failure of a component so that $\pi(\phi_3) = (0, 1, 0)$.

There are $n!$ permutations for the orderings of component failures. We can use Table 1 to calculate the **signatures** for the three-component coherent systems ϕ_4 and ϕ_5 .

So $\pi(\phi_4) = (0, \frac{2}{3}, \frac{1}{3})$ and $\pi(\phi_5) = (\frac{1}{3}, \frac{2}{3}, 0)$.

The signature leads to a method for comparison of two coherent systems. Let π_A and π_B denote the signatures of systems A and B , respectively, where each system has n components. Let T_A and T_B denote the lifetimes of the two systems. Let $\bar{F}_A(t) = \Pr\{T_A > t\}$ and $\bar{F}_B(t) = \Pr\{T_B > t\}$. Define

$$T_A \leq^{st} T_B \quad (43)$$

if and only if

$$\bar{F}_A(t) \leq \bar{F}_B(t), \text{ for all } t \geq 0 \quad (44)$$

That is, system A lifetime is stochastically less than or equal to system B lifetime if and only if the probability of survival beyond any time t is never greater for system A than for system B .

Define

$$\pi_A \leq^{st} \pi_B \quad (45)$$

if and only if

$$\sum_{k=j}^n \pi_{Ak} \leq \sum_{k=j}^n \pi_{Bk}, \text{ for all } j = 1, \dots, n \quad (46)$$

Theorem 3 $\pi_A \leq^{st} \pi_B$ implies $T_A \leq^{st} T_B$.

Table 1 Ordered component failures

$T_{(1)} < T_{(2)} < T_{(3)}$	$T(\phi_4)$	$T(\phi_5)$
$T_1 < T_2 < T_3$	$T_{(2)}$	$T_{(1)}$
$T_1 < T_3 < T_2$	$T_{(2)}$	$T_{(1)}$
$T_2 < T_1 < T_3$	$T_{(2)}$	$T_{(2)}$
$T_2 < T_3 < T_1$	$T_{(3)}$	$T_{(2)}$
$T_3 < T_1 < T_2$	$T_{(2)}$	$T_{(2)}$
$T_3 < T_2 < T_1$	$T_{(3)}$	$T_{(2)}$

Proof Suppose $\pi_A \leq^{st} \pi_B$. From equation (42),

$$\begin{aligned} \bar{F}_A(t) &= \sum_{k=1}^n \sum_{j=0}^{k-1} \pi_{Ak} \binom{n}{j} (F(t))^j (\bar{F}(t))^{n-j} \\ &= \sum_{j=0}^{n-1} \left(\sum_{k=j+1}^n \pi_{Ak} \right) \binom{n}{j} (F(t))^j (\bar{F}(t))^{n-j} \\ &\leq \sum_{j=0}^{n-1} \left(\sum_{k=j+1}^n \pi_{Bk} \right) \binom{n}{j} (F(t))^j (\bar{F}(t))^{n-j} \\ &= \sum_{k=1}^n \sum_{j=0}^{k-1} \pi_{Bk} \binom{n}{j} (F(t))^j (\bar{F}(t))^{n-j} \\ &= \bar{F}_B(t) \end{aligned} \quad (47)$$

Thus, $T_A \leq^{st} T_B$.

Not all pairs of coherent systems have signatures that are comparable by the signature criterion. Consider, for example, the four-component coherent systems with structure functions

$$\phi_A(x) = x_1 \left(x_2 \coprod x_3 \coprod x_4 \right) \quad (48)$$

and

$$\phi_B(x) = (x_1 x_2) \coprod (x_3 x_4) \quad (49)$$

One can easily derive $\pi_A = (\frac{1}{4}, \frac{1}{4}, \frac{1}{2}, 0)$ and $\pi_B = (0, \frac{2}{3}, \frac{1}{3}, 0)$. So,

$$\pi_{A3} + \pi_{A4} > \pi_{B3} + \pi_{B4} \quad (50)$$

but

$$\pi_{A2} + \pi_{A3} + \pi_{A4} < \pi_{B2} + \pi_{B3} + \pi_{B4} \quad (51)$$

However, the coherent systems of three components can be compared and ordered by the signature criterion as

$$T(\phi_2) \leq^{st} T(\phi_5) \leq^{st} T(\phi_3) \leq^{st} T(\phi_4) \leq^{st} T(\phi_1) \quad (52)$$

For the 20 systems of four components, there are $\binom{20}{2} = 190$ possible comparisons. In 180 of these cases, a stochastic inequality can be implied from the signature criterion.

Further Reading

Barlow, R.E. & Proschan, F. (1981). *Statistical Theory of Reliability and Life Testing*, McArdle Press, Silver Spring.

Boland, P.J. & Samaniego, F.J. (2001). The signature of a coherent system. *Mathematical Reliability: An Expository Perspective*, Kluwer Academic Publishers.

ANDREW A. NEATH

Parallel, Series, and Series-Parallel Systems

Introduction

In the area of reliability, often the objective is to determine the probability that a system will function for a certain period. Such systems consist of multiple components connected according to a certain structure. Whether a system functions depends on how the components are connected as well as the state of each component. Each component will have a state that can be classified as either “functioning” or “failed”. In most cases, there is an underlying probability distribution that is used to determine if the component is in the functioning state or the failed state at some time t .

Three important system configurations that are widely used in reliability are series systems, parallel systems, and series-parallel systems. In a series system, components are connected in such a way that the failure of a single component results in system failure. A series system is also referred to as a **competing risks system** since the failure of an individual (system) can be classified as one of n possible risks (component) that competes for the individual’s failure. An example of a series system is illustrated in Agustin and Peña [1] where a series system of softwares is considered with the failure of a software (component) resulting in the failure of the system. On the other hand, a system where the failure of all components will result in the failure of the system is called a *parallel system*. Some papers that have considered parallel systems include Kvam and Peña [2] on an equal load-sharing dynamic model and Sinkovic *et al.* [3] in the area of communication. A common interest in the aforementioned papers is in modeling how the load can be shared by the components in parallel whenever critical jobs or tasks are given to the system to complete. In addition, standby or redundant systems can also be associated with parallel systems (see Osaki and Nakagawa [4]). For such systems, the operational unit can be considered to be connected in parallel with all the standby units. As soon the unit fails, one of the standby units goes into operation. Finally, a series-parallel system is a system in which m subsystems are connected in series and

each subsystem consists of k components connected in parallel (see Baxter and Harche [5]). For such a system, failure is observed if a subsystem has failed. Note that the failure of a subsystem occurs when all of the k components have failed, hence the name *series-parallel system*.

One of the goals in reliability is to ensure that at any given instant, we observe a functioning system with a high probability. The probability that the system is functioning is defined as the system reliability. In this article, we will obtain the system reliability for the series, parallel, and series-parallel systems. The system reliability will depend on the reliability of each component as well as on how the components are connected. Moreover, we will assume that the reliability of each component is a function of how long it is observed. The time to failure of a component will depend on a known probability distribution.

System Reliability

Structure Function

Consider a system of n components in which a component is either in a “functioning” state or in a “failed” state. For $i = 1, \dots, n$, let X_i be the binary random variable defined to be

$$X_i = \begin{cases} 1, & \text{component } i \text{ is functioning} \\ 0, & \text{otherwise} \end{cases} \quad (1)$$

The random variables $X_1, \dots, X_n$ are assumed to be independent. Thus, the components function independent of one another. We will define $\mathbf{X} = (X_1, \dots, X_n)$ to be the state vector of the system. The state of the system depends on a function that considers how the components are connected as well as the corresponding state of each component. This is called the *structure function* of the system and is denoted by $\phi(\mathbf{X})$. The function $\phi(\mathbf{X})$ is a binary random variable where

$$\phi(\mathbf{X}) = \begin{cases} 1, & \text{system is functioning} \\ 0, & \text{otherwise} \end{cases} \quad (2)$$

To describe the systems of interest in this article, we need to define the following: vector inequality, relevant components, and a nondecreasing structure function.

Definition Let $\mathbf{X} \in \mathbf{R}^n$ and $\mathbf{Y} \in \mathbf{R}^n$ be two state vectors.

2 Parallel, Series, and Series-Parallel Systems

1. We say that $X \leq Y$ if $X_i \leq Y_i$ for all $i = 1, \dots, n$. We have $X < Y$ if $X_j < Y_j$ for at least one j .
2. Component i is irrelevant if for all values of X_j , where $j = 1, \dots, n$ and $j \neq i$,

$$\begin{aligned} &\phi(X_1, \dots, X_{i-1}, 0, X_{i+1}, \dots, X_n) \\ &= \phi(X_1, \dots, X_{i-1}, 1, X_{i+1}, \dots, X_n) \end{aligned} \quad (3)$$

3. A structure function is a nondecreasing function if $X \leq Y$ implies that $\phi(X) \leq \phi(Y)$.

The primary interest in reliability are systems where replacing a failed component by a functioning component does not result in the system being “worse off”. In other words, the structure function of the system is nondecreasing. Such a system is called a *coherent system* (see **Coherent Systems**).

Definition A system of n components where all components are relevant and whose structure function is nondecreasing is called a *coherent system*.

In what follows, we will show that a series system, a parallel system, and a series-parallel system are all examples of coherent systems.

Series System. A series system is defined to be a system that functions if and only if all components are functioning. One can also consider such a system as having failed as soon as one component has failed. Thus, the structure function of a series system is

$$\phi(X) = X_1 X_2 \dots X_n = \prod_{i=1}^n X_i \quad (4)$$

Since the random variables $X_1, \dots, X_n$ are all binary, the structure function of a series system can be written as $\phi(X) = \min(X_1, \dots, X_n)$. Figure 1 illustrates a series system with three components.

As an illustration, consider a series system where component 1 has failed while the remaining $n - 1$ components are functioning. Thus, the state vector is

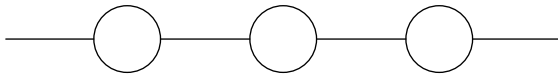


Figure 1 Diagram for a series system with three components

$X = (X_1, X_2, \dots, X_n) = (0, 1, \dots, 1)$, which results in $\phi(X) = \min(0, 1, \dots, 1) = 0$. On the other hand, if all components are functioning, then $\phi(X) = \min(1, \dots, 1) = 1$. Clearly, $\phi(X) = 0$ whenever $X_i = 0$ for some i . Thus, to show that a series system is a coherent system, we need to consider two state vectors X and Y such that $X_i = 0$ for some $i = 1, \dots, n$ while $Y_j = 1$ for all $j = 1, \dots, n$. Thus, $X < Y$ implies that $\phi(X) \leq \phi(Y)$, which means that ϕ is nondecreasing. In addition, it is easily seen that a series system with state vector $Y = (1, 1, \dots, 1)$ will fail as soon as one component fails. Hence, each component in a series system is relevant. It follows that a series system is coherent.

Parallel System. A parallel system is a system that functions if and only if at least one component is functioning. Since $1 - X_j = 0$ whenever component j has failed, the structure function of a parallel system is

$$\begin{aligned} \phi(X) &= 1 - (1 - X_1) \dots (1 - X_n) \\ &= 1 - \prod_{i=1}^n (1 - X_i) \end{aligned} \quad (5)$$

It should be noted that an alternative way of expressing the structure function of a parallel system is $\phi(X) = \max(X_1, \dots, X_n)$. Figure 2 presents a diagram for a parallel system with two components.

To illustrate, suppose that all components have failed except for component 1. Thus, $X_1 = 1$ while $X_2 = 0, \dots, X_n = 0$. It follows that $\phi(X) = \max(1, 0, \dots, 0) = 1$. On the other hand, if $X_1 = 0, \dots, X_n = 0$, then $\phi(X) = \max(0, \dots, 0) = 0$. To show that a parallel system is coherent, it suffices to consider two state vectors X and Y such that $X_j = 0$ for all $j = 1, \dots, n$ and $Y_i = 1$ for some $i = 1, \dots, n$. Clearly, $X < Y$ and it follows that $\phi(X) < \phi(Y)$. Moreover, every component is relevant since the structure function is equal to zero whenever all components have failed whereas the structure function is one if the state of one component changes its state from failed to functioning. Hence, a parallel system is coherent.

Series-Parallel System. We define a series-parallel system with mk components as a system comprised of m subsystems connected in series. Each subsystem consists of k components connected in

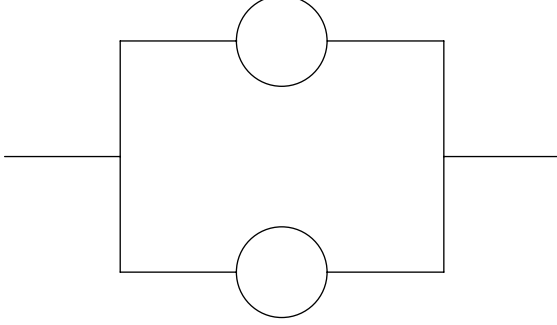


Figure 2 Diagram for a parallel system with two components

parallel. Thus, if at least one component in each subsystem is functioning, then the system is functioning. On the other hand, if all components in one subsystem have failed, then the system has failed regardless of the state of the components in the other subsystems. To obtain the structure function of a series-parallel system, let us define the structure function of the j th subsystem to be $\phi_j(X_j)$, for $j = 1, \dots, m$. The vector $X_j = (X_{j1}, \dots, X_{jk})$ represents the state vector of the j th subsystem, where X_{jl} is the random variable that represents the state of the l th component in the subsystem, for $l = 1, \dots, k$. One can obtain

$$\phi_j(X_j) = 1 - \prod_{l=1}^k (1 - X_{jl}) = \max(X_{j1}, \dots, X_{jk}) \quad (6)$$

Hence, by setting $X = (X_1, \dots, X_m)$, the structure function of the series-parallel system is

$$\phi(X) = \prod_{j=1}^m \phi_j(X_j) = \prod_{j=1}^m \left[1 - \prod_{l=1}^k (1 - X_{jl}) \right] \quad (7)$$

An alternative form for the structure function of the series-parallel system is

$$\phi(X) = \min_{1 \leq j \leq m} \left[\max_{1 \leq l \leq k} X_{jl} \right] \quad (8)$$

To provide an illustration of a series-parallel system, consider a series-parallel system with $m = 3$ subsystems with each subsystem containing $k = 2$ components. Hence, the system has a total of six components. An illustration for this system is given in Figure 3. For example, suppose that $X_{11} = 0$, $X_{12} = 0$, and $X_{22} = 0$ while the remaining X_{jl} are all equal to 1. One can easily see that subsystem 1 has failed while the remaining subsystems are still functioning. It follows that

$$\phi_1(X_{11}, X_{12}) = \max(0, 0) = 0 \quad (9)$$

$$\phi_2(X_{21}, X_{22}) = \max(1, 0) = 1 \quad (10)$$

$$\phi_3(X_{31}, X_{32}) = \max(1, 1) = 1 \quad (11)$$

which results in $\phi(X) = \min(0, 1, 1) = 0$. To show that a series-parallel system is coherent, consider a state vector Y such that all subsystems are functioning. In other words, for all $j = 1, \dots, m$, there exists a component, denoted by l , such that $Y_{jl} = 1$. In addition, consider a state vector X such that there exists a subsystem j where $X_{j1} = \dots = X_{jk} = 0$. It follows that $X < Y$ and $\phi(X) < \phi(Y)$. Moreover, the relevance of each component follows from the fact that components in series systems and parallel systems are all relevant.

System Reliability

Let $p_i = P\{X_i = 1\}$ be the probability that the i th component is functioning and suppose $\mathbf{p} = (p_1, \dots, p_n)$. We define

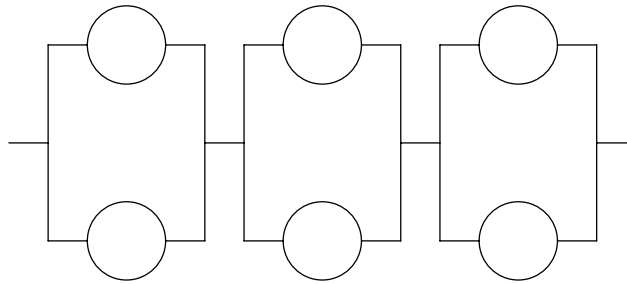


Figure 3 Diagram for a series-parallel system with $m = 3$ subsystems and $k = 2$ components per subsystem

4 Parallel, Series, and Series-Parallel Systems

$$h(\mathbf{p}) = E[\phi(X)] = P\{\phi(X) = 1\} \quad (12)$$

to be the reliability of the system. Recall that $X_1, \dots, X_n$ are independent random variables. Thus, for a series system, the reliability of the system is

$$\begin{aligned} h(\mathbf{p}) &= P\{X_1 = 1, \dots, X_n = 1\} \\ &= P\{X_1 = 1\} \dots P\{X_n = 1\} = \prod_{i=1}^n p_i \end{aligned} \quad (13)$$

On the other hand, to find the reliability of a parallel system, it is more convenient to write the reliability as

$$h(\mathbf{p}) = 1 - P\{\phi(X) = 0\} \quad (14)$$

By the independence of the X_i 's, we get

$$\begin{aligned} h(\mathbf{p}) &= 1 - P\{X_1 = 0, \dots, X_n = 0\} \\ &= 1 - P\{X_1 = 0\} \dots P\{X_n = 0\} \end{aligned} \quad (15)$$

It follows that the reliability of a parallel system is

$$\begin{aligned} h(\mathbf{p}) &= 1 - [1 - P\{X_1 = 1\}] \dots [1 - P\{X_n = 1\}] \\ &= 1 - \prod_{i=1}^n [1 - p_i] \end{aligned} \quad (16)$$

For a series-parallel system, we first define for $j = 1, \dots, m$ and $l = 1, \dots, k$, $p_{jl} = P\{X_{jl} = 1\}$ to be the probability that the l th component in the j th subsystem is functioning. We will set $\mathbf{p} = (\mathbf{p}_1, \dots, \mathbf{p}_m)$ where $\mathbf{p}_j = (p_{j1}, \dots, p_{jk})$. Since the subsystems are connected in series, it follows from equation (13) that

$$\begin{aligned} h(\mathbf{p}) &= P\{\phi_1(X_1) = 1, \dots, \phi_m(X_m) = 1\} \\ &= \prod_{j=1}^m P\{\phi_j(X_j) = 1\} \end{aligned} \quad (17)$$

Since the components in a subsystem are connected in parallel, we use equation (16) to show that

$$P\{\phi_j(X_j) = 1\} = 1 - \prod_{l=1}^k (1 - p_{jl}) \quad (18)$$

Thus, the reliability of a series-parallel system is

$$\begin{aligned} h(\mathbf{p}) &= \left[1 - \prod_{l=1}^k (1 - p_{1l}) \right] \dots \left[1 - \prod_{l=1}^k (1 - p_{ml}) \right] \\ &= \prod_{j=1}^m \left[1 - \prod_{l=1}^k (1 - p_{jl}) \right] \end{aligned} \quad (19)$$

Time-Dependent System Reliability

In this section, we will consider a system of components where the lifetime of the i th component is a random variable Y_i . For example, a car (system) will have parts (components) such as the wheel, engine, or air conditioning system where the time to failure of each part is a random variable that depends on a number of factors (e.g., driving conditions, manufacturing process of the part, etc.) We will assume that the components operate independently of one another. In other words, the component lifetime variables $Y_1, Y_2, \dots, Y_n$ are independent. Our main interest is to obtain the system reliability on the basis of the underlying assumption on the distribution of the component lifetimes. In doing so, we need to provide the appropriate representation of the distribution of the lifetime of a system. Leemis [6] provides an extensive discussion on five functions that are typically used in reliability. We will focus on three of these functions, namely, the probability density function, the survivor function, and the hazard function.

Let Y_i be a random variable with **probability density function** $f_i(y; \theta)$ and cumulative distribution function $F_i(y; \theta) = P\{Y_i \leq y\}$, for $i = 1, \dots, n$ and some parameter θ . It follows that the reliability of component i is

$$\begin{aligned} S_i(y; \theta) &\equiv P\{Y_i > y\} = \int_y^\infty f_i(u; \theta) du \\ &= 1 - F_i(y; \theta) \end{aligned} \quad (20)$$

The reliability of component i given in equation (20) is also known as the survivor function of Y_i . Since $F_i(y; \theta) = \int_0^y f_i(t; \theta) dt$, the density function of the i th component, in terms of the reliability of component i , is

$$f_i(y; \theta) = -\frac{d}{dy} S_i(y; \theta) = -S'_i(y; \theta) \quad (21)$$

Another representation for the lifetime of a component is *via* the hazard function $h_i(y; \theta)$. This is especially important due to its interpretation as the amount of risk associated with a component. In terms of the **component reliability**, the hazard function can be expressed as

$$\begin{aligned} h_i(y; \theta) &= \lim_{\Delta y \searrow 0} \frac{S_i(y; \theta) - S_i(y + \Delta y; \theta)}{S_i(y; \theta) \Delta y} \\ &= -\frac{S'_i(y; \theta)}{S_i(y; \theta)} = \frac{f_i(y; \theta)}{S_i(y; \theta)} \end{aligned} \quad (22)$$

Let Y be the random variable associated with the system lifetime and $S(y; \theta)$ be the system reliability. The corresponding density function, cumulative distribution function, and hazard function of the system lifetime are $f(y; \theta)$, $F(y; \theta)$, and $h(y; \theta)$, respectively. In terms of the density function and cumulative distribution function, the hazard function is of the form

$$h(y; \theta) = \frac{f(y; \theta)}{S(y; \theta)} \quad (23)$$

Our main interest is to obtain the appropriate expressions for the series, parallel, and series-parallel systems.

For a series system, note that the system lifetime can be expressed as $Y = \min(Y_1, \dots, Y_n)$. This results in a system reliability equal to

$$\begin{aligned} S(y; \theta) &= P\{Y > y\} = P\{\min(Y_1, \dots, Y_n) > y\} \\ &= P\{Y_1 > y, \dots, Y_n > y\} \end{aligned} \quad (24)$$

Since the component lifetimes are assumed to be independent, the system reliability simplifies to

$$S(y; \theta) = \prod_{i=1}^n P\{Y_i > y\} = \prod_{i=1}^n S_i(y; \theta) \quad (25)$$

Following equation (21), the density function of Y is obtained as

$$f(y; \theta) = -\sum_{i=1}^n S'_i(y; \theta) \prod_{\substack{j=1 \\ j \neq i}}^n S_j(y; \theta) \quad (26)$$

The hazard function for a series system is obtained by substituting the system reliability given in equation (25) and the density function given in equation (26) into equation (23).

For a parallel system, the random variable Y can be expressed as $Y = \max(Y_1, \dots, Y_n)$. It follows that the system reliability for a parallel system is

$$\begin{aligned} S(y; \theta) &= P\{\max(Y_1, \dots, Y_n) > y\} \\ &= 1 - P\{\max(Y_1, \dots, Y_n) \leq y\} \\ &= 1 - P\{Y_1 \leq y, \dots, Y_n \leq y\} \end{aligned} \quad (27)$$

By the independence of the component lifetimes, the resulting system reliability is

$$S(y; \theta) = 1 - \prod_{i=1}^n P\{Y_i \leq y\} = 1 - \prod_{i=1}^n [1 - S_i(y; \theta)] \quad (28)$$

The corresponding density function for the parallel system is

$$f(y; \theta) = -\sum_{i=1}^n S'_i(y; \theta) \prod_{\substack{j=1 \\ j \neq i}}^n [1 - S_j(y; \theta)] \quad (29)$$

For a series-parallel system with m subsystems connected in series with each subsystem containing k components connected in parallel, let Y_j be the lifetime of the j th subsystem. The random variable Y_{jl} denotes the lifetime of the l th component in the j th subsystem, for $j = 1, \dots, m$ and $l = 1, \dots, k$. Since the components are all independent, it follows that the subsystem lifetimes are also independent. By following the same arguments used to obtain equation (25), it follows that

$$S(y; \theta) = P\{Y > y\} = \prod_{j=1}^m P\{Y_j > y\} \quad (30)$$

Since $P\{Y_j \leq y\} = P\{\max(Y_{j1}, \dots, Y_{jk}) \leq y\}$, we can apply equation (28) for each subsystem to obtain the system reliability. Thus, after straightforward computations, equation (30) simplifies to

$$\begin{aligned} S(y; \theta) &= \prod_{j=1}^m \left[1 - \prod_{l=1}^k P\{Y_{jl} \leq y\} \right] \\ &= \prod_{j=1}^m \left[1 - \prod_{l=1}^k [1 - S_{jl}(y; \theta)] \right] \end{aligned} \quad (31)$$

6 Parallel, Series, and Series-Parallel Systems

From equation (31), the density function of Y for the series-parallel system is

$$f(y; \theta) = - \sum_{j=1}^m \left[\sum_{l=1}^k S'_{jl}(y; \theta) \prod_{\substack{l^*=1 \\ l^* \neq l}}^k [1 - S_{jl^*}(y; \theta)] \right] \times \left[\prod_{\substack{j^*=1 \\ j^* \neq j}}^m \left(1 - \prod_{l=1}^k (1 - S_{j^*l}(y; \theta)) \right) \right] \quad (32)$$

Example As an illustration, consider the case when each component lifetime follows an exponential distribution with mean $\theta = \lambda_i$, that is, $Y_i \sim \exp(\lambda_i)$. In other words, the random variable Y_i has density function

$$f_i(y; \lambda_i) = \frac{1}{\lambda_i} \exp\left(-\frac{y}{\lambda_i}\right), \quad y > 0 \quad (33)$$

Using equation (20), the survivor function of Y_i is

$$\begin{aligned} S_i(y; \lambda_i) &= \int_y^\infty \frac{1}{\lambda_i} \exp\left(-\frac{u}{\lambda_i}\right) du \\ &= \exp\left(-\frac{y}{\lambda_i}\right), \quad y > 0 \end{aligned} \quad (34)$$

while the hazard function is

$$h_i(y; \lambda_i) = \frac{f_i(y; \lambda_i)}{S_i(y; \lambda_i)} = \frac{1}{\lambda_i}, \quad y > 0 \quad (35)$$

Let $\lambda = (\lambda_1, \dots, \lambda_n)$ be the parameter vector of interest. For a series system with n components, equation (25) simplifies to

$$S(y; \lambda) = \prod_{i=1}^n \exp\left(-\frac{y}{\lambda_i}\right) = \exp\left(-y \sum_{i=1}^n \frac{1}{\lambda_i}\right) \quad (36)$$

Since the density function of Y is

$$\begin{aligned} f(y; \lambda) &= -\frac{d}{dy} S(y; \lambda) = \left(\sum_{i=1}^n \frac{1}{\lambda_i} \right) \\ &\times \exp\left(-y \sum_{j=1}^n \frac{1}{\lambda_j}\right) \end{aligned} \quad (37)$$

it follows that for a series system, the random variable Y is exponentially distributed with mean $[\sum_{i=1}^n (\lambda_i)^{-1}]^{-1}$. It can be easily seen that by using equation (23), the hazard function is

$$h(y; \lambda) = \sum_{i=1}^n \frac{1}{\lambda_i} \quad (38)$$

On the other hand, for a parallel system, using the survivor function of an exponential random variable given in equation (34) in equation (28) results in a system reliability equal to

$$S(y; \lambda) = 1 - \prod_{i=1}^n \left[1 - \exp\left(-\frac{y}{\lambda_i}\right) \right], \quad y > 0 \quad (39)$$

The resulting density function of the system lifetime for a parallel system is

$$\begin{aligned} f(y; \lambda) &= -\frac{d}{dy} S(y; \lambda) \\ &= \frac{d}{dy} \left[\prod_{i=1}^n \left(1 - \exp\left(-\frac{y}{\lambda_i}\right) \right) \right] \\ &= \sum_{i=1}^n \left(\frac{1}{\lambda_i} \right) \exp\left(-\frac{y}{\lambda_i}\right) \\ &\times \left[\prod_{\substack{j=1 \\ j \neq i}}^n \left(1 - \exp\left(-\frac{y}{\lambda_j}\right) \right) \right] \end{aligned} \quad (40)$$

This can also be obtained by replacing the appropriate expressions in equation (29) by the survivor function of an exponential random variable given in equation (34).

Finally, for a series-parallel system, let λ_{jl} be the mean lifetime of the l th component in the j th subsystem, for $j = 1, \dots, m$ and $l = 1, \dots, k$. The parameter vector of interest is $\lambda = (\lambda_{11}, \dots, \lambda_{1k}, \lambda_{21}, \dots, \lambda_{2k}, \dots, \lambda_{m1}, \dots, \lambda_{mk})$. Equation (31) simplifies to

$$S(y; \lambda) = \prod_{j=1}^m \left[1 - \prod_{l=1}^k \left(1 - \exp\left(-\frac{y}{\lambda_{jl}}\right) \right) \right], \quad y > 0 \quad (41)$$

Using equations (32) and (34), the density function of Y for the series-parallel system is

$$f(y; \lambda) = -\frac{d}{dy} S(y; \lambda)$$

$$\begin{aligned}
 &= \sum_{j=1}^m \left\{ \left[\prod_{\substack{j^*=1 \\ j^* \neq j}}^m \left(1 - \prod_{l=1}^k \left(1 - \exp \left(-\frac{y}{\lambda_{jl^*}} \right) \right) \right) \right] \right. \\
 &\quad \times \left[\sum_{l=1}^k \left(\prod_{\substack{l^*=1 \\ l^* \neq l}}^k \left(1 - \exp \left(-\frac{y}{\lambda_{jl^*}} \right) \right) \right) \right] \\
 &\quad \left. \times \left(\frac{1}{\lambda_{jl}} \right) \exp \left(-\frac{y}{\lambda_{jl}} \right) \right\} \quad (42)
 \end{aligned}$$

As a special case, suppose $Y_i \sim \exp(\alpha)$ for all $i = 1, \dots, n$. For ease of notation, we will let α be the parameter of interest. The system reliability of a series system simplifies to

$$S(y; \alpha) = \exp \left(-\frac{n}{\alpha} y \right) = \exp \left(-\frac{y}{\alpha/n} \right), \quad y > 0 \quad (43)$$

which shows that the system lifetime Y is exponentially distributed with mean α/n . On the other hand, the system reliability of a parallel system can be written as

$$S(y; \alpha) = 1 - \left[1 - \exp \left(-\frac{y}{\alpha} \right) \right]^n, \quad y > 0 \quad (44)$$

while the system reliability of a series-parallel system simplifies to

$$S(y; \alpha) = \left[1 - \left(1 - \exp \left(-\frac{y}{\alpha} \right) \right)^k \right]^m, \quad y > 0 \quad (45)$$

The appropriate density function for each system can then be obtained by differentiating the corresponding system reliability.

Example Whenever the risk of failure rapidly increases with time, a failure time distribution that is commonly used is the Weibull distribution. For instance, Ferdous *et al.* [7] assumed **software failure** times that follow the Weibull distribution. Their goal was to predict failure rates and mean time to failure of the software.

In this example, we consider the case where the lifetime for component i follows a Weibull distribution with parameters α_i and β_i . Hence, the

random variable Y_i has density function

$$f_i(y; \alpha_i, \beta_i) = \frac{\beta_i}{\alpha_i} \left(\frac{y}{\alpha_i} \right)^{\beta_i-1} \exp \left[-\left(\frac{y}{\alpha_i} \right)^{\beta_i} \right], \quad y > 0 \quad (46)$$

It is easily seen that $\beta_i = 1$ will result in $Y_i \sim \exp(\alpha_i)$. The survivor function of Y_i is

$$\begin{aligned}
 S_i(y; \alpha_i, \beta_i) &= \int_y^\infty \frac{\beta_i}{\alpha_i} \left(\frac{u}{\alpha_i} \right)^{\beta_i-1} \exp \left[-\left(\frac{u}{\alpha_i} \right)^{\beta_i} \right] du \\
 &= \exp \left[-\left(\frac{y}{\alpha_i} \right)^{\beta_i} \right], \quad y > 0 \quad (47)
 \end{aligned}$$

while the hazard function is

$$h_i(y; \alpha_i, \beta_i) = \frac{\beta_i}{\alpha_i} \left(\frac{y}{\alpha_i} \right)^{\beta_i-1}, \quad y > 0 \quad (48)$$

Let $\alpha = (\alpha_1, \dots, \alpha_n)$ and $\beta = (\beta_1, \dots, \beta_n)$ denote the two parameter vectors of interest. For a series system, the system reliability is

$$\begin{aligned}
 S(y; \alpha, \beta) &= \prod_{i=1}^n \exp \left[-\left(\frac{y}{\alpha_i} \right)^{\beta_i} \right] \\
 &= \exp \left[-\sum_{i=1}^n \left(\frac{y}{\alpha_i} \right)^{\beta_i} \right], \quad y > 0 \quad (49)
 \end{aligned}$$

In the special case where $\alpha_i = \alpha$ and $\beta_i = \beta$ for $i = 1, \dots, n$, equation (49) simplifies to

$$\begin{aligned}
 S(y; \alpha, \beta) &= \exp \left[-\sum_{i=1}^n \left(\frac{y}{\alpha} \right)^{\beta} \right] = \exp \left[-n \left(\frac{y}{\alpha} \right)^{\beta} \right] \\
 &= \exp \left[-\left(\frac{n^{1/\beta} y}{\alpha} \right)^{\beta} \right], \quad y > 0 \quad (50)
 \end{aligned}$$

which implies that Y is a Weibull random variable with parameters $\alpha/(n^{1/\beta})$ and β . On the other hand, by using equations (28) and (47), the system reliability for a parallel system is

$$S(y; \alpha, \beta) = 1 - \prod_{i=1}^n \left[1 - \exp \left[-\left(\frac{y}{\alpha_i} \right)^{\beta_i} \right] \right], \quad y > 0 \quad (51)$$

8 Parallel, Series, and Series-Parallel Systems

Finally, for a series-parallel system with k components in each of the m subsystems, the system reliability is

$$S(y; \alpha, \beta) = \prod_{j=1}^m \left[1 - \prod_{l=1}^k \left(1 - \exp \left[- \left(\frac{y}{\alpha_{jl}} \right)^{\beta_{jl}} \right] \right) \right], \quad y > 0 \quad (52)$$

References

- [1] Agustin, M. & Peña, E. (1999). A dynamic competing risks model, *Probability in the Engineering and Informational Sciences* **13**, 333–358.
- [2] Kvam, P. & Peña, E. (2005). Estimating load-sharing properties in a dynamic reliability system, *Journal of the American Statistical Association* **100**, 262–272.
- [3] Sinkovic, V., Lovrek, I. & Nemeth, G. (1999). Load balancing in distributed parallel systems for telecommunications, *Computing* **63**, 201–218.
- [4] Osaki, S. & Nakagawa, T. (1971). On a two-unit standby redundant system with standby failure, *Operations Research* **19**, 510–523.
- [5] Baxter, L. & Harche, F. (1992). On the optimal assembly of series-parallel systems, *Operations Research Letters* **11**, 153–157.
- [6] Leemis, L. (1995). *Reliability: Probabilistic Models and Statistical Methods*, Prentice Hall, Englewood Cliffs.
- [7] Ferdous, J., Borhan Uddin, M. & Pandey, M. (1995). Reliability estimation with Weibull inter failure times, *Reliability Engineering and System Safety* **50**, 285–296.

MARCUS A. AGUSTIN

***k*-out-of-*n* Systems**

Definition of *k*-out-of-*n* System Model

A system with n components is referred to as an n -component system. Each component is either working or not working (failed). The design of the system determines whether the system as a whole is working. A series system is considered working or “good” if all n components are working. This type of system is seen in decorative tree lights or a high-voltage multicell battery. If one component in a series system fails, then the entire system will fail. A parallel system is a system that will continue to work if at least one of the n components is working. A k -out-of- n system is considered working if at least k out of the n components are working. Examples of this type of system include an airplane with four engines that will work (fly) if at least two of the engines work, a stereo with two speakers that will work (may be heard) if at least one of the two speakers work, or a car with eight cylinders that will work (operate) if at least four of the cylinders are working (firing). A k -out-of- n system may be described in two ways, depending on the definition of the parameter, k , as follows. The parameter k may represent the number of components in the system that must work in order for the entire system to work, referred to as a k -out-of- n :G system. Or k may represent the number of components in the system that must fail before the entire system fails, referred to as a k -out-of- n :F system. The k -out-of- n :F system is equivalent to a $n - (k - 1)$ -out-of- n :G system. The series system and the parallel system are both special cases of a k -out-of- n system with $k = 1$ and $k = n$, respectively, for the k -out-of- n :F system.

Reliability Computations for *k*-out-of-*n* System (Assuming Independent and Identical Components with Known Reliabilities)

When the components are independent and identically distributed (i.i.d.), calculating the reliability of the system is easily done using the binomial distribution. In the context of the k -out-of- n :G system model, the random variable X represents the number of components out of the n components that work and the

probability that each component works is p . In other words, $P(X = x)$ means the probability that exactly x components work out of the n components in the system. The probability mass function of the binomial random variable X is

$$P(X = x) = \binom{n}{x} p^x (1 - p)^{n-x}, \quad x = 0, 1, 2, \dots, n \quad (1)$$

and the cumulative distribution function (CDF) is given by

$$P(X \leq x) = \sum_{i=0}^x P(X = i) \quad (2)$$

The reliability of the k -out-of- n :G system is the probability that at least k out of the n components work and is given as

$$P(X \geq k) = R(k, n, p) = \sum_{i=k}^n \binom{n}{i} p^i (1 - p)^{n-i} \quad (3)$$

where n is the number of components, p is the reliability of each component, and k is the minimum number of components that need to work for the system to work. The reliability of the system may be written as a function of the CDF, that is,

$$\begin{aligned} R(k, n, p) &= \sum_{i=k}^n \binom{n}{i} p^i (1 - p)^{n-i} \\ &= 1 - P(X \leq k - 1) \end{aligned} \quad (4)$$

The reliability of the system changes as the parameters n , p , and k change. If two of the parameters are fixed and the remaining parameter changes, the reliability of the system changes as follows:

- If p and k remain fixed, the reliability of the system increases as n increases.
- If p and n remain fixed, the reliability of the system decreases as k increases.
- If n and k remain fixed, the reliability of the system increases as p increases.

The changes in the reliability of the system can be shown for each case described above. The increase in reliability of the system for a unit increase in n can be found by calculating the difference

2 k -out-of- n Systems

$R(k, n, p) - R(k, n - 1, p)$. To derive an expression for the difference, Kuo and Zuo [1] show that pivotal decomposition to component n may be used as follows:

$$\begin{aligned}
 R(k, n, p) &= R(k - 1, n - 1, p)p + R(k, n - 1, p)(1 - p) \\
 &= P(\text{at least } k - 1 \text{ components work out of } n - 1 \text{ components}) P(\text{the } n\text{th component works}) + P(\text{at least } k \text{ components work out of } n - 1 \text{ components}) P(\text{the } n\text{th component fails}) \\
 &= \sum_{i=k-1}^{n-1} \left\{ \binom{n-1}{i} p^i (1-p)^{(n-1)-i} \right\} p \\
 &\quad + \sum_{i=k}^{n-1} \left\{ \binom{n-1}{i} p^i (1-p)^{(n-1)-i} \right\} (1-p) \\
 &= \sum_{i=k-1}^{n-1} \left\{ \binom{n-1}{i} p^i (1-p)^{(n-1)-i} \right\} p \\
 &\quad - \sum_{i=k}^{n-1} \left\{ \binom{n-1}{i} p^i (1-p)^{(n-1)-i} \right\} p \\
 &\quad + \sum_{i=k}^{n-1} \left\{ \binom{n-1}{i} p^i (1-p)^{(n-1)-i} \right\} \\
 &= \binom{n-1}{k-1} p^k (1-p)^{n-k} \\
 &\quad + \sum_{i=k}^{n-1} \left\{ \binom{n-1}{i} p^i (1-p)^{(n-1)-i} \right\} p \\
 &\quad - \sum_{i=k}^{n-1} \left\{ \binom{n-1}{i} p^i (1-p)^{(n-1)-i} \right\} p \\
 &\quad + \sum_{i=k}^{n-1} \left\{ \binom{n-1}{i} p^i (1-p)^{(n-1)-i} \right\} \quad (5)
 \end{aligned}$$

The second and third terms cancel and the reliability of the system may be written as,

$$\begin{aligned}
 R(k, n, p) &= \binom{n-1}{k-1} p^k (1-p)^{n-k} \\
 &\quad + \sum_{i=k}^{n-1} \left\{ \binom{n-1}{i} p^i (1-p)^{(n-1)-i} \right\}
 \end{aligned}$$

$$= \binom{n-1}{k-1} p^k (1-p)^{n-k} + R(k, n-1, p) \quad (6)$$

Solving for the difference, $R(k, n, p) - R(k, n - 1, p)$, we see that the increase in the system reliability for an increase of one component is given by

$$\begin{aligned}
 R(k, n, p) - R(k, n - 1, p) &= \binom{n-1}{k-1} p^k (1-p)^{n-k}, \quad n \geq k \quad (7)
 \end{aligned}$$

Notice the difference is positive, meaning that the reliability of the system does increase as n increases, but as n increases, the *amount of improvement* in system reliability decreases.

As the parameter k decreases, the system reliability will increase. This can be shown by computing the difference in a unit decrease in k , $R(k - 1, n, p) - R(k, n, p)$.

$$\begin{aligned}
 R(k - 1, n, p) - R(k, n, p) &= P(\text{at least } k - 1 \text{ components work out of } n) \\
 &\quad - P(\text{at least } k \text{ components work out of } n) \\
 &= P(\text{exactly } k - 1 \text{ components work out of } n) \\
 &= \binom{n}{k-1} p^{k-1} (1-p)^{n-k+1} \quad (8)
 \end{aligned}$$

Therefore, for a unit decrease in k , the reliability of the system increases since the difference $R(k - 1, n, p) - R(k, n, p)$ is positive.

Since the parameter p is continuous, we can take the derivative of the reliability function with respect to p to find how the system reliability changes with respect to the **component reliability** p if the components are i.i.d.

$$\begin{aligned}
 \frac{d}{dp} R(k, n, p) &= \frac{d}{dp} \sum_{i=k}^n \binom{n}{i} p^i (1-p)^{n-i} \\
 &= \frac{d}{dp} \left\{ \binom{n}{k} p^k (1-p)^{n-k} \right. \\
 &\quad + \binom{n}{k+1} p^{k+1} (1-p)^{n-k-1} \\
 &\quad + \binom{n}{k+2} p^{k+2} (1-p)^{n-k-2}
 \end{aligned}$$

$$\begin{aligned}
& + \cdots + \binom{n}{n-1} p^{n-1} (1-p)^{n-(n-1)} \\
& + \binom{n}{n} p^n (1-p)^{n-n} \Big\} \\
= & \binom{n}{k} (-p^k (n-k) (1-p)^{n-k-1} \\
& + kp^{k-1} (1-p)^{n-k}) \\
& + \binom{n}{k+1} (-p^{k+1} (n-k-1) (1-p)^{n-k-2} \\
& + (k+1)p^k (1-p)^{n-k-1}) \\
& + \binom{n}{k+2} (-p^{k+2} (n-k-2) (1-p)^{n-k-3} \\
& + (k+2)p^{k+1} (1-p)^{n-k-2}) + \cdots + \binom{n}{n-1} \\
& \times (-p^{n-1} (n-(n-1)) (1-p)^{n-n} \\
& + (n-1)p^{n-2} (1-p)^{n-(n-1)}) + \binom{n}{n} (np^{n-1}) \\
= & k \binom{n}{k} p^{k-1} (1-p)^{n-k} \\
& + \left\{ (k+1) \binom{n}{k+1} - \binom{n}{k} (n-k) \right\} \\
& \times p^k (1-p)^{n-k-1} \\
& + \left\{ (k+2) \binom{n}{k+2} - \binom{n}{k+1} (n-k-1) \right\} \\
& \times p^{k+1} (1-p)^{n-k-2} \\
& - \binom{n}{k+2} (n-k-2) p^{k+2} (1-p)^{n-k-3} \\
& + \cdots + (n-1) \binom{n}{n-1} p^{n-2} (1-p)^{n-(n-1)} \\
& + \left\{ n \binom{n}{n} - \binom{n}{n-1} (n-(n-1)) \right\} \\
& \times p^{n-1} (1-p)^{n-n} \tag{9}
\end{aligned}$$

Since any expression of the form $(k+1) \binom{n}{k+1}$ is equivalent to $(n-k) \binom{n}{k}$, the coefficients of $p^i (1-p)^{n-i-1}$ for $i = k, k+1, \dots, n-1$ are zero.

That is,

$$\begin{aligned}
& (k+1) \binom{n}{k+1} - (n-k) \binom{n}{k} \\
& = (k+1) \frac{n!}{(n-k-1)!(k+1)!} \\
& \quad - (n-k) \frac{n!}{(n-k)!(k)!} \\
& = \frac{n!}{(n-k-1)!(k)!} - \frac{n!}{(n-k-1)!(k)!} \\
& = 0 \tag{10}
\end{aligned}$$

Therefore, the derivative of the reliability function with respect to p is given as

$$\frac{dR(k, n, p)}{dp} = k \binom{n}{k} p^{k-1} (1-p)^{n-k} \tag{11}$$

Note that the above derivative is positive meaning that the reliability function is increasing with respect to the parameter p .

Example 1 Evaluate the performance of a 2-out-of-5:G system with $p = 0.5$.

Solution:

$$\begin{aligned}
R(2, 5, 0.5) &= \sum_{i=2}^5 \binom{5}{i} 0.5^i (1-0.5)^{5-i} \\
&= P(X \geq 2) \\
&= 1 - P(X < 2) \\
&= 1 - P(X \leq 1) \\
&= 0.8125
\end{aligned}$$

Reliability Computations Assuming Exponential Distributions of Lifetimes

In the above section, the time to failure of a component was not considered in the calculation of the reliability of the system. Systems of any kind are either repairable or nonrepairable. In either case, the time until failure or repair must be considered. The time, T_i , until the failure of the i th component will now be included in the system reliability calculations. The time until failure of each of the components is a random variable. If the time until failure of all

4 k -out-of- n Systems

of the components is identically and independently distributed, the reliability function of the system may be expressed as in the first section with p replaced by $R(t)$, the reliability of each component. The reliability of the system, $R_s(t)$, may be expressed as

$$R_s(t) = \sum_{i=k}^n \binom{n}{i} R(t)^i (1 - R(t))^{n-i} \quad (12)$$

$R_s(t)$ may be thought of as the probability of at least k units working until time t . Similarly, $F_s(t) = 1 - R_s(t)$ is the probability of $k - 1$ or fewer units working until time t , where $F_s(t)$ is the cumulative distribution function of the system lifetime. $F_s(t)$ may be referred to as the *unreliability of the system*.

One of the main interests involving the system reliability is the mean time to failure of the system, denoted by $MTTF_s$. It can be shown, in general, that the expected value of a random variable defined on the interval $[0, \infty)$, is equivalent to the integral of the survival function of that random variable over the interval $[0, \infty)$. That is,

$$\begin{aligned} E(X) &= \int_x^\infty xf(x) dx = \int_x^\infty S(x) dx \\ &= \int_x^\infty (1 - F(x)) dx \end{aligned} \quad (13)$$

The mean time to failure of the system, $MTTF_s$, may be found by integrating the reliability of the system from zero to infinity, that is,

$$\int_0^\infty R_s(t) dt = MTTF_s \quad (14)$$

When the lifetimes of the components are i.i.d. and follow an exponential distribution, the $MTTF_s$ of a k -out-of- n :G system may be found explicitly. If the time until failure of each component is exponentially distributed with mean $1/\lambda$, then the reliability of each component is $R(t) = e^{-\lambda t}$ and the CDF is $F(t) = 1 - e^{-\lambda t}$. Since the time until failure for each component is independently and identically distributed, the probability that at least k components work may be expressed as

$$\begin{aligned} R_s(t) &= P(\text{exactly } k \text{ components work out of } n) \\ &\quad + P(\text{exactly } k + 1 \text{ components work out of } n) \\ &\quad + \cdots + P(\text{exactly } n \text{ components work out of } n) \end{aligned}$$

$$\begin{aligned} &= \binom{n}{k} (e^{-\lambda t})^k (1 - e^{-\lambda t})^{n-k} \\ &\quad + \binom{n}{k+1} (e^{-\lambda t})^{k+1} (1 - e^{-\lambda t})^{n-k-1} \\ &\quad + \cdots + \binom{n}{n} (e^{-\lambda t})^n (1 - e^{-\lambda t})^{n-n} \\ &= \sum_{i=k}^n \binom{n}{i} (e^{-\lambda t})^i (1 - e^{-\lambda t})^{n-i} \end{aligned} \quad (15)$$

Kuo and Zuo [1] show that if $k \geq 2$ the reliability of the system may be written recursively as

$$\begin{aligned} R_s(t; k, n) &= P(\text{at least } k \text{ components work out of } n) \\ &= P(\text{at least } k - 1 \text{ components work out of } n) \\ &\quad - P(\text{exactly } k - 1 \text{ components work out of } n) \\ &= R_s(t; k - 1, n) - \binom{n}{k-1} \\ &\quad \times (e^{-\lambda t})^{k-1} (1 - e^{-\lambda t})^{n-k+1} \end{aligned} \quad (16)$$

For $k \geq 2$, the mean time to failure of the system, $MTTF_s$, may be found by integrating both sides of equation (16) with respect to t as follows:

$$\begin{aligned} MTTF_s(k, n) &= MTTF_s(k - 1, n) - \binom{n}{k-1} \\ &\quad \times \int_0^\infty (e^{-\lambda t})^{k-1} (1 - e^{-\lambda t})^{n-k+1} dt \end{aligned} \quad (17)$$

The last term may be evaluated using integration by substitution and recognizing that the resulting integrand is the β **probability density function** divided by the constant $(\Gamma(\alpha)\Gamma(\beta))/\Gamma(\alpha + \beta)$, where $\alpha = n - k + 2$ and $\beta = k - 1$. That is, let $u = 1 - e^{-\lambda t}$, then $du = \lambda e^{-\lambda t} dt = \lambda(1 - u) dt$. Thus,

$$\begin{aligned} &\int_0^\infty \binom{n}{k-1} \frac{1}{\lambda} (e^{-\lambda t})^{k-1} (1 - e^{-\lambda t})^{n-k+1} dt \\ &= \binom{n}{k-1} \frac{1}{\lambda} \int_0^1 (1 - u)^{-1} (1 - u)^{k-1} (u)^{n-k+1} du \end{aligned}$$

$$\begin{aligned}
 &= \binom{n}{k-1} \frac{1}{\lambda} \frac{\Gamma(n-k+2) \Gamma(k-1)}{\Gamma(n-k+2+k-1)} \int_0^1 \\
 &\quad \frac{\Gamma(n-k+2+k-1)}{\Gamma(n-k+2)\Gamma(k-1)} (1-u)^{k-2} (u)^{n-k+1} du \\
 &= \binom{n}{k-1} \frac{1}{\lambda} \frac{\Gamma(n-k+2)\Gamma(k-1)}{\Gamma(n-k+2+k-1)} (1) \\
 &= \frac{n!}{(k-1)!(n-k+1)! \lambda} \frac{1}{\lambda} \frac{(n-k+1)!(k-2)!}{n!} \\
 &= \frac{(k-2)!}{\lambda(k-1)!} = \frac{1}{\lambda(k-1)} \quad (18)
 \end{aligned}$$

and

$$MTTF_s(k, n) = MTTF_s(k-1, n) - \frac{1}{\lambda(k-1)} \quad (19)$$

If $k = 1$, then

$$\begin{aligned}
 R_s(t; 1, n) &= P \left(\begin{array}{l} \text{at least 0 components} \\ \text{work out of } n \text{ components} \end{array} \right) \\
 &\quad - P \left(\begin{array}{l} \text{exactly 0 components} \\ \text{work out of } n \text{ components} \end{array} \right) \\
 &= 1 - \binom{n}{0} (e^{-\lambda t})^0 (1 - e^{-\lambda t})^{n-0} \\
 &= 1 - (1 - e^{-\lambda t})^n \quad (20)
 \end{aligned}$$

Expanding the last term, we see that the reliability of the system may be written as

$$\begin{aligned}
 R_s(t; 1, n) &= 1 - (1 - e^{-\lambda t})^n \\
 &= 1 - \left\{ \binom{n}{0} (-e^{-\lambda t})^n + \binom{n}{1} (-e^{-\lambda t})^{n-1} \right. \\
 &\quad + \binom{n}{2} (-e^{-\lambda t})^{n-2} + \cdots + \binom{n}{n-1} \\
 &\quad \times (-e^{-\lambda t})^1 + \left. \binom{n}{n} (-e^{-\lambda t})^0 \right\} \\
 &= - \left\{ \binom{n}{0} (-e^{-\lambda t})^n + \binom{n}{1} (-e^{-\lambda t})^{n-1} \right. \\
 &\quad + \binom{n}{2} (-e^{-\lambda t})^{n-2} \\
 &\quad + \cdots + \left. \binom{n}{n-1} (-e^{-\lambda t})^1 \right\} \quad (21)
 \end{aligned}$$

Integrating with respect to t from zero to infinity to calculate the mean time to failure, $MTTF_s$, gives

$$\begin{aligned}
 &MTTF_s(t; 1, n) \\
 &= \int_0^\infty R_s(t; 1, n) dt \\
 &= - \int_0^\infty \left\{ \binom{n}{0} (-e^{-\lambda t})^n + \binom{n}{1} (-e^{-\lambda t})^{n-1} \right. \\
 &\quad + \binom{n}{2} (-e^{-\lambda t})^{n-2} + \cdots + \binom{n}{n-2} (-e^{-\lambda t})^2 \\
 &\quad + \left. \binom{n}{n-1} (-e^{-\lambda t})^1 \right\} dt \\
 &= \frac{1}{\lambda} \left\{ \binom{n}{0} (-1)^n \frac{(e^{-\lambda t})^n}{n} \right. \\
 &\quad + \binom{n}{1} (-1)^{n-1} \frac{(e^{-\lambda t})^{n-1}}{n-1} + \binom{n}{2} (-1)^{n-2} \\
 &\quad \times \frac{(e^{-\lambda t})^{n-2}}{n-2} + \cdots + \binom{n}{n-2} (-1)^2 (e^{-\lambda t})^2 \\
 &\quad + \left. \binom{n}{n-1} (-1)^1 (e^{-\lambda t})^1 \right\} \Big|_0^\infty \\
 &= -\frac{1}{\lambda} \left\{ \binom{n}{0} (-1)^{n-0} \frac{1}{n} + \binom{n}{1} (-1)^{n-1} \frac{1}{(n-1)} \right. \\
 &\quad + \binom{n}{2} (-1)^{n-2} \frac{1}{(n-2)} \\
 &\quad + \cdots + \left. \binom{n}{n-2} (-1)^2 \frac{1}{2} + \binom{n}{n-1} (-1)^1 \right\} \\
 &= \frac{1}{\lambda} \left\{ \binom{n}{0} (-1)^{n-1} \frac{1}{n} + \binom{n}{1} (-1)^{n-1-1} \frac{1}{(n-1)} \right. \\
 &\quad + \binom{n}{2} (-1)^{n-2-1} \frac{1}{(n-2)} + \cdots + \\
 &\quad \times \left. \binom{n}{n-2} (-1)^{2-1} \frac{1}{2} + \binom{n}{n-1} (-1)^{1-1} \right\} \\
 &= \frac{1}{\lambda} \left\{ \binom{n}{n} (-1)^{n-1} \frac{1}{n} + \binom{n}{n-1} (-1)^{n-1-1} \right. \\
 &\quad \times \left. \frac{1}{(n-1)} + \binom{n}{n-2} (-1)^{n-2-1} \frac{1}{(n-2)} \right.
 \end{aligned}$$

$$\begin{aligned}
& + \cdots + \binom{n}{2} (-1)^{2-1} \frac{1}{2} + \binom{n}{1} (-1)^{1-1} \Big\} \\
& = \frac{1}{\lambda} \Big\{ \binom{n}{1} (-1)^{1-1} + \binom{n}{2} (-1)^{2-1} \frac{1}{2} \\
& \quad + \cdots + \binom{n}{n-2} (-1)^{n-2-1} \\
& \quad \times \frac{1}{(n-2)} + \binom{n}{n-1} (-1)^{n-1-1} \frac{1}{(n-1)} \\
& \quad + \binom{n}{n} (-1)^{n-1} \frac{1}{n} \Big\} \\
& = \frac{1}{\lambda} \sum_{i=1}^n \binom{n}{i} (-1)^{i-1} \frac{1}{i} \tag{22}
\end{aligned}$$

The sum $\sum_{i=1}^n \binom{n}{i} (-1)^{i-1} (1/i)$ in the above calculation is known as the *binomial harmonic identity* and is equivalent to $\sum_{i=1}^n (1/i)$, therefore, the mean time to failure of the system for $k = 1$ is

$$MTTF_s(1, n) = \frac{1}{\lambda} \sum_{i=1}^n \frac{1}{i} \tag{23}$$

Using the above result for $MTTF_s(1, n)$ and the recursive formula given for $MTTF_s(k, n)$, for $k \geq 2$, in equation (19), a closed-form expression for $MTTF_s(k, n)$ with a system of i.i.d. components with exponential distributions to lifetimes can easily be shown to be

$$MTTF_s(k, n) = \frac{1}{\lambda} \sum_{i=k}^n \frac{1}{i} \tag{24}$$

Example 2 Find the reliability and the $MTTF_s$ of a 2-out-of-4:G system in which the components have identical exponential lifetimes with mean 3.

Solution: The reliability of the system is given by

$$\begin{aligned}
R_s(t) &= \sum_{i=2}^4 \binom{4}{i} (e^{-(1/3)t})^i (1 - e^{-(1/3)t})^{n-i} \\
&= 6e^{-(2/3)t} (1 - e^{-(1/3)t})^2 \\
&\quad + 4e^{-t} (1 - e^{-(1/3)t}) + e^{-(4/3)t}
\end{aligned}$$

The $MTTF_s$ is

$$\begin{aligned}
MTTF_s &= 3 \sum_{i=2}^4 \frac{1}{i} \\
&= 3 \left(\frac{1}{2} + \frac{1}{3} + \frac{1}{4} \right) \\
&= \frac{13}{4}
\end{aligned}$$

Consecutive *k-out-of-n* Systems

Here we are talking about a *k-out-of-n* system that must have at least *k* consecutive components that fail before the system fails or at least *k* consecutive components that work for the system to continue to work. Only the circular and linear systems will be considered here. If the components are connected in a circular fashion, such as satellite communication networks, then the system is called a *circular k-out-of-n system*. Similarly, if the system is set up with components connected linearly, then the system is called a *linear k-out-of-n system*. The system will fail if at least *k* consecutive components of the *n* components fail. An example of this type of system is a communication relay system in which the signal is received and relayed at each station. A common shorthand method to refer to the *linear* and *circular* F and G systems are *Lin/Con/k/n:F*, *Lin/Con/k/n:G*, *Cir/Con/k/n:F*, and *Cir/Con/k/n:G*, which refer to linear consecutive *k-out-of-n:F* system, linear consecutive *k-out-of-n:G* system, circular consecutive *k-out-of-n:F* system, and circular consecutive *k-out-of-n:G* system, respectively. As in the previous sections, the components can only occupy one of the two states, either working or not working (failed).

Reliability Computations for Consecutive *k-out-of-n* System

In this section, the reliability computations will be calculated for systems with i.i.d. components. The linear consecutive *k-out-of-n:F* system will be considered first. There are several methods of computing the reliability of a *Lin/Con/k/n:F* system. The challenge in computing the reliability of this type of system is in counting the ways in which *j*, failed

components, can be arranged out of the n components so that there are at least k failed consecutive components in the system. First, let j represent the number of components that do not work so that $n - j$ represents the number of working components. If a combinatorial approach is used and the system is a *Lin/Con/2/n:F*, that is, $k = 2$, and the system works if at least every other component is working, it is easy to show that the reliability of the system is

$$\begin{aligned}
 R_{\text{linear}}(2, n) &= \binom{n+1}{0} (1-p)^0 p^{n+1} \\
 &\quad + \binom{n}{1} (1-p)^1 p^n \\
 &\quad + \cdots + \binom{n - \left\lfloor \frac{n+1}{2} \right\rfloor + 1}{\left\lfloor \frac{n+1}{2} \right\rfloor} \\
 &\quad \times (1-p)^{\left\lfloor \frac{n+1}{2} \right\rfloor} p^{n - \left\lfloor \frac{n+1}{2} \right\rfloor} \\
 &= \sum_{i=1}^{\lfloor (n+1)/2 \rfloor} \binom{n-i+1}{i} p^{n-i} (1-p)^i \\
 &= R_{\text{linear}}(2, n) = \sum_{i=1}^{\lfloor (n+1)/2 \rfloor} \binom{n-i+1}{i} p^{n-i} (1-p)^i \quad (25)
 \end{aligned}$$

where $p = P(\text{the component works})$ and $1 - p = P(\text{the component fails})$.

To calculate the reliability for a linear system for any k , $R_{\text{linear}}(k, n)$, it is necessary to find the number of ways to arrange the failed components so that there are no more than k consecutive failed components. A common symbol for this number is $N(j, k, n)$ and $R_{\text{linear}}(k, n)$ may be written as

$$R_{\text{linear}}(k, n) = \sum_{i=1}^n N(i, k, n) p^{n-i} (1-p)^i \quad (26)$$

Notice that some of the terms in equation (26) are zero because at some point as i approaches n there will be more than k consecutive failed components, meaning that the entire system would have already failed. There are several approaches to calculating an expression for $N(i, k, n)$. Bollinger [2]

developed a method to construct a table to calculate $N(i, k, n)$. Derman *et al.* [3] developed a recursive formula to calculate $N(i, k, n)$. To use equation (26) to find the reliability of the *Lin/Con/k/n:F* system, Satam [4] derived an expression for the maximum number of terms in the sum $R_{\text{linear}}(k, n)$ that are not zero, in other words, the maximum value of the index i . Let $\max$ represent the largest number of nonzero terms in the sum; then $\max$ is given by

$$\begin{aligned}
 \max &= \begin{cases} n - \frac{n+k}{k} + 1 & \text{if } n+1 \text{ is a multiple of } k \\ n - \left\lfloor \frac{n+1}{k} \right\rfloor & \text{if } n+1 \text{ is not a multiple of } k \end{cases} \quad (27)
 \end{aligned}$$

The reliability of the system may then be written as

$$R_{\text{linear}}(k, n) = \sum_{i=1}^{\max} N(i, k, n) p^{n-i} (1-p)^i \quad (28)$$

Two additional closed-form expressions of the reliability for the *Lin/Con/k/n:F* system are given by Goulden [5] and Hwang [6] as,

$$\begin{aligned}
 R_{\text{linear}}(k, n) &= \sum_{i=0}^{\lfloor n/(k+1) \rfloor} (-1)^i p^i (1-p)^{ki} \left[\binom{n-ki}{i} \right. \\
 &\quad \left. - (1-p)^k \binom{n-k(i+1)}{i} \right] \quad (29)
 \end{aligned}$$

and

$$\begin{aligned}
 R_{\text{linear}}(k, n) &= \sum_{i=0}^{\lfloor (n+1)/(k+1) \rfloor} (-1)^i p^{i-1} (1-p)^{ki} \\
 &\quad \times \left[\binom{n-ki+1}{i} - (1-p) \binom{n-ki}{i} \right] \quad (30)
 \end{aligned}$$

Example 3 Find the reliability of a *Lin/Con/k/n:F* system with $n = 15$, $k = 4$, and $p = 0.8$.

8 k -out-of- n Systems

Solution: Using the closed-form expression by Goulden, the reliability is calculated by the sum:

$$\begin{aligned}
 R_{\text{linear}}(3, 11) &= \sum_{i=0}^{\lfloor 15/5 \rfloor} (-1)^i (1 - 0.8)^{4i} (0.8)^i \left[\binom{15-3i}{i} \right. \\
 &\quad \left. - (1 - 0.8)^3 \binom{15-3(i+1)}{i} \right] \\
 &= (-1)^0 (0.2)^0 (0.8)^0 \left[\binom{15}{0} - (0.2)^3 \binom{12}{0} \right] \\
 &\quad + (-1)^1 (0.2)^4 (0.8)^1 \left[\binom{12}{1} - (0.2)^3 \binom{9}{1} \right] \\
 &\quad + (-1)^2 (0.2)^8 (0.8)^2 \left[\binom{9}{2} - (0.2)^3 \binom{6}{2} \right] \\
 &\quad + (-1)^3 (0.2)^{12} (0.8)^3 \left[\binom{6}{3} - (0.2)^3 \binom{3}{3} \right] \\
 &\approx 0.9768
 \end{aligned}$$

Circular Consecutive k -out-of- n :F System

The reliability of the $\text{Cir}/\text{Con}/k/n$:F system may be expressed in the same way as that of the $\text{Lin}/\text{Con}/k/n$:F system by

$$R_{\text{circular}}(k, n) = \sum_{i=0}^{\max} N_{\text{circular}}(i, k, n) p^{n-i} (1-p)^i \quad (31)$$

where, p is the probability that the component works, $(1-p)$ is the probability that the component fails, $N_{\text{circular}}(i, k, n)$ is the number of ways to obtain i failed components arranged in a circle such that no more than k failed components are consecutive and $\max$ is the maximum number of failed components that can occur until the system fails. The value of $\max$ may be calculated by

$$\max = \begin{cases} n - \frac{n}{k} & \text{if } n \text{ is a multiple of } k \\ n - \left\lfloor \frac{n}{k} \right\rfloor - 1 & \text{if } n \text{ is not a multiple of } k \end{cases} \quad (32)$$

Other equations to evaluate the reliability of a $\text{Cir}/\text{Con}/k/n$:F system are given respectively by

Goulden [5] and Lambiris and Papastavridis [7] as

$$\begin{aligned}
 R_{\text{circular}}(k, n) &= 1 - q^n + n \sum_{i=1}^{\lfloor n/(k+1) \rfloor} \frac{(-1)^i}{i} p^i q^{ki} \\
 &\quad \times \binom{n-ik-1}{i-1}, \quad n \geq k \quad (33)
 \end{aligned}$$

$$\begin{aligned}
 R_{\text{circular}}(k, n) &= R_{\text{circular}}(k, n-1) - (pq^k) \\
 &\quad \times R_{\text{circular}}(k, n-k-1), \quad n \geq 2k+1 \quad (34)
 \end{aligned}$$

Equation (34) is the most efficient way to evaluate the reliability of the $\text{Cir}/\text{Con}/k/n$:F system with a computational complexity of $O(n)$.

The $\text{Cir}/\text{Con}/k/n$:F system may be treated as a special case of the reliability function of the $\text{Lin}/\text{Con}/k/n$:F system. Derman *et al.* [3] developed the following recursive relationship for $R_{\text{circular}}(k, n)$.

$$\begin{aligned}
 R_{\text{circular}}(k, n) &= \sum_{i=0}^{k-1} (i+1) p^2 q^i R_{\text{linear}}(k, n-i-2) \\
 &= \sum_{i=0}^{k-1} (i+1) P(N+N'=i) \\
 &\quad P(\text{at least } k \text{ consecutive components} \\
 &\quad \text{work out of } n-i-2) \quad (35)
 \end{aligned}$$

N and N' are computed as follows. Choose a component in the circular system. Arrange the components linearly with the chosen component as the first component in a linear arrangement. Let N be the number of failed components until the first working component is found starting from the first component. Next, arrange the components linearly, but now let the chosen component be the last component and let N' be the number of failed components until the first working component is found starting the count from the last component. Notice, $P(N=i) = P(N'=i) = p(1-p)^i$ for $i = 1, 2, \dots, n-1$ and if $i \leq n-2$,

then

$$\begin{aligned}
 P(N + N' = i) &= \sum_{j=0}^i P(N = j)P(N = i - j) \\
 &= \sum_{j=0}^i p(1 - p)^j p(1 - p)^{i-j} \\
 &= p^2 q^i \quad \text{for } i = 0, 1, 2, \dots, n - 2
 \end{aligned} \tag{36}$$

Example 4 Find the reliability of a *Cir/Con/k/n:F* system with $n = 15$, $k = 4$, and $p = 0.8$.

Solution: Using the expression (15) by Goulden the reliability is calculated by

$$\begin{aligned}
 R_{\text{circular}}(k, n) \\
 = 1 - q^n + n \sum_{i=1}^{\lfloor n/(k+1) \rfloor} \frac{(-1)^i}{i} p^i q^{ki} \binom{n - ik - 1}{i - 1}
 \end{aligned} \tag{37}$$

$$\begin{aligned}
 R_{\text{circular}}(4, 15) \\
 = 1 - 0.2^{15} + 15 \sum_{i=1}^3 \frac{(-1)^i}{i} (0.8)^i (0.2)^{4i} \\
 \times \binom{15 - 4i - 1}{i - 1}
 \end{aligned}$$

$$\begin{aligned}
 &= 1 - 0.2^{15} + 15 \left\{ - (0.8)^1 (0.2)^4 \binom{10}{0} \right. \\
 &\quad \left. + \frac{1}{2} (0.8)^2 (0.2)^8 \binom{6}{1} - \frac{1}{3} (0.8)^3 (0.2)^{12} \binom{2}{2} \right\} \\
 &= 0.9809
 \end{aligned} \tag{38}$$

References

- [1] Kuo, W. & Zuo, M. (2003). *Optimal Reliability Modeling Principles and Applications*, John Wiley & Sons, Hoboken, p. 223.
- [2] Bollinger, R.C. (1982). Direct computation for consecutive-*k*-out-of-*n*:F systems, *IEEE Transactions on Reliability* **R-31**(5), 444–446.
- [3] Derman, C., Lieberman, G.J. & Ross, S.M. (1982). On the consecutive-*k*-out-of-*n*:F system, *IEEE Transactions on Reliability* **R-31**(1), 57–63.
- [4] Satam, M.R. (1991). Consecutive-*k*-out-of-*n*:F systems: a comment, *IEEE Transactions on Reliability* **R-40**(1), 62.
- [5] Goulden, I.P. (1987). Generating functions and reliabilities for consecutive *k*-out-of-*n*:F systems, *Utilitas Mathematica* **32**(1), 141–147.
- [6] Hwang, F.K. (1986). Simplified reliabilities for consecutive-*k*-out-of-*n*:F system, *SIAM Journal on Algebraic and Discrete Methods* **7**(2), 258–264.
- [7] Lambiris, M. & Papastavridis, S.G. (1985). Exact reliability formulas for linear and circular consecutive-*k*-out-of-*n*:F systems, *IEEE Transactions on Reliability* **R-34**(2), 124–126.

DEBORAH K. SHEPHERD

Path Sets and Cut Sets in System Reliability Modeling

Background and Preliminaries

Graphs

Definitions. A *graph* is an ordered pair of sets $(\mathcal{N}, \mathcal{L}) = \mathcal{G}$ with $\mathcal{L} \subset \mathcal{N} \times \mathcal{N}$. $\mathcal{N}$ is called the *set of nodes* of the graph and $\mathcal{L}$ is called the *set of links* of the graph. Typically, a graph is identified with a drawing in which the nodes are represented as points in the plane and the links are represented as lines drawn to join two nodes. In other terminology in common use, the nodes may be called *vertices* and the links may be called *arcs* or *edges*.

A *labeled graph* is a graph in which the nodes and/or links have names. That is, there is a one-to-one correspondence between the nodes of the graph and a set of $|\mathcal{N}|$ objects (the node labels) and/or between the links of the graph and a set of $|\mathcal{L}|$ objects (the link labels).

A *directed graph* is a graph in which each link is assigned an orientation or direction. In a directed graph, the links (i, j) and (j, i) are different, whereas in an ordinary (undirected) graph, they are identical. The concept of “a link from i to j ” makes sense in a directed graph; in an undirected graph it would be proper to say, rather, “a link between i and j ”.

Examples. The cities in a state, together with the roads joining them, may be described as a graph. In this example, the cities form the nodes of the graph and the roads joining them form the links. A natural labeling of the nodes and links exists. A natural gas pipeline may be modeled as a graph with the terminals and transfer stations as the nodes and the pipes as the links. Graphs find application in the social sciences also. A graph may be constructed from a set of individuals (who form the nodes) in which two individuals are joined by a link if they are known to each other. Indeed, any binary reflexive relation between discrete entities gives rise to a graph in which the entities are the nodes and a link joins two nodes if the relation holds for those two nodes. If the relation is not reflexive, then the result is a

directed graph. If there is a notion of “strength” of the relationship and we wish to incorporate this notion into the model, a more appropriate approach is to model the graph as a capacitated network; *see* the section titled “Paths, Cuts, and Demand Satisfaction in Capacitated Networks”.

References. Other useful graph concepts, such as degree, are not covered in this article because they are not relevant to the reliability modeling application. For a review of the basic facts of graph theory, see [1–3].

Paths and Connectedness

Definitions. Two nodes i and j are *adjacent* if $(i, j) \in \mathcal{L}$. The *adjacency matrix* of $\mathcal{G}$ is an $|\mathcal{N}| \times |\mathcal{N}|$ matrix whose ij entry is one if $(i, j) \in \mathcal{L}$ and is zero otherwise. In other terminology in common use, the adjacency matrix may be called the *incidence matrix*. A *path* in a graph is a sequence of alternating adjacent nodes and the links joining them, beginning and ending with a node. Two nodes i and j are said to be *connected* if there is a path having i as its initial node and j as its terminal node. That is, the path takes the form $\{i, (i, v_1), v_1, (v_1, v_2), \dots, v_k, (v_k, j), j\}$ for some $v_1, \dots, v_k \in \mathcal{N}$ and $(i, v_1), (v_1, v_2), \dots, (v_k, j) \in \mathcal{L}$. There is no loss in abbreviating this to $\{(i, v_1), (v_1, v_2), \dots, (v_k, j)\}$. When it is necessary or desirable to explicitly indicate the nodes being connected, the path will be called an (i, j) *path*. Clearly, adjacent nodes are connected but connected nodes need not be adjacent. If the graph is directed, the links in the path must be considered with the proper orientation.

A *cut* for two given nodes is a set of nodes and/or links whose removal from the graph disconnects the two nodes. When it is necessary or desirable to explicitly indicate the nodes being disconnected, the cut will be called an (i, j) *cut*.

A path connecting two given nodes is called *minimal* if it contains no proper subset that is also a path connecting the same two nodes. A cut for two given nodes is called *minimal* if it contains no proper subset that is also a cut disconnecting those nodes.

Examples. The Washington, DC, Metro subway system <http://www.wmata.com/metro/rail/systemmap.cfm> may be modeled as a graph with the stations as

the nodes. In this graph, the Pentagon and College Park – University of Maryland are connected but not adjacent. DuPont Circle and Farragut North are both connected and adjacent. The Manhattan (NY) city streets may be modeled as a graph with the intersections as the nodes. For vehicle traffic, this graph is directed because many streets in Manhattan are one-way. For pedestrian traffic, the graph need not be directed.

Random Graphs

Definitions. A *random graph* is a graph in which a stochastic indicator variable is attached to each node and link. When the variable is zero, it indicates that, that node or link is not present in the graph. When it is one, it indicates that, that node or link is present in the graph. Each choice of values for these indicator variables, by whatever random mechanism is at play, produces a different graph (the choice is not completely unrestricted; if the indicator of a node is zero, the indicators of all the links emanating from that node must be zero also). The adjacency matrix of a random graph is a random matrix.

In a reliability engineering application, the indicator variable for a link or node describes the functioning or nonfunctioning of the link or node. The usual convention is that the indicator variable is one when the link or node functions and zero if it does not function. There is no fundamental reason why the opposite assignment could not be used, and it is sometimes seen.

References. For purposes of the reliability modeling described in this article, this definition is as much of the subject of random graphs as we will use. Additional material on random graphs may be found in [4] and [5].

Paths, Cuts, and System Reliability Modeling

Introduction: Structure Functions and Reliability Block Diagrams

For the purposes of this article, reliability is considered an all-or-nothing condition: either a unit or system functions properly and completely at a stated time, or it does not function at all at that time. Thus,

we may imagine the functioning-or-failed condition of a unit or system as a zero-one, or indicator, random variable that, in the general case, will be allowed to depend on time. The system *structure function* is a Boolean function that maps $\{0, 1\}^c$ into $\{0, 1\}$, where the number of components in the system is c . The structure function is the indicator variable of the system functioning when the arguments in the structure function are replaced by the indicator variables of the components' functioning. Let X_k indicate the functioning of component k and $x_k \in \{0, 1\} (k = 1, \dots, c)$. The system structure function is $\varphi(x_1, \dots, x_c)$ and the system reliability is $P\{\varphi(X_1, \dots, X_c) = 1\}$. For example, the structure function for a series system of c components is $\varphi(x_1, \dots, x_c) = x_1 x_2 \cdots x_c$.

A structure function is called *coherent* if it is nondecreasing in each variable separately and each component is relevant (that is, the state of the system depends on the states of all the components in the structure function).

The system reliability block diagram is a labeled random graph whose links represent the components or subsystems whose reliability description is known or provided (for analysis of a system reliability block diagram, it is sufficient to allow the nodes to be perfectly reliable). The system reliability block diagram expresses the reliability logic of a system in the sense that it shows how the system fails when constituent components and subsystems fail. It is a pictorial representation of the system structure function. Two special nodes are called out: a source or origin node and a terminal or destination node. The system functions if and only if in the random graph there is a path connecting the source node to the terminal node.

In many cases, the system reliability block diagram is a *series-parallel* structure. In such cases, the probability that the system functions is easily concluded from nesting of the standard formulas for the reliability of series systems and parallel systems (see **Parallel, Series, and Series-Parallel Systems**). Other structures, such as the *k-out-of-n* hot standby and *k-out-of-n* cold standby structures, are also amenable to similar probabilistic analysis. Some other structures, such as the bridge structure shown in Figure 1, lend themselves less readily to this type of analysis. In such cases, it may be convenient to use the path set or the cut set methods that are the subject of this article.

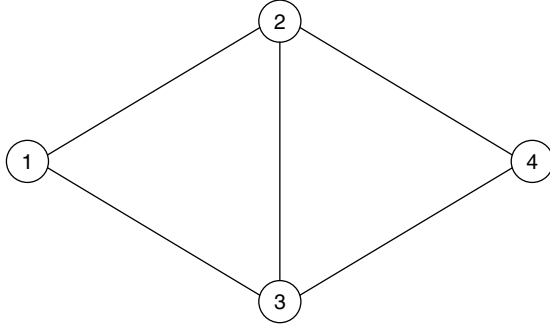


Figure 1 Bridge network

Path Sets and Cut Sets

Given two nodes in a graph, the *path set* for those two nodes is the union of all paths connecting those two nodes. In other terminology in common use, a path set may be called a *tie set*. The *cut set* for those two nodes is the union of all cuts for those two nodes. The *minimal path set* is the union of all minimal paths. The *minimal cut set* is the union of all minimal cuts.

Some authors use the phrases “path set” for what we call a *path* and “cut set” for what we call a *cut*. This makes the naming of the path set or cut set (as defined here) less natural and we avoid that difficulty by using this scheme.

The key concept is that the system functions if and only if there is at least one minimal path whose components are all in a functioning condition. Similarly, the system does not function if and only if there is at least one minimal cut whose components are all in a failed (nonfunctioning) condition. The random graph model provides a framework for computing probabilities of system functioning and failure (nonfunctioning) based on these concepts.

Examples

Consider the reliability block diagram depicted in Figure 2. The small circles represent the nodes of this random graph and these are assumed not to fail. The links representing subsystems that can fail individually in this system have been labeled for convenience by the letters A, B, C, D, and E. The two unlabeled links in the center of the diagram are not associated with any subsystems and are assumed not to fail. Note that, like the bridge structure of Figure 1, this is not a series-parallel graph. In this model, the

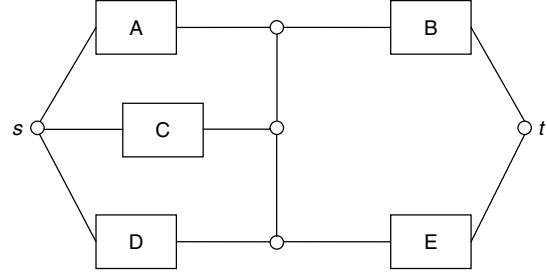


Figure 2 Example of a reliability block diagram

system functions if the node s at the left-hand edge of the diagram and the node t at the right-hand edge of the diagram are connected. This representation indicates that the system functions if any one of the sets $\{A, B\}$, $\{D, E\}$, $\{C, B\}$, or $\{C, E\}$ of units consists entirely of functioning units. Each of these is an (s, t) path. The union of these five paths constitutes the path set for the node pair (s, t) . Note that all these paths are minimal, and these are the only minimal paths, so the path set $\{A, B\} \cup \{D, E\} \cup \{C, B\} \cup \{C, E\}$ is the minimal path set.

Similarly, the system fails to function if any of the sets $\{B, E\}$, $\{A, C, D\}$, $\{C, B, D\}$, or $\{C, A, E\}$ consists entirely of nonfunctioning, or failed, units. Any one of these is a cut for (s, t) , or an (s, t) cut. There are other (s, t) cuts in this graph, such as $\{A, B, E\}$, but the four cuts enumerated above are the only minimal cuts. The minimal cut set for (s, t) is their union $\{B, E\} \cup \{A, C, D\} \cup \{C, B, D\} \cup \{C, A, E\}$.

References

See [4, 6, 7] for additional discussion of path sets and cut sets as models for reliability of complex systems.

System Reliability Modeling and Computation Using Path Sets and Cut Sets

Path Set Method

Because the path set contains all paths connecting s to t , for the system to function it suffices that at least one path be made up entirely of functioning units. Therefore, the probability that the system functions is given by the probability of the path set in the labeled random graph representing the system reliability block diagram. However, only minimal paths need be considered because if a path is not a minimal path,

4 Path Sets and Cut Sets in System Reliability Modeling

then it has a proper subset that is still a path and is a member of the minimal path set. In other words, the union of all (s, t) paths is equal to the union of all (s, t) minimal paths. Consequently, we have the following result:

Proposition 1. The probability that the system functions is given by the probability of the system's minimal path set.

Example. Consider again the system shown in Figure 2. Letting $p_A = P\{A = 1\}$ (where we have abused notation slightly by identifying the indicator random variable's letter with the unit's label) and similarly for B, C, D, and E, the probability that the system functions is given by

$$P(\{A = 1, B = 1\} \cup \{D = 1, E = 1\} \cup \{C = 1, B = 1\} \cup \{C = 1, E = 1\}) \quad (1)$$

This equation shows the strength and weaknesses of the path set method. Its strength is that it is completely straightforward and mechanical to write the expression for the probability that the system functions once the path sets are known. Its weaknesses are that (a) enumerating the paths connecting s and t is tedious for all but the simplest of graphs and (b) the expression that results is the probability of a large union of events that are not, in general, disjoint. However, these weaknesses pertain mainly to manual execution; the algorithmic nature of the procedure means that software for path set reliability analysis is within reach, and indeed has been available for some time [8].

General Representation of System Reliability via the Minimal Path Set. As usual, let $\mathbf{x} = (x_1, \dots, x_c)$ denote the vector of indicators of the functioning of the c components of the system. Enumerate the minimal paths of the system; suppose there are m of them called $\pi_1, \dots, \pi_m$. The structure function for the series system represented by the minimal path π_k is

$$\varphi_k(\mathbf{x}) = \prod_{i \in \pi_k} x_i \quad (2)$$

for $k = 1, \dots, m$. The system functions if and only if at least one of the minimal paths consists entirely of functioning units, so it follows that the structure function for the system may be written as

$$\varphi(\mathbf{x}) = 1 - \prod_{k=1}^m [1 - \varphi_k(\mathbf{x})] = 1 - \prod_{k=1}^m [1 - \prod_{i \in \pi_k} x_i] \quad (3)$$

This equation shows how the system structure function may be represented in terms of the structure functions of the system's minimal paths.

Cut Set Method

Because the cut set contains all (s, t) cuts, for the system to fail it suffices that at least one cut be made up entirely of nonfunctioning units. However, only minimal cuts need be considered because if a cut is not a minimal cut, then it has a proper subset that is still a cut and is a member of the minimal cut set. In other words, the union of all (s, t) cuts is equal to the union of all (s, t) minimal cuts. Therefore, the probability that the system fails to function is given by the probability of the minimal cut set in the labeled random graph representing the system reliability block diagram. Consequently, we have the following result:

Proposition 2. The probability that the system fails is given by the probability of the system's minimal cut set.

Example. Consider again the system shown in Figure 2. The probability that the system fails to function is given by

$$P(\{B = 0, E = 0\} \cup \{A = 0, C = 0, D = 0\} \cup \{C = 0, B = 0, D = 0\} \cup \{C = 0, A = 0, E = 0\}) \quad (4)$$

Again, equation (4) represents the probability of a union of events that are not, in general, disjoint so manual computation can become cumbersome. An algorithm for system reliability evaluation using cut sets may be found in [3].

General Representation of System Reliability via the Minimal Cut Set. Enumerate the minimal cuts of the system; suppose there are n of them called $\chi_1, \dots, \chi_n$. The structure function for the series system represented by the minimal cut χ_k is

$$\varphi_k(\mathbf{x}) = 1 - \prod_{i \in \chi_k} (1 - x_i) \quad (5)$$

for $k = 1, \dots, n$. The system fails if and only if at least one of the minimal cuts consists entirely of nonfunctioning units, so it follows that the structure function for the system may be written as

$$\varphi(\mathbf{x}) = \prod_{k=1}^n \varphi_k(\mathbf{x}) = \prod_{k=1}^n \left[1 - \prod_{i \in \chi_k} (1 - x_i) \right] \quad (6)$$

This equation shows how the system structure function may be represented in terms of the structure functions of the system's minimal cuts.

References. Additional information on the use of path sets and cuts sets for system reliability modeling and computation may be found in [6] and [9].

Bounds

The minimal path set and minimal cut set representations for the system reliability lend themselves readily to the development for bounds on the system reliability. The first such bounds were developed by Esary and Proschan [3]. Letting C (resp., W) denote the minimal cut (resp., path) set for the system, i.e., for the nodes (s, t) , Esary and Proschan's lower bound for the system reliability is

$$\prod_{c \in C} \left(1 - \prod_{i \in c} P\{X_i = 0\} \right) \quad (7)$$

and their upper bound is

$$1 - \prod_{w \in W} \left(1 - \prod_{i \in w} P\{X_i = 1\} \right) \quad (8)$$

The lower bound gives good approximations for highly reliable systems, while the upper bound works better for systems whose components have low reliability. Numerous improvements have been developed; we refer especially to [10] and [11] for later developments.

Paths, Cuts, and Demand Satisfaction in Capacitated Networks

A *network* is a directed graph containing two distinguished subsets $\mathcal{O}$ and $\mathcal{D}$ of $\mathcal{N}$, the sets of origin nodes and destination nodes, respectively. A network is *capacitated* if there is an assignment of nonnegative real numbers to each link and node that represents the *capacity* of that link or node. Origin nodes and destination nodes are called so because the network is required to transport given quantities

of some commodity, discrete (package shipments, telecom/datacom packets) or continuum (oil, gas), tangible (postal letters) or intangible (sensor readings) from the origin node(s) to the destination node(s). The matrix of required quantities is the *exogenous demand*. In other terminology in common use, the origin nodes may be called *sources* and the destination nodes may be called *sinks*.

A *flow* in a network is a map that assigns to each link a nonnegative real number that represents the quantity of the commodity present on that link; see [10]. Sometimes, we speak of the flow on a link (or at a node) to mean the real number in the flow that is attached to that link or node. In a capacitated network, this is required to be no greater than the capacity of the link or node. Determining the flow from the origin-to-destination demands usually requires solution of some optimization problem (such as minimum cost) and is beyond the scope of this article. Nonetheless, in seeking the probability that the demand is satisfied, that is, that the required number of units is transported to the destination nodes, the concepts of path and cut may be generalized to yield useful methods.

The flow into a node is equal to the sum of the flows on the links terminating on that node. Similarly, the flow out of a node is equal to the sum of the flows on the links originating from that node. To say that a demand is satisfied is to say that the flows into the destination nodes are all at least as large as the amount of the commodity required to be delivered to the destinations (in the exogenous demand). Demand satisfaction is a generalization of connectivity. Connectivity is necessary for demand satisfaction: if two nodes are not connected by any paths, then the flow into the terminal node will be zero. Ramirez-Marquez *et al.* [12] developed a generalization of the cut set method described above that is appropriate for demand satisfaction in capacitated networks.

References

- [1] Harary, F. (1971). Graph Theory, *Addison-Wesley*, Reading.
- [2] Bollobas, B. (1998). *Modern Graph Theory*, Springer-Verlag, New York.
- [3] Esary, J.D. & Proschan, F. (1963). Coherent structures of non-identical components, *Technometrics* **5**, 191–209.
- [4] Bollobas, B. (2001). *Random Graphs*, 2nd Edition, Cambridge University Press, New York.

6 Path Sets and Cut Sets in System Reliability Modeling

- [5] Spencer, J. (2001). *The Strange Logic of Random Graphs*, Springer-Verlag, New York.
- [6] Barlow, R.E. & Proschan, F. (1981). *Statistical Theory of Reliability and Life Testing: Probability Models*, To Begin With, Silver Spring.
- [7] Birolini, A. (2004). *Reliability Engineering Theory and Practice*, Springer-Verlag, Berlin.
- [8] Shen, Y. (1995). A New Simple Algorithm for Enumerating All Minimal Paths and Cuts of a graph, *Microelectronics & Reliability*, **35**(6), 973–976.
- [9] Pages, A. & Gondran, M. (1986). *System Reliability Evaluation and Prediction in Engineering*, Springer-Verlag, New York.
- [10] Ford, L.R. & Fulkerson, D.R. (1962). *Flows in Networks*, Princeton University Press, Princeton.
- [11] Koutras, M.V. & Papastavridis, S.G. (1993). Application of the Stein-Chen method for bounds and limit theorems in the reliability of coherent structures, *Naval Research Logistics* **40**, 617–631.
- [12] Ramirez-Marquez, J.E., Coit, D. & Tortorella, M. (2005). Multistate two-terminal reliability: a generalized cut-set approach, *IEEE Transactions on Reliability* (under review).

Further Reading

- Billinton, R. & Allan, R.N. (1983). *Reliability Evaluation of Engineering Systems*, Plenum Press, New York.
- Elsayed, E.A. (1996). *Reliability Engineering*, Addison-Wesley, Reading.
- Fu, J.C. & Koutras, M.V. (1995). Reliability bounds for coherent structures with independent components, *Statistics and Probability Letters* **22**, 137–148.
- Harary, F. (1969). *Graph Theory*, Addison-Wesley, Reading.
- Kolchin, V.F. (1999). *Random Graphs*, Cambridge University Press, New York.
- Rabinowitz, L. (1968). *Reachability and connectedness in random graphs and digraphs*, Ph. D. thesis, Rutgers University, New Brunswick.

MICHAEL TORTORELLA

Component Reliability Importance

Introduction

Consider an n -component coherent system having components whose reliabilities can be expressed in the vector $\mathbf{p} = (p_1, p_2, \dots, p_n)$. The system reliability is then a function of the component reliabilities, so we write

$$r(\mathbf{p}) = P(\text{system functions if components have reliabilities } \mathbf{p}) \quad (1)$$

The *reliability importance* $I(k)$ of component k is defined to be the rate of change in system reliability with respect to the reliability p_k of component k ; in other words

$$I(k) = \frac{\partial r(\mathbf{p})}{\partial p_k} \quad k = 1, 2, \dots, n \quad (2)$$

An alternative formula for the reliability importance of component k can be obtained from decomposition:

$$\begin{aligned} r(\mathbf{p}) &= P(\text{system works}) \\ &= P(\text{(component } k \text{ works and system works)} \\ &\quad \text{or (component } k \text{ fails and system works)}) \\ &= P(\text{component } k \text{ works})P(\text{system works} \mid \\ &\quad \text{component } k \text{ works}) \\ &\quad + P(\text{component } k \text{ fails})P(\text{system works} \mid \\ &\quad \text{component } k \text{ fails}) \end{aligned} \quad (3)$$

If we let $(1_k, \mathbf{p}) = (p_1, p_2, \dots, p_{k-1}, 1, p_{k+1}, \dots, p_n)$ denote the vector $\mathbf{p}$ with 1 in place of p_k , and similarly $(0_k, \mathbf{p}) = (p_1, p_2, \dots, p_{k-1}, 0, p_{k+1}, \dots, p_n)$, then we can write

$$r(\mathbf{p}) = p_k r(1_k, \mathbf{p}) + (1 - p_k) r(0_k, \mathbf{p}) \quad (4)$$

Differentiating both sides with respect to p_k gives

$$I(k) = \frac{\partial r(\mathbf{p})}{\partial p_k} = r(1_k, \mathbf{p}) - r(0_k, \mathbf{p}) \quad (5)$$

This gives an alternative interpretation of the reliability importance: the difference in the reliabilities if component k does and does not work.

Reliability importance is described in detail in Barlow and Proschan [1].

Examples

Series System

For a **series system** with component reliabilities $\mathbf{p} = (p_1, p_2, \dots, p_n)$, the reliability function is

$$r(\mathbf{p}) = \prod_{i=1}^n p_i = p_1 p_2 \cdots p_k \cdots p_n \quad (6)$$

so that

$$\begin{aligned} I(k) &= \frac{\partial r(\mathbf{p})}{\partial p_k} = \frac{\partial}{\partial p_k} [(p_1 p_2 \cdots p_{k-1} p_{k+1} \cdots p_n) p_k] \\ &= (p_1 p_2 \cdots p_{k-1} p_{k+1} \cdots p_n) \\ &= \frac{p_1 p_2 \cdots p_{k-1} p_{k+1} \cdots p_n}{p_k} \\ &= \frac{r(\mathbf{p})}{p_k} \end{aligned} \quad (7)$$

Thus, the reliability importance of component k is the product of the reliabilities of all components except component k , or equivalently, it is equal to the system reliability divided by the component reliability p_k . Thus, the least reliable component is the one with the greatest reliability importance. This makes intuitive sense: the system reliability is improved the most when the least reliable component is improved.

For illustration, suppose that the components have reliabilities

$$p_1 = 0.9 \quad (8)$$

$$p_2 = 0.99 \quad (9)$$

$$p_3 = 0.95 \quad (10)$$

The reliability importance for the three components are

$$I(1) = p_2 p_3 = 0.99 \times 0.95 = 0.9405 \quad (11)$$

$$I(2) = p_1 p_3 = 0.9 \times 0.95 = 0.855 \quad (12)$$

$$I(3) = p_1 p_2 = 0.9 \times 0.99 = 0.891 \quad (13)$$

For a series system, the structure is symmetric in the components, meaning that interchanging components leads to the same system structure, so if the

2 Component Reliability Importance

component reliabilities are all equal, this would imply that all components have equal importance. This is clearly true since

$$I(k) = \frac{r(\mathbf{p})}{p_k} = \frac{p^n}{p} = p^{n-1} \quad (14)$$

for all k .

Parallel System

For a **parallel system**, the reliability function is

$$r(\mathbf{p}) = 1 - \prod_{i=1}^k (1 - p_i) = 1 - (1 - p_1) \times (1 - p_2) \cdots (1 - p_k) \cdots (1 - p_n) \quad (15)$$

so the reliability importance of the k th component is

$$\begin{aligned} \frac{\partial r(\mathbf{p})}{\partial p_k} &= \frac{\partial}{\partial p_k} [1 - (1 - p_1) \times (1 - p_2) \cdots (1 - p_k) \cdots (1 - p_n)] \\ &= -(1 - p_1)(1 - p_2) \cdots (1 - p_{k-1}) \\ &\quad \times (1 - p_{k+1}) \cdots (1 - p_n)(-1) \\ &= (1 - p_1)(1 - p_2) \cdots (1 - p_{k-1}) \\ &\quad \times (1 - p_{k+1}) \cdots (1 - p_n) \\ &= \frac{(1 - p_1)(1 - p_2) \cdots (1 - p_n)}{1 - p_k} \\ &= \frac{1 - r(\mathbf{p})}{1 - p_k} \end{aligned} \quad (16)$$

Thus, since the numerator is a constant, independent of k , the component with the highest reliability is the one with the greatest reliability importance. In a parallel system, the system reliability is improved the most when the most reliable component is improved.

If the component reliabilities are

$$p_1 = 0.8 \quad (17)$$

$$p_2 = 0.9 \quad (18)$$

$$p_3 = 0.95 \quad (19)$$

$$p_4 = 0.99 \quad (20)$$

then the system reliability is

$$r(\mathbf{p}) = 0.99999 \quad (21)$$

and the reliability importances are

$$I(1) = \frac{1 - 0.99999}{1 - 0.8} = 0.00005 \quad (22)$$

$$I(2) = \frac{1 - 0.99999}{1 - 0.9} = 0.0001 \quad (23)$$

$$I(3) = \frac{1 - 0.99999}{1 - 0.95} = 0.0002 \quad (24)$$

$$I(4) = \frac{1 - 0.99999}{1 - 0.99} = 0.001 \quad (25)$$

Series-Parallel Systems

Consider the **series-parallel system** shown in Figure 1.

For this system, the reliability function is

$$r(\mathbf{p}) = [1 - (1 - p_1)(1 - p_2)][1 - (1 - p_3)(1 - p_4)] \times [1 - (1 - p_5)(1 - p_6)] \quad (26)$$

The reliability importances are

$$I(1) = \frac{\partial r(\mathbf{p})}{\partial p_1} = (1 - p_2) \times [1 - (1 - p_3)(1 - p_4)] \times [1 - (1 - p_5)(1 - p_6)] \quad (27)$$

$$I(2) = \frac{\partial r(\mathbf{p})}{\partial p_2} = (1 - p_1) \times [1 - (1 - p_3)(1 - p_4)] \times [1 - (1 - p_5)(1 - p_6)] \quad (28)$$

$$I(3) = \frac{\partial r(\mathbf{p})}{\partial p_3} = (1 - p_4) [1 - (1 - p_1)(1 - p_2)] \times [1 - (1 - p_5)(1 - p_6)] \quad (29)$$

$$I(4) = \frac{\partial r(\mathbf{p})}{\partial p_4} = (1 - p_3) [1 - (1 - p_1)(1 - p_2)] \times [1 - (1 - p_5)(1 - p_6)] \quad (30)$$

$$I(5) = \frac{\partial r(\mathbf{p})}{\partial p_5} = (1 - p_6) [1 - (1 - p_1)(1 - p_2)] \times [1 - (1 - p_3)(1 - p_4)] \quad (31)$$

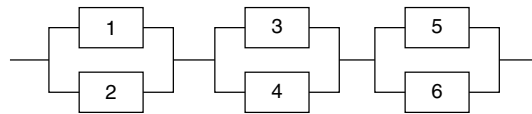


Figure 1 Series-parallel system with six components

$$I(6) = \frac{\partial r(\mathbf{p})}{\partial p_6} = (1 - p_5) [1 - (1 - p_1)(1 - p_2)] \times [1 - (1 - p_3)(1 - p_4)] \quad (32)$$

In the special case where $p_1 = p_2 = 0.9$, $p_3 = p_4 = 0.95$, and $p_5 = p_6 = 0.99$, the component importances are

$$I(1) = I(2) = (1 - p_1) [1 - (1 - p_3)^2] \times [1 - (1 - p_5)^2] = 0.09974 \quad (33)$$

$$I(3) = I(4) = (1 - p_3) [1 - (1 - p_1)^2] \times [1 - (1 - p_5)^2] = 0.0494951 \quad (34)$$

$$I(5) = I(6) = (1 - p_5) [1 - (1 - p_1)^2] \times [1 - (1 - p_3)^2] = 0.00987525 \quad (35)$$

Structural Importance

The structural importance of a component is related to the structure of the system, but does not depend on the system reliabilities. Consider for instance the system whose block diagram is shown in Figure 2. It would seem that component 1 is most important, because the system does not function without it.

In general, we let $\mathbf{x} = (x_1, x_2, \dots, x_n)$ be the vector of zeros and ones representing the state of the components, where $x_i = 1$ if component i is working and 0 if it is not working. The structure function (for a given system) is

$$\phi(\mathbf{x}) = \begin{cases} 1, & \text{if the system works for the state vector } \mathbf{x} \\ 0, & \text{if the system fails for the state vector } \mathbf{x} \end{cases} \quad (36)$$

Then let $(1_k, \mathbf{x}) = (x_1, x_2, \dots, x_{k-1}, 1, x_{k+1}, \dots, x_n)$ and $(0_k, \mathbf{x}) = (x_1, x_2, \dots, x_{k-1}, 0, x_{k+1}, \dots, x_n)$. If for a given vector $\mathbf{x}$ and component k , the system works if $x_k = 1$ and fails if $x_k = 0$, that is,

$$\phi(1_k, \mathbf{x}) - \phi(0_k, \mathbf{x}) = 1 \quad (37)$$

then component k would seem important. If, on the other hand,

$$\phi(1_k, \mathbf{x}) - \phi(0_k, \mathbf{x}) = 0 \quad (38)$$

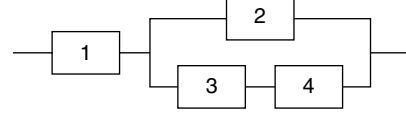


Figure 2 A four-component system

Table 1 Computation of $I_\phi(1)$

$\mathbf{x}$	$\phi(1_1, \mathbf{x})$	$\phi(0_1, \mathbf{x})$	$\phi(1_1, \mathbf{x}) - \phi(0_1, \mathbf{x})$
(*, 0, 0, 0)	0	0	0
(*, 0, 0, 1)	0	0	0
(*, 0, 1, 0)	0	0	0
(*, 0, 1, 1)	1	0	1
(*, 1, 0, 0)	1	0	1
(*, 1, 0, 1)	1	0	1
(*, 1, 1, 0)	1	0	1
(*, 1, 1, 1)	1	0	1

then system works (fails) whether component k works (fails), suggesting that component k is not so important.

The *structural reliability importance* of component k is obtained by averaging the differences $\phi(1_k, \mathbf{x}) - \phi(0_k, \mathbf{x})$ over all 2^{n-1} possible state vectors $\mathbf{x}$ (without component k). Specifically,

$$I_\phi(k) = \frac{1}{2^{n-1}} \sum_{\{\mathbf{x}: x_k=1\}} [\phi(1_k, \mathbf{x}) - \phi(0_k, \mathbf{x})] \quad (39)$$

For illustration, consider the four component system in Figure 2. Table 1 shows all eight state vectors (excluding the first component) and the values $\phi(1_k, \mathbf{x})$ and $\phi(0_k, \mathbf{x})$. From this table, we can compute the structural reliability importance for component 1:

$$\begin{aligned} I_\phi(1) &= \frac{1}{2^{4-1}} \sum_{\{\mathbf{x}: x_1=1\}} [\phi(1_1, \mathbf{x}) - \phi(0_1, \mathbf{x})] \\ &= \frac{1}{8} [0 + 0 + 0 + 1 + 1 + 1 + 1 + 1] = \frac{5}{8} \end{aligned} \quad (40)$$

The corresponding tables for components 2 and 3 are shown in Tables 2 and 3. From this we can compute the structural importance of component 2 and 3 (by symmetry, components 3 and 4 have the same structural reliability importance):

$$I_\phi(2) = \frac{1}{8} [0 + 0 + 0 + 0 + 1 + 1 + 1 + 0] = \frac{3}{8}$$

4 Component Reliability Importance

Table 2 Computation of $I\phi(2)$

$\mathbf{x}$	$\phi(1_2, \mathbf{x})$	$\phi(0_2, \mathbf{x})$	$\phi(1_2, \mathbf{x}) - \phi(0_2, \mathbf{x})$
(0, *, 0, 0)	0	0	0
(0, *, 0, 1)	0	0	0
(0, *, 1, 0)	0	0	0
(0, *, 1, 1)	0	0	0
(1, *, 0, 0)	1	0	1
(1, *, 0, 1)	1	0	1
(1, *, 1, 0)	1	0	1
(1, *, 1, 1)	1	1	0

Table 3 Computation of $I\phi(3) = I\phi(4)$

$\mathbf{x}$	$\phi(1_3, \mathbf{x})$	$\phi(0_3, \mathbf{x})$	$\phi(1_3, \mathbf{x}) - \phi(0_3, \mathbf{x})$
(0, 0, *, 0)	0	0	0
(0, 0, *, 1)	0	0	0
(0, 1, *, 0)	0	0	0
(0, 1, *, 1)	0	0	0
(1, 0, *, 0)	0	0	0
(1, 0, *, 1)	1	0	1
(1, 1, *, 0)	1	1	0
(1, 1, *, 1)	1	1	0

(41)

$$I_{\phi}(3) = I_{\phi}(4) = \frac{1}{8}[0 + 0 + 0 + 0 + 0 + 1 + 0 + 0] = \frac{1}{8} \quad (42)$$

For the system in Figure 2, component 3 (and also by symmetry, component 4) has structural reliability importance equal to $\frac{1}{8}$. This means that there is only one configuration out of the eight possible in which it matters whether component 3 is working. This is the case where $x_1 = 1$, $x_2 = 0$, and $x_4 = 1$. In all

seven of the other cases, it does not matter whether component 3 works or fails; the state of the system will be the same.

Relationship between Reliability Importance and Structural Reliability Importance

We saw in the section titled “Examples” that the reliability importance can be written as

$$I(k) = r(1_k, \mathbf{x}) - r(0_k, \mathbf{x}) \quad (43)$$

If all components except component k have reliability $\frac{1}{2}$, and if components act independently, then all 2^{n-1} possible configurations are equally likely and have probability $1/2^{n-1}$. If $\mathbf{X}$ denotes the random state vector, then we can write

$$E[\phi(1_k, \mathbf{X}) - \phi(0_k, \mathbf{X})] = \sum_{\{\mathbf{x}: x_k=1\}} \frac{1}{2^{n-1}} [\phi(1_k, \mathbf{x}) - \phi(0_k, \mathbf{x})] = I(k) \quad (44)$$

Thus, for the special case of $p_i = \frac{1}{2}$, $i \neq k$, the structural reliability importance and the reliability importance coincide.

Reference

- [1] Barlow, R.E. & Proschan, F. (1975). *Statistical Theory of Life Testing: Probability Models*, Holt, Reinhart & Winston, Austin.

STEVEN E. RIGDON

Modules and Modular Decomposition

Introduction

In engineering terminology, a *module* is simply a cluster of components that is treated as a single entity in a system. For example, an engine is often treated as a module in the airline industry. If an engine has failed or has signs of serious **degradation**, it is taken off the airplane as a whole and sent to the engine manufacturer for repair and/or rebuild. A self-contained assembly of electronic components and circuitry such as video card, a hard disk, or a motherboard is also referred to as a module. In software engineering, a module is defined as a software entity that groups a set of subprograms and data structures. These software modules are units that can be compiled separately, which makes them reusable and allows multiple programmers to work on different modules simultaneously. A module may also be an integrated unit including both hardware and software.

In system reliability theory, a module represents a group of components that has a single input from and a single output to the rest of the system in the system's reliability block diagram (*see Coherent Systems*). The state of the module can be represented by a random variable. The contribution of all components in a module to the performance of the whole system can be represented by the state of the module. Once the state of the module is known, one does not need to know the states of the components within the module to determine the state of the system. If we know the reliabilities of the modules in a system, we can use them to find the reliability of the system.

With the advancement of technology, today's engineering systems are becoming more and more sophisticated. For example, a spacecraft may consist of millions of components. In design optimization of such a system subject to budget, reliability, physical weight, and other constraints, efficient reliability evaluation algorithms are needed. Appropriate identification of modules within a system and use of modular decomposition will enhance the efficiency of system reliability evaluation and, as a result, that of design optimization.

In this article, we will provide the formal definitions of modules and modular decomposition for system reliability analysis, illustrate the technique of modular decomposition, and outline reliability evaluation algorithms for several well-studied modules. References for further reading in this area will also be provided. Unless otherwise stated, we adopt the following assumptions in this article:

1. The systems to be discussed are coherent systems (*see Coherent Systems*).
2. The system and its components may be in only two possible states, either working or failed.
3. The states of the components are independent random variables.
4. The mission time of the system is implicitly specified.

On the basis of these assumptions, we deal primarily with system and component reliabilities rather than their reliability functions of time t .

Definitions and General Guidelines

In reliability analysis, a *system* is defined as a set of components that work together to perform a specific function. The reliability of the system is determined by the reliabilities of its components and how the components are logically connected together to form the system. If there exists a subset of components in the system which work together to perform a sub-function of the system, we say that this subset of components forms a subsystem. The reliability of this subsystem is completely determined by its components and how they are logically connected together to form the subsystem. In reliability evaluation of the original system, if we can treat this subsystem as a single "component" with a reliability value of the subsystem and the reliability of the system obtained this way is the same as that obtained when all the components of the subsystem are explicitly considered, then we say that this subsystem of components is actually a "module". This is an intuitive explanation of what a module is. If we use the definition of a coherent system (see Barlow and Proschan [1]), a module can be explained to be a subset of components which form a coherent system structure such that the system as a whole is coherent even if the components in the module are treated as a single

2 Modules and Modular Decomposition

“component”. The formal definition of a module is given below.

Definition 1 (Barlow and Proschan [1]) The coherent system (A, χ) is a module of the coherent system (C, ϕ) if $\phi(\mathbf{x}) = \psi(\chi(\mathbf{x}^A), \mathbf{x}^{A^c})$, where ψ is a coherent structure function, A is a subset of C , A^c is the subset of C that is complementary to A , and $\mathbf{x}^A$ is a vector representing the states of the components in set A . The set $A \subseteq C$ is called a modular set of (C, ϕ) .

On the basis of this definition, each component by itself is a module of the system. The system itself can also be considered to be a module. However, these two extreme cases do not provide any advantage in system reliability evaluation.

If A is a module, A^c must be a module too. If a module A has been identified in a system, we can say that the system has been decomposed into two modules, A and A^c . The system can be considered as consisting of these two modules. Each module has a structure function. The structure function of the system can be expressed as a function of the structure functions of the two modules.

More than two modules may exist in a coherent system. Modular decomposition is a technique that can be used to decompose a coherent system into several disjoint modules. Such a decomposition is useful if the structure function of each module can be easily derived and the structure function relating the system state to the states of these disjoint modules can be easily derived. A formal definition of modular decomposition is given below.

Definition 2 (Barlow and Proschan [1]) A modular decomposition of a coherent system (C, ϕ) is a set of disjoint modules, $\{(A_1, \chi_1), (A_2, \chi_2), \dots, (A_r, \chi_r)\}$ together with an organizing structure ψ such that

$$C = \bigcup_{i=1}^r A_i, \text{ and } A_i \cap A_j = \emptyset \text{ for } i \neq j \quad (1)$$

$$\phi(\mathbf{x}) = \psi[\chi_1(\mathbf{x}^{A_1}), \chi_2(\mathbf{x}^{A_2}), \dots, \chi_r(\mathbf{x}^{A_r})] \quad (2)$$

To apply the modular decomposition technique in the derivation of system structure functions, we need to identify disjoint modules within the system structure. The structure function of each disjoint module must be easy to derive. These modules can

then be treated as “supercomponents” or subsystems, and we only need to find the relationship between the state of the system and the states of these modules. Hopefully, the number of such modules is much smaller than the number of components in the system.

We may use the following two approaches individually or in combination to identify modules in a complex system. The first approach defines modules based on a clear understanding of the functions of some components in the system. For example, in a chemical refinery, a primary pump is used to provide the required flow rate and pressure of a certain liquid. A secondary pump is installed as a cold standby unit. These two pumps together with their auxiliary instruments such as the sensing and switching mechanism can be called the *pump module*. The reliability of the pump module can be derived using the standby structure and used in further evaluation of the chemical refinery’s reliability. Another example is a power system including generation, transmission, and distribution subsystems. A single power station including several generators may be considered to be a module if all these generators use a single point of delivering power to the transmission network. This module of generators may have a *k-out-of-n* structure (see *k-out-of-n Systems*).

The second approach assumes that the reliability block diagram of the system has already been developed. From the system reliability block diagram, we then identify modules following Definition 1. We will illustrate the technique of modular decomposition using this second approach in the section titled “Modular Decomposition”.

Well-studied modules include the series structure, the parallel structure, the series–parallel structure, the parallel–series structure, the bridge structure, the *k-out-of-n* structure, the standby structure, and the consecutive-*k-out-of-n* structure. Efficient reliability evaluation algorithms have been reported for these modules. These modules will be further elaborated in the section titled “Well-Studied Modules”.

In subsequent discussions, the following notation will be used:

p_i : reliability of component i

q_i : unreliability of component i ,

$$q_i = 1 - p_i$$

R_a : reliability of a where

a is a string of letters or

numerals that may represent

a module or subsystem or the system

$Q_a : 1 - R_a$.

Modular Decomposition

Consider the reliability block diagram depicted in Figure 1(a). The nine components are numbered from 1 to 9. We are interested in finding the reliability of this system.

Examining Figure 1(a), we can identify the following modules:

- Module I consists of components 1 and 2 connected in series.
- Module II consists of components 3 and 4 connected in series.
- Module III consists of components 6, 7, and 8, wherein components 7 and 8 connected in parallel form a submodule which is then connected to component 6 in series.

The three modules listed above are the only meaningful modules in Figure 1(a) that may be used for efficient evaluation of system reliability. They are the only clusters of components that have a single input from and a single output to the rest of the system in the reliability block diagram. The

reliabilities of these modules are

$$R_I = p_1 p_2 \quad (3)$$

$$R_{II} = p_3 p_4 \quad (4)$$

$$R_{III} = p_6(1 - q_7 q_8) \quad (5)$$

With these identified modules, the system reliability block diagram shown in Figure 1(a) is transformed into the one shown in Figure 1(b). This resulting reliability block diagram is actually called the *bridge structure*.

If direct identification of known modules does not yield a simple organizing structure for system reliability evaluation, there are other techniques one may apply for further identification of modules. These techniques include the so-called pivotal decomposition, Δ -star transformation, and star- Δ transformation. For details on these techniques, readers may refer to Kuo and Zuo [2] and Misra [3].

Well-Studied Modules

Many special system structures have been studied by researchers. Efficient reliability evaluation algorithms have been reported for these structures. In this section, we call these special structures modules and provide their definitions and reliability evaluation algorithms (*see Parallel, Series, and Series-Parallel Systems; k-out-of-n Systems*).

Any n components connected in series form the *series module*. The reliability of a series module can be calculated as

$$R_{\text{series}} = p_1 p_2 \cdots p_n \quad (6)$$

Any n components connected in parallel form the *parallel module*. The reliability of a parallel module can be calculated as

$$R_{\text{parallel}} = 1 - q_1 q_2 \cdots q_n \quad (7)$$

A parallel-series system can be considered to consist of m disjoint modules that are connected in parallel and module i ($1 \leq i \leq m$) consists of n_i components that are connected in series. If the number of components in each of the m modules is not too large and m is not too large, the whole parallel-series system may be considered to be a module. The reliability of such a *parallel-series*

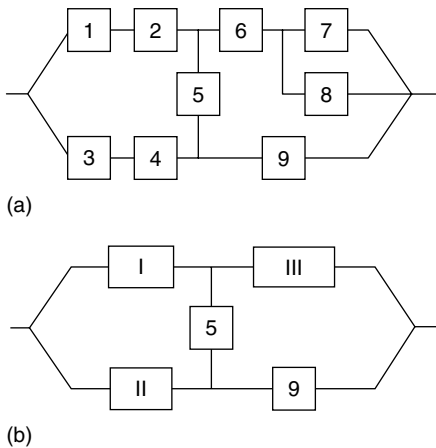


Figure 1 Illustration of modular decomposition

4 Modules and Modular Decomposition

module can be expressed as

$$R_{\text{parallel-series}} = 1 - \prod_{i=1}^m \left(1 - \prod_{j=1}^{n_i} p_{ij} \right) \quad (8)$$

where p_{ij} is the reliability of component j in series module i for $1 \leq i \leq m$ and $1 \leq j \leq n_i$.

A series-parallel system can be considered to consist of m disjoint modules that are connected in series and module i ($1 \leq i \leq m$) consists of n_i components that are connected in parallel. If the number of components in each of the m modules is not too large and m is not too large, the whole series-parallel system may be considered to be a module. The reliability of such a *series-parallel module* can be expressed as

$$R_{\text{series-parallel}} = \prod_{i=1}^m \left(1 - \prod_{j=1}^{n_i} q_{ij} \right) \quad (9)$$

where q_{ij} is the unreliability of component j in parallel module i for $1 \leq i \leq m$ and $1 \leq j \leq n_i$.

A standby structure can also be considered to be a module. Consider a *standby module* with a primary unit and one cold standby unit. There is a sensing and switching mechanism that switches the standby unit into operation when it detects the failure of the primary unit. To evaluate the reliability of such a module, one needs to know the lifetime distribution function of each unit. Under the assumptions that each component being used follows the exponential lifetime distribution with parameter λ , the standby unit is perfect, and the sensing and switching mechanism is perfect, the reliability function of the module can be expressed as

$$R_{\text{standby}}(t) = e^{-\lambda t} (1 + \lambda t) \quad (10)$$

With this equation, one can evaluate the reliability of the module for any specified mission time t . For more advanced analysis of the standby module under more relaxed conditions, readers are referred to Kuo and Zuo [2].

The reliability block diagram of the bridge structure is shown in Figure 2. It can be called the *bridge module*. The reliability of this module can be expressed as

$$R_{\text{bridge}} = p_3(1 - q_1q_4)(1 - q_2q_5) + q_3 \times [1 - (1 - p_1p_2)(1 - p_4p_5)] \quad (11)$$

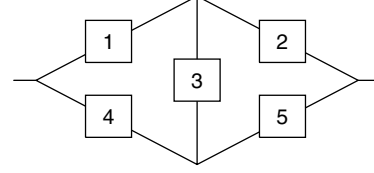


Figure 2 The bridge module

The k -out-of- n structure may also appear in a sophisticated reliability system. A k -out-of- n system works if and only if at least k of the n components work. A recursive formula for reliability evaluation of such a k -out-of- n module is given by Rushdi [4] as follows:

$$R(k, n) = p_n R(k-1, n-1) + q_n R(k, n-1) \quad (12)$$

where $R(a, b)$ denotes the probability that at least a components out of the b components work. The boundary conditions for equation (12) are

$$R(0, j) = 1, \quad j \geq 0 \quad (13)$$

$$R(j+1, j) = 0, \quad j \geq 0 \quad (14)$$

It would be advantageous to treat a k -out-of- n structure as a module in a complex system for its reliability evaluation only if the number of components in the k -out-of- n structure is not too large.

A consecutive- k -out-of- n system consists of linearly arranged components and it is failed if and only if at least k consecutive components are failed. Such a system structure may appear as a module in a sophisticated system. An efficient recursive equation is provided by Hwang [5] for reliability evaluation of such a *consecutive- k -out-of- n module*:

$$R_c(k, n) = R_c(k, n-1) - R_c(k, n-k-1) \times p_{n-k} \prod_{j=n-k+1}^n q_j \quad (15)$$

where $R_c(a, b)$ is the probability that at least a consecutive components have failed among the first b components. The boundary conditions for equation (15) are

$$R_c(a, b) = 1, \quad b < a \quad (16)$$

$$p_0 \equiv 0 \quad (17)$$

Advanced coverage of the consecutive- k -out-of- n structure and its variations can be found in Chang *et al.* [6].

Summary and Further Reading

A key step in modular decomposition is the definition of modules. Each identified module should be simple enough to enable us to obtain its various reliability characteristics relatively easily. Consider the example reliability block diagram given in Figure 3. There are 10 components in the system. One option is to decompose the system into three modules with module I including components 1 through 4, module II including components 5 through 7, and module III including components 8 through 10. Module I has a parallel structure, module II has a series-parallel structure, while module III has a 2-out-of-3 structure. The organization structure of the system consisting of these three modules is a series structure.

To evaluate a large and complex reliability system, we often have to apply the so-called hierarchical modular decomposition. The system is decomposed into modules (this is called the *level 1 decomposition*). Each module in the system may be further decomposed into submodules (this is called the *level 2 decomposition*). The process continues until the modules or submodules in the lowest level of decomposition are simple enough.

The technique of modular decomposition can be used not only to evaluate the exact system reliability but also to derive bounds on system reliability. This can be achieved through hierarchical modular decomposition. Starting from the lowest decomposition level, we first find the bounds on the reliabilities of all the modules at this level. Once these bounds are obtained, these modules at the lowest level are treated as supercomponents with known reliability bounds. Higher-level modules consisting of these supercomponents are then analyzed. Bounds

on these higher-level modules may be obtained. This process is repeated until we reach the system level wherein the upper and/or lower bounds on the system reliability can be obtained from the bounds of its immediate modules. As illustrated by Barlow and Proschan [1], utilization of modular decomposition can improve the bounds obtained.

The modules summarized in this article are widely seen in engineering systems. The series, parallel, parallel-series, and series-parallel modules are most commonly seen. The bridge module is used in networks. The standby module is often used in mechanical systems. The k -out-of- n module is seen in modeling power generation stations and the triple-modular redundancy (TMR) and N -modular redundancy (NMR) schemes in software engineering and fault tolerance control (see Hecht [7], Xie *et al.* [8]), and N -version programming (NVP) (see Avizienis [9]). The consecutive- k -out-of- n module has been used in modeling pipeline transportation systems (see Chang *et al.* [6]).

Advanced studies of the k -out-of- n and the consecutive- k -out-of- n modules can be found in Chang *et al.* [6], Kuo and Zuo [2], and Lisnianski and Levitin [10]. Both components and the system may assume binary states or multiple states (more than two states). There may be common-cause failures. The system and the components may be repairable or nonrepairable. Other reliability characteristics such as mean time to failure (MTTF), mean time between failure (MTBF), and availability have also been studied. Another topic involving modules that has been well studied is the assembly of components into different modules of a system. This involves allocation of components having different reliabilities to different modules within the system. El-Newehi *et al.* [11] provided optimal allocations of components to parallel-series and **series-parallel systems**. Du and Hwang [12] provided optimal allocations of components to an s -stage k -out-of- n system, in which the system can be hierarchically decomposed into s stages and each stage has the k -out-of- n system structure.

References

- [1] Barlow, R.E. & Proschan, F. (1975). *Statistical Theory of Reliability and Life Testing: Probabilities*, Holt, Rinehart and Winston, New York.

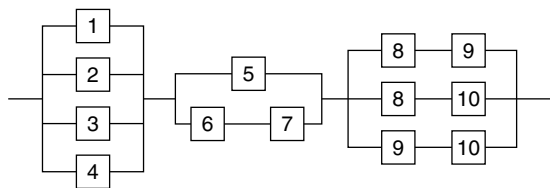


Figure 3 An example system reliability block diagram

- [2] Kuo, W. & Zuo, M.J. (2003). *Optimal Reliability Modeling: Principles and Applications*, John Wiley & Sons, New York.
- [3] Misra, K.B. (1992). *Reliability Analysis and Prediction: A Methodology Oriented Treatment*, Elsevier, Amsterdam.
- [4] Rushdi, A.M. (1986). Utilization of symmetric switching functions in the computation of k -out-of- n system reliability, *Microelectronics and Reliability* **26**(5), 973–987.
- [5] Hwang, F.K. (1982). Fast solutions for consecutive- k -out-of- $n:F$ system, *IEEE Transactions on Reliability* **R-31**(5), 447–448.
- [6] Chang, G.J., Cui, L. & Hwang, F.K. (2000). *Reliabilities of Consecutive- k Systems*, Kluwer, Boston.
- [7] Hecht, H. (2004). *Systems Reliability and Failure Prevention*, Artech House, Boston.
- [8] Xie, M., Dai, Y.S. & Poh, K.L. (2004). *Computing Systems Reliability: Models and Analysis*, Kluwer Academic, New York.
- [9] Avizienis, A. (1995). The methodology of N -version programming, in *Software Fault Tolerance*, M.R. Lyu ed, John Wiley & Sons, Chapter 2, pp. 23–46.
- [10] Lisnianski, A. & Levitin, G. (2003). *Multi-state System Reliability: Assessment, Optimization and Applications*, World Scientific, Singapore.
- [11] El-Newehi, E., Proschan, F. & Sethuraman, J. (1986). Optimal allocation of components in parallel-series and series-parallel systems, *Journal of Applied Probability* **23**(3), 770–777.
- [12] Du, D.Z. & Hwang, F.K. (1990). Optimal assembly of an s -stage k -out-of- n system, *SIAM Journal of Discrete Mathematics* **3**(3), 349–354.

Related Articles

Coherent Systems; k -out-of- n Systems; Parallel, Series, and Series-Parallel Systems; Reliability of Redundant Systems.

MING J. ZUO

Repairable Systems Reliability

Introduction

We begin by defining a repairable system, after some preliminary concepts and definitions.

1. *Part*: an item not subject to disassembly and hence discarded when it fails.
2. *Socket*: a circuit or equipment position which, at any given time, holds a part of a given type.
3. *System*: a collection of two or more sockets and their associated parts, interconnected to perform one or more functions.
4. *Nonrepairable System*: a system that is discarded the first time that it ceases to perform satisfactorily.
5. *Repairable System*: a system which, after failing to perform at least one of its required functions, can be returned to performing all of its required functions satisfactorily by any method other than replacement of the entire system.

Three points must be made. First, since small appliances are systems, many systems, perhaps even a majority, are nonrepairable. Nevertheless, the overwhelming majority of systems of interest in reliability applications are designed to be repaired, rather than discarded, after their first failure. Henceforth, therefore, the term *system* will be used to denote a repairable system.

Secondly, given that a system contains n parts, the definition of a repairable system allows up to $n - 1$ part replacements during a single repair. In practice, however, when the repair would require many new parts, it usually is more cost effective to replace the entire system. Most repairs involve the replacement of only a minute fraction of a system's constituent parts and this has major implications for probabilistic modeling and therefore, for statistical analysis as well. Some repairs, e.g., cleaning contacts or adjusting internal potentiometers, do not involve replacement of any parts.

Thirdly, if a system can be returned to satisfactory operation by "occasional" adjustments of external, i.e. "front panel" controls, it hasn't failed and therefore, doesn't require repair. However, "too frequently"

required adjustments necessitate repair to reduce the rate of such adjustments.

Probabilistic Modeling

Most of the literature concerning methods for predicting the time to failure of a system assumes that system failure is an absorbing state, i.e., that the system of interest is nonrepairable. In many cases, the nonrepairable system can include one or more groups of repairable redundant subsystems. Once a system-level failure occurs, however, the system is assumed to be discarded. Such "reliability with repair" models will not be considered in this article.

This section concentrates on "black-box" modeling and analysis. That is, models are postulated for the pattern of system-level failures, regardless of the system's design, and the postulated models can be tested against even small datasets.

Model for Parts

The time to failure of a part is a random variable, X , described either with the cumulative distribution function

$$F_x(x) \equiv \Pr\{X \leq x\} \quad (1)$$

or the force of mortality (FOM)

$$h_x(x) \equiv \frac{F'_x(x)}{1 - F_x(x)} \quad (2)$$

Intuitively, $h_x(x)$ measures how likely it is that *one* failure will occur soon after time x , given that *it* has not occurred by time x .

Model for Systems

The failures of a system are described by a stochastic point process $T_1, T_2, \dots$, where T_i denotes the arrival time to the i th failure, measured from the instant at which a system was first put into operation. Equivalently, the process can be represented by the interarrivals times, $X_1, X_2, X_3, \dots$, where $X_i \equiv T_i - T_{i-1}$ and $T_0 \equiv 0$. Downtimes, usually very small compared to most interarrivals times, are ignored throughout, except for a brief mention in the section titled "Misconceptions and their Causes".

Let $N(t)$ be the observed number of system failures in the interval $(0, t]$. The expected number

2 Repairable Systems Reliability

of failures is $V(t) \equiv E[N(t)]$. The rate of occurrence of failures (ROCOF) is defined as

$$v(t) \equiv \frac{d}{dt} E[N(t)] \quad (3)$$

The ROCOF is an absolute rate, whereas the FOM is a relative rate. Hence, even though ROCOF and FOM can be, and sometimes must be, represented by the same mathematical function, with numerically equal parameters, the section titled “Misconceptions and Their Causes” shows that they have up-to-infinitely different interpretations.

We will consider only the homogeneous **Poisson process** (HPP) and some generalizations of it: the renewal process (RP), the superimposed renewal process (SRP), and the nonhomogeneous Poisson process (NHPP). The HPP can be defined as a *nonterminating sequence of exponential interarrivals times* which are independent and identically distributed (i.i.d.). The ROCOF of an HPP is a constant for all $t \geq 0$. The RP is a direct generalization, since the i.i.d. interarrivals times can be distributed according to any nonnegative distribution.

Consider two or more independent RPs, each of which represents the pattern of failures in a socket. Then the union of all failures, including the instants at which they occurred, is an SRP. In general, the SRP is not an RP; in fact, if the superposition of two independent RPs is an RP, then all three are HPPs [1].

For a formal definition of an NHPP, (see **Intensity Functions for Nonhomogeneous Poisson Processes**). Here, an NHPP can be considered to be any process where the ROCOF is a nonnegative deterministic function of time, independent of the history of the process.

Most repairs involve replacement of only a minute fraction of a system’s parts. It is very implausible, therefore, to assume the RP as a model, i.e. to assume that a system’s effective age is reduced to zero by every repair. In fact, under the universal engineers’ and laymen’s usage of the term *repair*, a well-performed repair returns a failed system to satisfactory performance, without *even intending* to reduce its age to zero. Nevertheless, the RP is often assumed to be *the* model for a system; see, e.g., [2, 3].

A plausible system model can be developed by assuming renewal in each of the system’s sockets. After all, the failed part in the socket is replaced with a new one. Given an RP for each socket, and a

series system, the resulting model is the SRP. There may be cases where one can work with an SRP directly, but usually an approximation is needed. The model for a finite number of superimposed **renewal processes** is unknown, but there are limit theorems (see [4]) indicating that when the number n of sockets in a series system increases without bound, the SRP converges to an NHPP. (If, in addition $t \rightarrow \infty$, then the SRP converges to an HPP [5].)

Many systems have some redundant paths, so it might appear that the NHPP approximation would not apply to them. However, the series parts often dominate the system’s reliability, so that redundant paths often can be ignored in developing the $SRP \rightarrow NHPP$ model.

A more pragmatic way of “deriving” the NHPP as the first-order system model is to imagine someone trying to sell you a used car. The first things you want to know are the total miles/kilometers on the odometer and the year the car was manufactured. Since each of these measures is independent of its failure/repair history, the universal first-order model for a car is the NHPP. You also want to know about major repairs and overhauls, so this model is not “exact”. However, the RP, which often is claimed to be *the* system model, is absurd in this scenario; if a salesman tried to convince you that a 10-year-old car was only 2 days old because it had been repaired 2 days earlier, you would seek a car elsewhere!

The HPP (usually and erroneously called *exponentiality* by both practitioners and theorists) is often portrayed as *the* system model, based on Drenick’s [6] asymptotic theorem. However, systems often do not operate long enough to have such asymptotic results hold, even to a reasonable approximation. For example, a 10-year-old car modeled by an HPP would have no effective age, not even the 2 days since the last repair. The salesman, therefore, could claim that it was “brand new”!

Fortunately, the NHPP is a very tractable model as well. For example, the superposition of independent NHPPs (SNHPP) is an NHPP, so – unlike the RP – the model for the SNHPP is known. In addition, the NHPP is much more tractable than the RP or the SRP.

The NHPP is not an exact model for a system, but it often is a good working approximation. Even far more complicated models ignore most – or all – of the 18 shortcomings of probabilistic modeling of systems listed in [7].

Statistical Analysis

The interarrivals times between failures of a system appear in natural order on a time line. The *first* step in an analysis, therefore, is to test for trend. If a trend toward larger interarrivals times (**reliability growth** or improvement) or toward smaller interarrivals times (deterioration) exists, the interarrivals times are not identically distributed. Hence, little or nothing in many reliability and statistics books is applicable, since many books concentrate solely on i.i.d. data. If an assumption is dropped, it usually is that of independence, rather than the even more important assumption of identically distributed interarrivals times.

Parametric Models

Given a trend, the NHPP is the first choice as a model. Consider the power-law process where $v(t) = \gamma\beta t^{\beta-1}$.

Then if failures occurred at arrival times $T_1 = t_1, T_2 = t_2, \dots, T_m = t_m$, over an observation interval $(0, t^*]$, the maximum-likelihood estimators (MLEs) of β and γ are

$$\hat{\beta} = \frac{m}{\sum_{i=1}^m \ln(t^*/t_i)} \quad (4)$$

$$\hat{\gamma} = \frac{m}{\hat{\beta} t^*} \quad (5)$$

If observation is till time T_m , then T_m is substituted for t^* in both equations. Many other techniques for the power-law process and other NHPPs are provided in [8] and the references therein. Lawless and Thiagarajah [9, 10] show how to include explanatory factors or covariates, e.g. different environmental stresses experienced by different system copies, in an analysis.

If there is no evidence of trend, the data can be considered to be identically distributed, but the interarrivals times of a copy may still be dependent. In practice, however, one seldom has enough failures to test the independence assumption [11].

If there is no evidence that the interarrivals times are not i.i.d., one fits an RP to them. If an exponential distribution provides adequate fit, the resulting model is the HPP; otherwise, a more general RP must be selected. The techniques (a) for fitting an exponential distribution to times to failure of parts and (b) for

fitting an HPP to system interarrivals times are interchangeable. The *interpretations* of the results, however, often are drastically different. This is because a failed part is discarded at failure, whatever the magnitude of its constant FOM, whereas a failed system modeled by an HPP is repaired to the ROCOF it had when it was new, which may be very large. For further information on the major differences in interpretations, see [12, 13.]

Nonparametric Models

Cumulative Plots. If we plot cumulative failures *versus* cumulative operating time, t , for a single system copy, we can perform trend testing visually. An increasing, constant, or decreasing slope is an indication of deterioration, noncommittal, or improving reliability, respectively, of that system copy. Deteriorating (improving) systems have become known as *sad* (*happy*) systems to help distinguish these concepts from wear out (**burn-in**) of parts.

Mean Cumulative Function. In most applications, data are available from two or more copies and in many cases the number of copies is far greater than two. The mean cumulative function (MCF) is constructed incrementally at each instant a failure occurs by considering the number of copies at risk at that instant. Copies may not be in the risk set because of left or right censoring. Left censoring occurs when a copy is operated for some time without knowledge of its failure history and then comes under observation. Right censoring occurs when a copy is no longer observed after a time $t = t^*$ for any reason, ranging from the copy being “totaled” to simply not accumulating more than t^* hours at a time when the number at risk is being calculated. The key reference to this approach is [14].

Misconceptions and Their Causes

When parts are tested to failure, the i.i.d. assumptions are plausible and greatly simplifying. When the times between successive failures in a socket are analyzed, the i.i.d. assumptions, leading to the RP, are plausible but no longer tractable. For a system, the RP is neither plausible nor tractable *nor even desirable*; one hopes that successive interarrivals times tend to become larger, e.g., through reliability growth testing

or from improved operating/maintenance procedures. It is amazing, therefore, that some reliability texts still totally ignore **repairable systems** or assume that the RP is *the* only model for such a system. Many texts do not provide techniques for testing for nonstationarity and do not even hint that such trend tests are essential.

Even when the RP is not assumed to be the universal system model, sometimes it seems to be assumed, when it is not. For example, [15] does not restrict itself to the RP, but also uses the word *renewal* as a synonym for repair. Hence a “renewal” may return a system to the same-as-new condition of an RP – but it may also leave the system in *some other condition* after the repair. As discussed in [13], renewal sometimes is infinitely better described as “bad as new”, rather than the virtually universally used term, *good as new*.

It is widely believed, even among theorists, that it is too difficult to apply nonstationary models. However, the NHPP is much more tractable than the RP. For example, the renewal function, i.e. the expected number of failures over $(0, t]$, is unavailable in closed form, except for a few special cases; the corresponding quantity under an NHPP is $\int_0^t v(z) dz$, and one selects a function for v , for which the integral is known. Or, compare the simple MLE, equation (4), for the shape parameter of the power-law process with the messy transcendental MLE equation for the shape parameter of the Weibull distribution, which must be solved by trial and error.

It is also widely believed that an RP whose distribution has increasing (decreasing) FOM can model deterioration (improvement) of a system. The main reason for this misconception is that the FOM, $h_x(x)$, is usually called *failure rate*, especially by theorists. A natural interpretation of “increasing failure rate” is an increasing number of failures per unit time, see e.g. [16], but increasing FOM does *not* imply that. Consider, for example, the model for a system where $X_i \sim U(0, i]$, $i = 1, 2, 3, \dots$. Then each X_i has an FOM which strictly increases from i^{-1} to infinity, but the number of system failures per unit time decreases asymptotically to zero as $i \rightarrow \infty$. Even for parts that are tested to failure, increasing FOM does *not* imply an increasing number of part failures per unit time. (Consider a group of a million people all born in say, 1907. By now the number of deaths per year is stochastically *decreasing* rapidly since very few from the original million are still at risk of dying.)

The FOM should not be called *failure rate*, to avoid assigning flagrantly counterintuitive meanings to the misnomer “failure rate”.

There are other reasons for confusion between FOM and ROCOF. For example, under an NHPP the ROCOF of the process is equal to – but not equivalent to – the FOM of the distribution of time to first failure of the process. Therefore, under an HPP both FOM and ROCOF are equal to the same constant. However, when the area under any FOM increases to infinity, one failure will occur with probability one. In infinite contrast, when the area under any ROCOF increases to infinity, the expected number of failures increases to infinity. Hence, it is essential to distinguish between FOM and ROCOF, rather than to emphasize their superficially striking – but relatively unimportant – similarities.

Confusion about the distinction between FOM and ROCOF has led to the erroneous belief that there is only one bathtub curve. In the bathtub curve for parts, FOM is plotted against part age x . In the curve for a system, ROCOF is plotted against the total operating time t of the system, regardless of whether one or more failures have occurred. In practice, both curves usually are depicted as $\lambda(t) = \text{“failure rate”}$ plotted against “ t ”. This makes it impossible to tell which bathtub curve is portrayed just from the plot. Some authors refer to only one of these interpretations, and some refer to both as if they were equivalent, see e.g. [17]. In many cases, however, widespread poor and incorrect terminology and notation make it *impossible* to ascertain *which* bathtub curve is being discussed, even from the text accompanying the plot.

One further point about the system bathtub curve: the increasing ROCOF for large t implies that $v(t) \rightarrow \infty$ as $t \rightarrow \infty$. On the other hand, asymptotic theorems such as Drenick’s imply that $v(t)$ approaches a finite constant as $t \rightarrow \infty$. Theorists use the asymptotic theorems to justify the HPP, whereas practitioners erroneously assume that the bottom of the system bathtub curve implies an HPP. Neither group is aware of the other’s rationale.

For a system the interarrivals times are the times between failures. For parts, the spacings between order statistics *also* can be interpreted as “times between failures”. Moreover, when parts are put on test (or into operation) simultaneously and run continuously, their “times between failures” appear *exactly* in real time; because of nonzero repair times, this is never exactly true for the interarrivals times of

systems. It is essential to clearly distinguish between the two kinds of “times between failures”.

Most of the listed misconceptions and inherent ambiguities, but especially the use of the misnomer, “failure rate” for both FOM and ROCOF, caused great confusion in [18], as acknowledged in [19]. For example, two obviously different plots obtained from the same failure numbers, one actually estimating FOM and the other actually estimating ROCOF, were claimed to be plots of the *same* “failure rate”. Ascher [20] summarized the problem and its causes and showed how Juran and Gryna [21] also confused themselves by using the term *failure rate* interchangeably for both FOM and ROCOF.

The many inherent ambiguities in analyzing repairable systems failure data make it essential to avoid avoidable ambiguities, such as (a) using “*t*” for both the age of a part and for the total operating time of a system and (b) using “failure rate” for both FOM and ROCOF. Unfortunately, to the present writing, most practitioners have followed the lead of most theorists in embracing both types of ambiguities.

References

- [1] Cinlar, E. (1975). *Introduction to Stochastic Processes*, Prentice Hall, Englewood Cliffs, p. 88.
- [2] Barlow, R. & Proschan, F. (1975). *Statistical Theory of Reliability and Life Testing: Probability Models*, Holt, Rinehart and Winston, New York, pp. 161–162.
- [3] Zacks, S. (1992). *Introduction to Reliability Analysis: Probability Models and Statistical Methods*, Springer-Verlag, New York, pp. 67–83.
- [4] Barlow, R. & Proschan, F. (1975). *Statistical Theory of Reliability and Life Testing: Probability Models*, Holt, Rinehart and Winston, New York, pp. 245–253.
- [5] Barlow, R. & Proschan, F. (1975). *Statistical Theory of Reliability and Life Testing: Probability Models*, Holt, Rinehart and Winston, New York, pp. 246–251.
- [6] Drenick, R. (1960). The failure law of complex equipment, *Journal of the Society for Industrial Applications of Mathematics* **8**, 680–690.
- [7] Ascher, H. & Feingold, H. (1984). *Repairable Systems Reliability: Modeling, Inference, Misconceptions and their Causes*, Marcel Dekker, New York, pp. 63–69.
- [8] Rigdon, S. & Basu, A. (2000). *Statistical Methods for the Reliability of Repairable Systems*, John Wiley & Sons, New York.
- [9] Lawless, J. (2003). *Statistical Models and Methods for Lifetime Data*, 2nd Edition, Wiley-Interscience, Hoboken.
- [10] Lawless, J. & Thiagarajah, K. (1966). A point-process model incorporating renewals and time trends, with application to repairable systems, *Technometrics* **38**, 131–138.
- [11] Ascher, H. & Feingold, H. (1984). *Repairable Systems Reliability: Modeling, Inference, Misconceptions and their Causes*, Marcel Dekker, New York, pp. 88–91.
- [12] Ascher, H. & Feingold, H. (1984). *Repairable Systems Reliability: Modeling, Inference, Misconceptions and their Causes*, Marcel Dekker, New York, pp. 144–145, 151, 160.
- [13] Ascher, H. (2007). Different insights for improving part and system reliability obtained from *exactly same* DFOM “failure numbers.”, *Reliability Engineering and Systems Safety* **92**, 552–559.
- [14] Nelson, W. (2003). *Recurrent Events Data Analysis for Product Repairs, Disease Recurrences, and Other Applications*, SIAM, Philadelphia; ASA, Alexandria.
- [15] Ansell, J. & Phillips, M. (1989). Practical problems in the statistical analysis of reliability data, *Applied Statistics: Journal of the Royal Statistical Society, Series C* **38**, 205–231.
- [16] Elsayed, E. (1996). *Reliability Engineering*, Addison Wesley, Longman, Reading, p. 541.
- [17] Hoyland, A. & Rausand, M. (1994). *System Reliability Theory: Models and Statistical Methods*, John Wiley & Sons, New York, pp. 6, 23–24, 487.
- [18] Frees, E. (1986). Nonparametric renewal function estimation, *Annals of Statistics* **14**, 1366–1378.
- [19] Frees, E. (1998). Correction: nonparametric renewal function estimation, *Annals of Statistics* **16**, 1741.
- [20] Ascher, H. (1999). A set-of-numbers is NOT a data-set, *IEEE Reliability Transactions* **48**, 135–140.
- [21] Juran, J. & Gryna, F. (1988). *Quality Control Handbook*, 4th Edition, McGraw-Hill, New York.

Related Articles

Age-Dependent Minimal Repair and Maintenance; Analysis of Recurrent Events from Repairable Systems; General Minimal Repair Models; Imperfect Repair, Counting Processes; Intensity Functions for Nonhomogeneous Poisson Processes; Multivariate Imperfect Repair Models; Nonparametric Methods for Analysis of Repair Data; Reliability Growth Modeling; Repairable Systems: Statistical Inference; Repairable Systems: Bayesian Analysis; Software Failure Data Analysis; System Availability.

HAROLD E. ASCHER

Reliability of Redundant Systems

Introduction

Redundancy models have been a part of reliability theory almost since the inception of reliability theory. Perhaps the simplest reliability with redundancy model is the dynamic redundancy model. A dynamic redundant system often consists of several modules but with only one module that is required to operate at a time for the system to function. Various fault detection schemes are used to determine when a module has become faulty, and the fault location is used to determine exactly which module, if any, is faulty. If a fault is detected and located, the faulty module is removed from operation and replaced by a spare. The section titled “Dynamic Redundant Systems” discusses various dynamic redundant systems, including fault detection and fault recovery, noncoherent and dual working modes. The section titled “Redundant Systems with Multiple Failure Modes” discusses redundant systems with multiple failure modes, and finally in the section titled “Load-Sharing Redundant Systems”, a brief load-sharing redundant system will be discussed.

Dynamic Redundant Systems

Systems with Imperfect Coverage

Suppose the system consists of $(S + 1)$ modules where S is the total number of spare modules. The modules could be a processor in a multiprocessor system, a generator in an electrical power system, or a logic gate. Examples of duplex (when $S = 1$) configurations are the UDET 7116 telephone switch system [1], the COMTRAC railway traffic controller [2], and the Bell ESS [3]. It is assumed that the lifetimes of the modules are independent and identically distributed (i.i.d.), both modules and the system are of two states, i.e., either successful or unsuccessful (failed) and the reliability of each module, active or spare, is p . The reliability of the dynamic redundant system with S spares is given by

$$R(S) = 1 - (1 - p)^{S+1} \quad (1)$$

The above reliability function is obtained assuming that the fault detection and the switch-over mechanism are perfect. This implies that the function $R(S)$ is, of course, an increasing function of the number of spare modules. On the one hand, adding spare modules makes it more likely that the reliability of the system will increase, but on the other hand, and in practice, the dynamic redundant system may not be able to use the redundancy because it cannot identify that a module is faulty, remove that fault module, and replace it with a fault-free one.

Let f be the probability that the fault detection and switching mechanisms operate correctly, given the failure of one of the modules. The reliability of a dynamic redundant system with fault coverage is given by

$$R_f(S) = \sum_{i=1}^{S+1} \binom{S+1}{i} p^i [(1-p)f]^{S+1-i} \quad (2)$$

It can be shown that [4], for fixed p and f , the maximum value of $R_f(S)$ is attained at $S^* = \lfloor S_1 \rfloor + 1$ where

$$S_1 = \frac{\ln \left(\frac{(1-p)(1-f)}{1-(1-p)f} \right)}{\ln \left(\frac{(1-p)f}{p+(1-p)f} \right)} \quad (3)$$

If S_1 is an integer, S_1 and $S_1 + 1$ both maximize $R_c(S)$.

Fail-Safe Redundancy

This section discusses a fail-safe system (FSS) consisting of a fail-safe redundancy (FSR) core with an associated bank of S spare units in which when one of the $(N - 1)$ operating units fails, the redundant unit immediately starts operating [5]. The failed unit is then replaced by a spare, which then becomes the new redundant unit restoring the FSR core to the all-perfect state. The system survives provided at least $(N - 1)$ out of $(N + S)$ i.i.d units are operable. Applications of such systems can be found in the areas of safety monitoring and reactor trip function systems. For example, in the areas of safety monitoring systems, nuclear plants often use three identical and independently functioning counters to monitor the radioactivity of the air in ventilation systems, with the aim to initiating reactor shutdown

2 Reliability of Redundant Systems

when a dangerous level of radioactivity is present. When two or more counters register a dangerous level of radioactivity, the reactor automatically shuts down.

Suppose that each unit costs c dollars and a system failure costs d dollars. The average total system cost $E(T_s)$ is

$$E(T_s) = c(N + S) + d[1 - R(S)] \quad (4)$$

where the reliability of FSS is

$$R(S) = \sum_{i=N-1}^{N+S} \binom{N+S}{i} p^i (1-p)^{N+S-i} \quad (5)$$

and p is a unit reliability. The average total system cost of system size $(N + S)$ is the cost incurred when the system has failed plus the cost of all $N + S$ modules.

Pham and Galyean [5] show that for given values of p , N , d , and c , the optimal value S^* such that the FSS minimizes the average total system cost is as follows:

- (1) if $f(S_0) < A$ then $S^* = 0$
- (2) if $f(S_0) \geq A$ and $f(0) \geq A$ then $S^* = S_a$
- (3) if $f(S_0) \geq A$ and $f(0) < A$ then

$$S^* = \begin{cases} 0 & \text{if } E(T_0) \leq E(T_{S_a}) \\ S_a & \text{if } E(T_0) > E(T_{S_a}) \end{cases} \quad (6)$$

where

$$S_0 = \max \{0, [S_1] + 1\}; \quad (7)$$

$$S_1 = \frac{(N+1)(1-p) - 3}{p}; \quad (8)$$

$$A = \frac{c}{dp^{N-1}(1-p)^2}; \quad (9)$$

$$S_a = \inf \{S \in [S_0, \infty): f(S) < A\} \quad (10)$$

$$f(S) = \binom{N+S}{S+2} (1-p)^S \quad (11)$$

Consider a cooling water system for a commercial nuclear power plant (either service water or component cooling) typically consisting of a two-drain system, each with a dedicated pump. The system will usually have a third “swing” pump that provides an automatic backup to either of the two

trains. Failure to maintain both cooling water trains operating cold will result in a plant shutdown and require the utility to buy replacement power at a cost of \$1 000 000 per day. The cost of a redundant pump is estimated at \$20 000. The constant failure rate (from random, independent faults) of the pumps is estimated at $\lambda = 10^{-4} \text{ h}^{-1}$, a 30-day pump postulated system failure (i.e., failure of the second pump), if it were to occur, is assumed midway in the replacement/repair time of the first failed pump (i.e., $T_r/2$), and would result in the plant being shut down for 15 days. In this example, we obtain $N = 3$ pumps, $c = \$20\,000/\text{pump}$, $d = \$1\,000\,000 \times 15 \text{ days} = \1.5×10^7 ; $\lambda = 10^{-4} \text{ h}^{-1}$; $T_r = 30 \text{ days}$; $p = e^{-\lambda T_r} = 0.93$.

This can be easily shown as $S^* = 1$.

A Noncoherent System

A k -to- l -out-of- n system is a noncoherent system in which “no fewer than k and no more than l ” out of its n components are to function for the successful operation of the system. This class of noncoherent systems was first proposed by Heidtmann [6]. A noncoherent system is any system that is not coherent; a system is coherent if and only if: (a) it is good (bad) when all components are good (bad), and (b) it never gets worse (better) when the number of good components increases (decreases). Such a system is noncoherent because the system reliability can decrease when the number of good components increases. The generality of the k -to- l -out-of- n system can be demonstrated as follows: (a) when $k = l = n$, the k -to- l -out-of- n system becomes a series system; (b) when $k = 1$ and $l = n$, the k -to- l -out-of- n system becomes a parallel system.

The k -to- l -out-of- n system takes into account the failures caused by the underproduction of components or services and the failures that are caused by overproduction as well [7]. Examples of noncoherent systems can be found in multiprocessor and communication [8]. A multiprocessor system, for example, shares resources such as memory, input/output units, and buses among n processors. If less than k processors are being used, the system does not work to its maximum capacity and the system efficiency is poor. On the other hand, if more than l processors are being used then efficiency again is poor because of the traffic congestion caused by the too many processors in use. Both of these cases need to be avoided; this can

be accomplished by modeling the system as failed for these cases. This multiprocessor system is thus a k -to- l -out-of- n system.

Consider a system composed of n i.i.d. components; assume that both components and the system are of two states, i.e., either successful or failed and p is the reliability of component, then the reliability of the k -to- l -out-of- n system is given by

$$R(k; l, n) = \sum_{i=k}^l \binom{n}{i} p^i (1-p)^{n-i} \quad (12)$$

It can be shown that for given values of k , l , and p , there exists an optimal value n^* for n that maximizes the system reliability $R(n)$. The value n^* is given by [7] as

$$n^* = \inf \{n : f(n) < a^{l+1-k}\} \quad (13)$$

where

$$f(n) = \frac{\binom{n}{k-1}}{\binom{n}{l}} \quad (14)$$

and

$$a = \frac{p}{1-p} \quad (15)$$

Consider a computer installation where a certain number of terminals are to be connected to a noncoherent system. It is of interest to find out what is the optimal number of terminals that maximizes the mean profit of the terminal subsystem. If the number of terminals is greater than, say l , then the response time might not be tolerable, causing a loss of profit. This is overperformance of work. If, on the other hand, the number of terminals is less than, say k , then the computer is not being used to its maximum capacity. This is underperformance of work with a lower profit. Both cases are to be avoided.

The mean system profit, $M(n)$, is defined as follows [9]:

$$\begin{aligned} M(n) &= \sum_{i=0}^n \binom{n}{i} p^i (1-p)^{n-i} f_n(i) \\ &= E \{f_n(S_n)\} \end{aligned} \quad (16)$$

where

$f_n(i)$ profit when i units are functioning

S_n binomially distributed random variable with parameters p , n , and $f_n(i) < f_n(j)$ and $f_n(r) < f_n(j)$ for $0 < i < k$, $k \leq j \leq l$, $l < r \leq n$

Pham [9] showed that for fixed k , l , and p , if the profit function f_n is concave, then there is an optimal value of n such that $M(n)$ is maximized.

Given $k = 5$, $l = 8$ and the profit function

$$f_n(i) = \begin{cases} 100i & \text{for } 0 \leq i < k \\ 500 & \text{for } k \leq i \leq l \\ 500 - 10(i-l) & \text{for } l < i \leq n \end{cases} \quad (17)$$

the optimal number of units to maximize the mean system profit is given by 9 for $p = 0.9$. The mean system profit corresponding to this optimum value is 496. For a lower value of p e.g. $p = 0.7$, the optimal number of units, 11, is higher. If p is close enough to 1, then the optimal n is very close to l .

Degradable Redundant Systems with Dual Working Modes

In many practical applications with changing working environments, the reliability of the components perhaps will change subject to working conditions. In fact, there is a minimum number of units that must function to guarantee the functioning of the system, especially in fault-tolerant computing systems. Consider a network of electric generators, for example, for which at least 80% of all generators must work to ensure a certain service level of the entire station under normal working conditions, called *good mode*. However, under those stress working conditions, called **degradation mode** at least 90% of all generators must work to ensure a full service level of the station. In the latter case, owing to a higher workload of each generator, the reliability of working generator will be decreased.

Several models and methods have been proposed in the literature to evaluate the reliability of fault-tolerant systems [10, 11], degradable computing systems [12, 13], interconnected power systems [14], weighted one-failure mode models [15], two-failure modes [16, 17], and weighted two-failure modes [18]. Not much research, however, has been published on degradable systems with dual working modes. Nordmann and Pham [19] studied the degradable systems with dual working modes; however,

4 Reliability of Redundant Systems

the requirements for the number of working units in each mode are the same as the ***k-out-of-n*** systems. Differing from such existing models and formulations in the literature, Pham [20] recently has studied a degradable system with dual working modes in which he introduces an option that the *k-out-of-n* system can lead to an *m-out-of-n* system subject to a system degradation mode under a stressed condition.

In this study the degradable redundant system with dual working modes in consideration consists of *n* units. The system and the units are either in a successful state (working) or in a failed state (not working). The system operates in a good mode with probability α and in a degradation mode with probability $1 - \alpha$. For the system to be working in a good mode, at least *k* out of *n* units must work in which the reliability of the working units is p_1 . For the system to be working in a degradation mode, at least *m* out of *n* units must work with the reliability of the working units is p_2 . Units in degradation mode are generally put under more stress than under normal conditions. This generally occurs when units such as power substations have to take extra loads (i.e., more power) from other unavailable (or failed) units that fail to keep the system in a working system state. On the basis of such applications, the number of *m* working units (number of units required for the system to be working in a degradation mode) is considered to be greater than or equal to *k* (the number of units required for the system to be working in a good mode). In other words, it is assumed that $m \geq k$ and $p_1 \geq p_2$.

The system functions in a good mode if and only if a minimum number of *k* out of *n* units must work. Similarly, the system functions in a degradation mode if and only if at least *m* out of *n* units must work. The reliability of the system with dual working modes is given by [20]:

$$R(n) = \alpha \left[\sum_{i=k}^n \binom{n}{i} p_1^i (1 - p_1)^{n-i} \right] + (1 - \alpha) \left[\sum_{i=m}^n \binom{n}{i} p_2^i (1 - p_2)^{n-i} \right] \quad (18)$$

where $k \leq m \leq n$ and $p_1 \geq p_2$.

Given $k = 5, m = 7, \alpha = 0.7, p_1 = 0.5$ and $p_2 = 0.4$ for example. Table 1 presents the reliability calculations for various values of *n*.

Table 1 Reliability values vs. *n*

<i>n</i>	<i>R(n)</i>
8	0.2569
9	0.3575
10	0.4525
15	0.7757

Redundant Systems with Multiple Failure Modes

In some applications, a component can be subjected to failure in either open or closed modes [21]. Networks of relays, fuse systems for warheads, diode circuits, fluid flow valves, and so on are a few examples of such components. As discussed in previous sections, redundancy can be used to enhance the reliability of a system without any change in the reliability of the individual components that form the system. However, in a two-failure mode problem, redundancy may either increase or decrease the system's reliability. For example, a network consisting of *n* relays in parallel has the property that a closed-mode failure of any one relay would cause a system closed-mode failure, and an open-mode failure of all *n* relays would cause an open-mode failure of the system. Therefore, adding components to the system may decrease the system reliability. (The designations "closed mode" and "short mode" both appear in this chapter, and these two terms will be used interchangeably.)

System reliability where components have various failure modes is covered in [16, 17, 21–24]. This section briefly discusses the aspects of the **reliability optimization** of systems subject to two types of failure. Reliability optimization of several configurations such as parallel-series, series-parallel, and *k-out-of-n* systems subject to two types of failure will be discussed.

Consider a series system consisting of *n* components. In this series system, any one component failing in an open mode causes system failure, whereas all components of the system must malfunction in short mode for the system to fail. The reliability of series system is given by:

$$R_s(n) = (1 - q_0)^n - q_s^n \quad (19)$$

where *n* is the number of identical and independent components. For given values of q_0 and q_s , it can

be shown that there exists an optimum number of components, say n^* , that is

$$n^* = \begin{cases} \lfloor n_0 \rfloor + 1 & \text{if } n_0 \text{ is not an integer} \\ n_0 \text{ or } n_0 + 1 & \text{if } n_0 \text{ is an integer} \end{cases} \quad (20)$$

where

$$n_0 = \frac{\log\left(\frac{q_0}{1-q_s}\right)}{\log\left(\frac{q_s}{1-q_0}\right)} \quad (21)$$

that maximizes the reliability of series system.

Consider a parallel system consisting of n components. For a parallel configuration, all the components must fail in open mode or at least one component must malfunction in short mode to cause the system to fail completely. The system reliability is

$$R_p(n) = (1 - q_s)^n - q_0^n \quad (22)$$

where n is the number of components connected in parallel. Define

$$n_0 = \frac{\log\left(\frac{q_s}{1-q_0}\right)}{\log\left(\frac{q_0}{1-q_s}\right)} \quad (23)$$

For given values of q_0 and q_s , the parallel system reliability $R_p(n^*)$ is maximum for

$$n^* = \begin{cases} \lfloor n_0 \rfloor + 1 & \text{if } n_0 \text{ is not an integer} \\ n_0 \text{ or } n_0 + 1 & \text{if } n_0 \text{ is an integer} \end{cases} \quad (24)$$

Consider a system of components arranged so that there are m subsystems operating in parallel, each subsystem consisting of n i.i.d. components in a series. Such an arrangement is called a *parallel-series arrangement*. Define q_0 be the open-mode failure probability of each component ($p_0 = 1 - q_0$) and q_s the short-mode failure probability of each component ($p_s = 1 - q_s$). The system can be (a) good, (b) failed open (at least one component in each subsystem fails open), or (c) failed short (all the components in any subsystem fail short). The system reliability is given by

$$R_{ps}(n, m) = (1 - q_s^n)^m - [1 - (1 - q_0)^n]^m \quad (25)$$

Let n , q_0 , and q_s be fixed. The maximum value of $R_{ps}(m)$ is attained at $m^* = \lfloor m_0 \rfloor + 1$, where

$$m_0 = \frac{n(\log p_0 - \log q_s)}{\log(1 - q_s^n + \log(1 - p_0^n))} \quad (26)$$

If m_0 is an integer, then m_0 and $m_0 + 1$ both maximize $R_{ps}(m)$.

The series-parallel structure to be discussed is the dual of the parallel-series structure above. It is a system of components arranged in such a way that there are m subsystems operating in a series with each subsystem consisting of n i.i.d. in parallel. Such an arrangement is called a *series-parallel arrangement*. The series-parallel system reliability is given by

$$R(m) = (1 - q_0^n)^m - [1 - (1 - q_s^n)]^m \quad (27)$$

Barlow *et al.* [24] show that there exists no pair (m, n) maximizing system reliability. However, for fixed n , q_0 , and q_s , however, one can determine the value of m that maximizes the system reliability $R(m)$ where $m^* = \lfloor m_0 \rfloor + 1$, and

$$m_0 = \frac{n(\log p_s - \log q_0)}{\log(1 - q_0^n - \log(1 - p_s^n))} \quad (28)$$

If m_0 is an integer, then m_0 and $m_0 + 1$ both maximize $R(m)$.

The k -out-of- n Systems. Consider a model in which a k -out-of- n system is composed of n identical and independent components that can be either good or failed. The components are subject to two types of failure: failure in an open mode and failure in a closed mode. The system can fail when k or more components fail in a closed mode or when $(n - k + 1)$ or more components fail in an open mode. The k -out-of- n system reliability is given by

$$R(k, n) = \sum_{i=0}^{k-1} \binom{n}{i} q_s^i p_s^{n-i} - \sum_{i=0}^{k-1} \binom{n}{i} p_0^i q_0^{n-i} \quad (29)$$

For fixed k , q_0 , and q_s , the maximum value of $R(k, n)$ is attained at $n^* = \lfloor n_0 \rfloor$ where

$$n_0 = k \left[1 + \frac{\log\left(\frac{1-q_0}{q_s}\right)}{\log\left(\frac{1-q_s}{q_0}\right)} \right] \quad (30)$$

If n_0 is an integer, both n_0 and $n_0 + 1$ maximize $R(k, n)$. For any given q_0 and q_s , we can easily obtain

the following bounds for the optimal value of the threshold k : $nq_s < k^* < np_0$

This result shows that for given values of q_0 and q_s , an upper bound for the optimal threshold k^* is the expected number of components working in the open mode, and a lower bound for the optimal threshold k^* is the expected number of components failing in the closed mode.

Weighted Systems with Multiple Failure Modes.

In many applications, ranging from target detection to pattern recognition, including safety-monitoring protection, undersea communication, and human organization systems, a decision has to be made on whether or not to accept the hypothesis, based on the given information, so that the probability of making a correct decision is maximized. In safety monitoring protection systems, e.g., in a nuclear power plant, where the system state is monitored by a multichannel sensor system, various core groups of sensors monitor the status of neutron flux density, coolant temperature at the reaction core exit (outlet temperature), coolant temperature at the core entrance (inlet temperature), coolant flow rate, coolant level in pressurizer, and on-off status of coolant pumps. Hazard-preventive actions should be performed when an unsafe state is detected by the sensor system.

In undersea communication and decision-making systems, the system consists of n electronic sensors with each scanning for an underwater enemy target [25]. Some electronic sensors, however, might falsely detect a target when none is approaching. Therefore, it is important to determine a threshold level that maximizes the probability of making a correct decision. All these applications have the following working principles in common. (a) System units make individual decisions; thereafter, the system as an entity makes a decision based on the information from the system units. (b) The individual decisions of the system units need not be consistent and can even be contradictory; for any system, rules must be made on how to incorporate all information into a final decision. System units and their outputs are, in general, subject to different errors, which in turn affect the reliability of the system decision.

This section presents a brief summary of recent studies in reliability analysis of systems with three failure modes. Pham [16] studied a dynamic redundant system with three failure modes. Each unit is subject to stuck-at-0, stuck-at-1, and stuck-at- x failures. The system outcome is either good or failed.

Focusing on the dynamic majority and k -out-of- n systems, Pham derived **optimal design** policies for maximizing the system reliability. Nordmann and Pham [26] have presented a simple algorithm to evaluate the reliability of weighted dynamic-threshold voting systems, and Nordmann and Pham [18] also presented a general analytic method for evaluating the reliability of weighted-threshold voting systems. It is worth considering the reliability of weighted voting systems with time dependency.

Load-Sharing Redundant Systems

In reliability, it is common practice to use redundancy techniques to improve system reliability. In most cases, when analyzing redundancy, independence is assumed across the components within the system. In other words, it is assumed that the failure of a component does not affect the failure properties (failure rates) of the remaining components. In the real world, however, many systems are **load sharing**, where the assumption of independence is no longer valid. In a load-sharing system, if a component fails, the same workload has to be shared by the remaining components, resulting in an increased load shared by each surviving component [27]. Load-sharing models have a wide range of engineering applications. They are extensively used in the textile engineering (fibers) [28], material science and testing (fatigue and crack growth), mechanical engineering, and civil and structural engineering (welded joint on large support structures) disciplines. Further, load-sharing models have interesting applications in nuclear reactor safety [27], **software reliability** [29], and condition-based maintenance [30].

Now the author will briefly discuss the pattern of loads acting on the system, load-sharing policies that dictate how the load on the system is distributed among the components, and the relationship between the dynamic load and the failure behavior of the components over a period of time. In load-sharing systems, the components of the system are subjected to varying loads for numerous reasons, including changing demands of the system, failure of some components, changes in the operating conditions and so on. The most important element of the load-sharing model is the rule that governs how the loads on the working components change after some components in the system fail. For example:

- Equal load-sharing rule. A constant system load is distributed equally among the working components. Examples of this kind include yarn bundles and untwisted cables that spread the stresses uniformly after individual failures.
- Local load-sharing rule. A load on a failed component is transferred to adjacent components, and the proportion of the load that the surviving components inherit depends on their “distance” to the failed component. Examples of this kind include cables supporting bridges and other structures, composite materials with bounding matrix joins, and transmission systems that are modeled using consecutive k -out-of- n systems.

In load-sharing systems, it is worth to note that a component operates at different loads during different time intervals. Therefore, it is important to consider the effects of load history on the component failure rate.

References

- [1] Morganti, M.G., Coppadoro, G. & Ceru, S. (1878). UDET 7116 common control for PCM telephone exchange: diagnostic software design and availability evaluation, in *Proceedings of the International Symposium on Fault-Tolerant Computing*, pp. 16–23.
- [2] Ihara, H., Fukuoka, K., Kubo, Y. & Yokota, S. (1978). Fault-tolerant computer system with three symmetric computers, *Proceedings of the IEEE* **66**(10), 1160–1177.
- [3] Toy, W.N. (1978). Fault-tolerant design of local ESS processors, *Proceedings of the IEEE* **66**(10), 1136–1145.
- [4] Pham, H. & Pham, M. (1992). Reliability analysis of dynamic redundant systems with imperfect coverage, *Reliability Engineering and System Safety Journal* **35**(2), 173–176.
- [5] Pham, H. & Galyean, W.J. (1992). Reliability analysis of nuclear fail-safe redundancy, *Reliability Engineering and System Safety Journal* **37**(2), 109–112.
- [6] Heidtmann, K.D. (1981). A class of noncoherent systems and their reliability analysis, in *Proceedings of the 11th Annual Symposium on Fault Tolerant Computing*, June 1981, pp. 96–98.
- [7] Upadhyaya, S.J. & Pham, H. (1993). Analysis of noncoherent systems and an architecture for the computation of the system reliability, *IEEE Transactions on Computers* **42**(4), 484–493.
- [8] Pham, H. (1991a). Cost optimization of a class of noncoherent systems, *Mathematical and Computer Modelling Journal* **15**(6), 1–7.
- [9] Pham, H. (1991b). Optimal design for a class of noncoherent systems, *IEEE Transactions on Reliability* **40**(3), 361–363.
- [10] Beutry, M. (1978). Performance related reliability for computer systems, *IEEE Transactions on Computers* **27**, 540–547.
- [11] Furchtgott, D. & Meyer, J. (1984). A performability solution method for degradable non-repairable systems, *IEEE Transactions on Computers* **33**, 550–554.
- [12] Goyal, A. & Tantawi, A.N. (1987). Evaluation of performability for degradable computer systems, *IEEE Transactions on Computers* **36**, 738–744.
- [13] Wu, L.T. (1982). Operational modes for the evaluation of degradable computing systems, in *Proceedings of the ACM/SIGMETRICS Conference on Measurement Modeling Computing Systems*.
- [14] Gubbala, N. & Singh, C. (1994). A fast and efficient method for reliability evaluation of interconnected power systems – preferential decomposition method, *IEEE Transactions on Power Systems* **9**(2), 113–118.
- [15] Wu, J. & Chen, R. (1994). An algorithm for computing the reliability of a weighted k -out-of- n system, *IEEE Transactions on Reliability* **43**(3), 327–328.
- [16] Pham, H. (1999). Reliability analysis for dynamic configurations of systems with three failure modes, *Reliability Engineering and System Safety* **63**, 13–23.
- [17] Sah, R.K. & Stiglitz, J.E. (1988). Qualitative properties of profit-making k -out-of- n systems subject to two kinds of failures, *IEEE Transactions on Reliability* **37**(6), 515–520.
- [18] Nordmann, L. & Pham, H. (1999). Weighted voting systems, *IEEE Transactions on Reliability* **48**(1), 42–49.
- [19] Nordmann, L. & Pham, H. (1995). Performability and cost analysis of degradable systems, *International Journal of Reliability, Quality and Safety Engineering* **2**(3), 291–298.
- [20] Pham, H. (2007). *Performability Analysis of Degradable Systems with Dual Working Modes*, Working Paper, Rutgers Department of Industrial and Systems Engineering.
- [21] Pham, H. & Malon, D.M. (1994). Optimal designs of systems with competing failure modes, *IEEE Transactions on Reliability* **43**, 251–254.
- [22] Page, L.B. & Perry, J.E. (1988). Optimal series-parallel networks of 3-stage devices, *IEEE Transactions on Reliability* **37**, 388–394.
- [23] Pham, H. & Pham, M. (1991). Optimal designs of $(k, n - k + 1)$ out-of- n : F systems (subject to 2 failure modes), *IEEE Transactions on Reliability* **40**, 559–562.
- [24] Barlow, R.E., Hunter, L.C. & Proschan, F. (1963). Optimum redundancy when components are subject to two kinds of failure, *Journal of the Society for Industrial and Applied Mathematics* **11**(1), 64–73.
- [25] Pham, H. (1997). Reliability analysis of digital communication systems with imperfect voters, *Mathematical and Computer Modelling Journal* **26**, 103–112.
- [26] Nordmann, L. & Pham, H. (1997). Weighted voting human-organization systems, *IEEE Transactions on Systems, Man, and Cybernetics – Part A* **30**(1), 543–549.

8 Reliability of Redundant Systems

- [27] Pham, H. (1992). Reliability analysis of a high-voltage power system with dependence failure and imperfect coverage, *Reliability Engineering and System Safety* **37**, 25–28.
- [28] Daniels, H.E. (1945). The statistical theory of the strength bundles of threads I, *Proceedings of the Royal Society of London. Series A* **183**, 405–435.
- [29] Pham, H. & Wang, H. (2001). A quasi renewal process for software reliability and testing costs, *IEEE Transactions on Systems, Man, and Cybernetics – Part A* **31**(6), 623–631.
- [30] Li, W. & Pham, H. (2005). An inspection-maintenance model for systems with multiple competing processes, *IEEE Transactions on Reliability* **54**(2), 318–327.

HOANG PHAM

Stress–Strength Model

Introduction

Let X and Y be two continuous random variables with cumulative distribution functions $F(x)$ and $G(y)$, respectively. Let Y denote the strength of a component subject to stress X . Then the component fails if the applied stress or load is greater than its strength or resistance. The *reliability* of the component is given by

$$R_1 = P(X < Y) \quad (1)$$

The above *stress-strength model* has been useful in a number of areas, especially in the structural and aerospace industries. As an example consider the following problem. A solid propellant rocket engine is successfully fired, provided the chamber pressure (X) generated by ignition stays below the burst pressure (Y) of the rocket engine. If $X > Y$, the engine blows up and the operation is a failure. The probability of failure is

$$P_1 = 1 - R_1 = P(X > Y) \quad (2)$$

From practical considerations it is desirable to draw an inference about the reliability function R_1 . The problems of estimating R_1 have been considered by many using nonparametric, parametric, and Bayesian approach. We present a survey of some of these results.

For a bibliography of relevant literature see Basu [1–4], Basu and Ebrahimi [5], Basu and Tarmast [6], Basu and Thompson [7], Thompson and Basu [8], Sun, Ghosh and Basu [9], Johnson [10] and the references given there.

Nonparametric Approach

Let $X_1, X_2, \dots, X_m$ be a random sample of size m of stresses and $Y_1, Y_2, \dots, Y_n$ be a random sample of size n of strength. The Mann–Whitney statistic U is

$$U = \sum_{i=1}^m \sum_{j=1}^n U_{ij} \quad (3)$$

and

$$U_{ij} = \begin{cases} 1, & \text{if } X_i < Y_j \\ 0, & \text{otherwise} \end{cases} \quad (4)$$

Here, U is the number of pairs (X_i, Y_j) with $X_i < Y_j$. The point estimate

$$\hat{R}_1 = \frac{U}{mn} \quad (5)$$

was first proposed by Birnbaum [11]. $\hat{R}_1$ is the uniform minimum variance unbiased estimator of R_1 . Note that since probability of failure $P_1 = P(X > Y) = 1 - R_1 = 1 - P(X < Y)$, we can either estimate R_1 or P_1 . That is, $\hat{P}_1 = 1 - \hat{R}_1$.

Birnbaum obtained one-sided **confidence interval** for P_1 for the case where F is known, and G unknown ($m \rightarrow \infty$), and both F and G are unknown. Birnbaum and McCarty [12] considered a numerical procedure for computing the sample sizes needed for the confidence interval based on $\hat{R}_1$.

Owen *et al.* [13] showed that the assumption of continuity required in Birnbaum [11] was not essential and produced some tables for use in computing sample sizes needed and computing confidence intervals for the Birnbaum–McCarty procedure. Govindarajulu [14] has also explicitly derived one-sided and two-sided distribution-free confidence bounds for R_1 based on the asymptotic normality of

$$\hat{P}_1 = 1 - \frac{U}{mn} = \sum_{i=1}^m \sum_{j=1}^n \frac{(1 - U_{ij})}{mn} \quad (6)$$

These bounds are approximately one half of the corresponding bounds due to Birnbaum and McCarty. In particular, for all F and G and large m or n , the solution of ϵ of the equation $P(R_1 \leq \hat{R}_1 + \epsilon) = P(R_1 \geq \hat{R}_1 - \epsilon) \geq \gamma$ is $t_1 = (2r)^{-1/2} \Phi^{-1}(\gamma)$, $0 < \gamma < 1$ and the solution of ϵ of the equation $P(|\hat{R}_1 - R_1| \leq \epsilon) \geq \gamma$, $0 < \gamma < 1$, is given by $t_2 = (4r)^{-1/2} \Phi(1 + \gamma/2)$. Here $\Phi(x) = 1/\sqrt{2\pi} \int_{-\infty}^x e^{-u^2/2} du$, and $\Phi^{-1}(\cdot)$ is the inverse function of $\Phi(\cdot)$.

Johnson [10] has also given a survey of some of these results.

Parametric Approach

In many situations, the family of the distribution of X and/or Y will be known, and it is desirable to obtain parametric solutions. We consider a number of cases for estimating R_1 for specific parametric distributions.

2 Stress–Strength Model

Normal Distribution

Owen *et al.* [13], gave one-sided confidence intervals for R_1 when (a) X and Y follow a bivariate **normal distribution** and observations (X, Y) are in pairs, or (b) when X and Y are independent normal with a common unknown variance. Govindarajulu [14] obtained two-sided confidence intervals for R_1 . Basu [3] and Johnson [10] present some of these results.

Let us assume that $f(\cdot)$ is $N(\mu_1, \sigma_1^2)$ and $g(\cdot)$ is $N(\mu_2, \sigma_2^2)$. Then

$$R_1 = P(X < Y) = \Phi\left(\frac{\mu_2 - \mu_1}{\sqrt{\sigma_1^2 + \sigma_2^2}}\right) \quad (7)$$

where Φ is the standard normal **cdf**. The maximum-likelihood estimator of R_1 is

$$\hat{R}_1 = \Phi\left(\frac{\bar{Y} - \bar{X}}{\sqrt{\hat{\sigma}_1^2 + \hat{\sigma}_2^2}}\right) \quad (8)$$

Here $\bar{X} = \frac{1}{m} \sum X_i$, $\bar{Y} = \frac{1}{n} \sum Y_i$, and $\hat{\sigma}_1^2 = \sum (X_i - \bar{X})^2 / m$, $\hat{\sigma}_2^2 = \sum (Y_i - \bar{Y})^2 / n$.

For large m and n , an approximate $100(1 - \alpha)\%$ lower confidence bound for R_1 is given by

$$R_1 = \Phi\left(\frac{\bar{y} - \bar{x}}{\sqrt{\hat{\sigma}_1^2 + \hat{\sigma}_2^2}}\right) - \frac{z_\alpha \hat{\sigma}_{R_1}}{\sqrt{m + n}} \quad (9)$$

where $1 - \alpha = \Phi(z_\alpha)$.

In the special case of equal variances, $\sigma_1^2 = \sigma_2^2$, the lower confidence bound on R_1 is given by Johnson [10].

Exponential Distribution

Since in many physical situations, especially in reliability and life testing problems the exponential distribution is considered a more realistic model, we shall consider estimating R_1 when X and Y follow independent exponential distributions with density functions

$$f(x) = \frac{1}{\theta_1} e^{-x/\theta_1}, \quad x > 0, \quad \theta_1 > 0 \quad (10)$$

and

$$g(y) = \frac{1}{\theta_2} e^{-y/\theta_2}, \quad y > 0, \quad \theta_2 > 0 \quad (11)$$

respectively. Here

$$R_1 = P(X < Y) = \frac{\theta_2}{\theta_1 + \theta_2} \quad (12)$$

Given two independent samples $(X_1, X_2, \dots, X_m)$ and $(Y_1, Y_2, \dots, Y_n)$ from $f(\cdot)$ and $g(\cdot)$, respectively, the maximum-likelihood estimator of R_1 is given by

$$\hat{R}_1 = \frac{\bar{Y}}{(\bar{X} + \bar{Y})} = \frac{1}{(\bar{X}/\bar{Y} + 1)} \quad (13)$$

Since $\hat{R}_1$ is a function of $\bar{X}/\bar{Y}$, the exact distribution of $\hat{R}_1$ can be obtained. Note $(\bar{X}\theta_2)/(\bar{Y}\theta_1)$ follows the F distribution with $(2m, 2n)$ **degrees of freedom**, a $100(1 - \alpha)\%$ lower confidence bound on R_1 is given by

$$\frac{1}{\left(1 + \frac{\bar{X}}{\bar{Y}} F_{2n, 2m}^{(\alpha)}\right)} < R_1 \quad (14)$$

where $F_{2n, 2m}^{(\alpha)}$ is the upper α th point of the $F_{2n, 2m}$ distribution.

Basu [3] has also considered the cases where the underlying distributions are (a) γ with known shape parameters, (b) Weibull with known common shape parameter, and (c) the case where (X, Y) follow the bivariate exponential distribution of Marshall and Olkin [15].

Bayesian Analysis of Complex Systems

Enis and Geisser [16] and Johnson [10] have considered the Bayesian inference for R_1 where the underlying distributions are normal. We shall here primarily consider the case when X and Y follow independent exponential distributions. The conditional density functions, conditional on the parameters θ_1 and θ_2 , are

$$f(x|\theta_1) = \frac{1}{\theta_1} \exp(-x/\theta_1) \quad (15)$$

and

$$g(y|\theta_2) = \frac{1}{\theta_2} \exp(-y/\theta_2) \quad (16)$$

for $(x, y, \theta_1, \theta_2 > 0)$. Then

$$\begin{aligned} R_1(\theta_1, \theta_2) &= P(X < Y | \theta_1, \theta_2) \\ &= \theta_2 / (\theta_1 + \theta_2) \\ &= \left(\frac{\lambda}{1 + \lambda} \right) \end{aligned} \quad (17)$$

where $\lambda = \theta_2 / \theta_1$. We assume that the parameters θ_1 and θ_2 are independently distributed. Then the prior density of θ_1 and θ_2 is $p(\theta_1, \theta_2) = p_1(\theta_1)p_2(\theta_2)$ where $p_i(\cdot)$ is prior density of θ_i , $i = 1, 2$. Given independent random samples $X_1, X_2, \dots, X_m$ and $Y_1, Y_2, \dots, Y_n$ form $f(\cdot | \theta_1)$ and $g(\cdot | \theta_2)$, we can obtain the posterior distribution of (θ_1, θ_2) , R_1 , and λ as given by Basu and Ebrahimi [5]. We shall now consider Bayesian estimates of reliability for complex systems.

k-out-of-p Systems

Let us consider a *k-out-of-p* system where a system has p components and the system functions if and only if at least k of these p components function successfully. When $k = p$ (or $k = 1$) we obtain a series (or parallel) system. If $k = p = 1$, the system is as described in the section titled “Introduction”. Basu and Tarmast [6] describe such systems. First, let us consider a system of strength Y which is subjected to p different stresses $X_1, X_2, \dots, X_p$. An example of interest is the case where a beam of strength Y is subjected to p different stresses. We assume X_i ’s and Y are independently distributed where the common conditional density of X_i ’s is exponential $f(\cdot | \theta_1)$ and Y also follows the exponential distribution with conditional density $g(\cdot | \theta_2)$. Basu and Tarmast [6] show that the system reliability is

$$\begin{aligned} R_2 &= P(\text{at least } k \text{ of the } X_i \text{'s} < Y) \\ &= \frac{\Gamma(p+1)\Gamma\left(p + \frac{1}{\lambda} + 1 - k\right)}{\left\{ \Gamma(p+1-k)\Gamma\left(p + \frac{1}{\lambda} + 1\right) \right\}} \end{aligned} \quad (18)$$

Next consider a p -component system with strengths $Y_1, Y_2, \dots, Y_p$ respectively, and each component is subjected to the same stress X . Here X and Y_i ’s follow independent exponential distributions with $X \sim f(\cdot | \theta_1)$ and $Y_i \sim g(\cdot | \theta_2)$. Then the system reliability

is given by

$$R_3 = 1 - \frac{\Gamma(k+\lambda)\Gamma(p+1)}{\Gamma(p+\lambda+1)\Gamma(k)} \quad (19)$$

Finally, consider a more general p -component system where the i th component of strength Y_i is subject to stress (shock) X_i , $i = 1, 2, \dots, p$. Assuming that the X_i ’s and Y_i ’s are independent exponential distributions with conditional densities of X_i and Y_i are $f(\cdot | \theta_1)$ and $g(\cdot | \theta_2)$, respectively.

Then the reliability of the k -out-of- p system is given by

$$R_4 = \sum_{j=k}^p \binom{p}{j} \frac{\lambda^j}{(1+\lambda)^p} \quad (20)$$

Note R_1 is a special case of R_2 , R_3 , and R_4 when $p = k = 1$. Here

$$R_1 = P(X < Y) = R_2 = R_3 = R_4 = \frac{\lambda}{1 + \lambda} \quad (21)$$

Bayesian Estimation Using Noninformative Priors

Let us assume that X and Y follow independent exponential distributions. A reasonable noninformative prior distribution for θ_i is given by

$$h(\theta_i) = \frac{1}{\theta_i}, \theta_i > 0 (i = 1, 2) \quad (22)$$

Basu and Tarmast [6] obtain the following Bayesian estimates for R_2 , R_3 , and R_4 .

$$\begin{aligned} \hat{R}_2 &= \frac{\Gamma(m+n)\Gamma(p+1)}{\Gamma(m)\Gamma(n)\Gamma(p+1-k)} \\ &\times \int_0^1 \left(\frac{p+1+k+\frac{1-y}{uy}}{\Gamma\left(p+1+\frac{1-y}{uy}\right)} \right) y^{m-1}(1-y)^{n-1} dy \end{aligned} \quad (23)$$

$$\begin{aligned} \hat{R}_3 &= 1 - \frac{\Gamma(m+n)\Gamma(p+1)}{\Gamma(m)\Gamma(n)\Gamma(k)} \\ &\times \int_0^1 \frac{\Gamma\left(k+\frac{uy}{1-y}\right) y^{m-1}(1-y)^{n-1}}{\Gamma\left(p+1+\frac{uy}{1-y}\right)} dy \end{aligned} \quad (24)$$

and

$$\hat{R}_4 = \sum_{j=\kappa}^p \binom{p}{j} \frac{\Gamma(m+n)u^j}{\Gamma(m)\Gamma(n)} \times \int_0^1 \frac{y^{m+j-1}(1-y)^{n-p-j-1}}{(1+(u-1)y)^p} dy \quad (25)$$

Here $u = (n\bar{y}/m\bar{x})$.

Basu and Ebrahimi [5] and Basu and Tarmast [6] have also considered Bayesian estimates of R_i ($i = 1, 2, 3, 4$) when the θ_i 's have independent conjugate **prior distributions**. Sun *et al.* [9] have considered Bayesian estimation of R_1 using noninformative priors. Basu and Thompson [7] Thompson and Basu [8] also derived the posterior densities of the reliabilities of exponential stress-strength systems using subjective and noninformative priors.

References

- [1] Basu, A. (1977a). Estimate of reliability in the stress-strength model, in *Proceedings of the 22nd Conference of the Design of Experiments in Army Research Development and Testing*, Adelphi, Maryland, pp. 97–110.
- [2] Basu, A. (1977b). Reliability estimation for complex systems, in *Proceedings of the ARO Workshop on Reliability and Probabilistic Design*, Watervliet, New York, pp. 85–99.
- [3] Basu, A. (1981). The estimation of $P(X < Y)$ for distributions useful in life testing, *Naval Research Logistics Quarterly* **28**, 383–392.
- [4] Basu, A. (1985). Estimation of reliability of complex systems—a survey, in *The Frontiers of Modern Statistical Inference Procedures*, E.J. Dudewicz, ed, American Science Press, pp. 271–287.
- [5] Basu, A. & Ebrahimi, N. (1991). Bayesian approach to life testing and reliability estimation using asymmetric loss function, *Journal of Statistical Planning and Inference* **29**, 21–31.
- [6] Basu, A. & Tarmast, G. (1987). Reliability of a complex system from Bayesian viewpoint, in *Probability and Bayesian Statistics*, R. Viertl, ed, Plenum Publishers, pp. 31–38.
- [7] Basu, A. & Thompson, R. (1987). Bayesian estimation of stress-strength systems using reference priors. Proceedings of the Section on Bayesian Statistical Science, *American Statistical Association* 170–172.
- [8] Thompson, R. & Basu, A. (1993). Bayesian reliability of stress-strength systems, in *Advances in Reliability*, A. Basu, ed, North Holland, Amsterdam, pp. 411–421.
- [9] Sun, D., Ghosh, M. & Basu, A. (1988). Bayesian analysis for a stress-strength system under noninformative priors, *The Canadian Journal of Statistics* **2**, 323–332.
- [10] Johnson, R. (1988). Stress-strength models for reliability, in *Handbook of Statistics, Volume 7: Quality Control and Reliability*, P.R. Krishnaiah & C.R. Rao, eds, Elsevier Science, New York, pp. 27–54.
- [11] Birnbaum, Z. (1956). On a use of the Mann-Whitney statistic, in *Proceedings of the 3rd Berkeley Symposium*, Berkeley, California, Vol. 1, pp. 13–17.
- [12] Birnbaum, A. & McCarty, R. (1958). A distribution free upper confidence bound for $\Pr\{Y < X\}$ based on independent sample of X and Y. *Annals of Mathematical Statistics* **29**, 558–562.
- [13] Owen, D., Craswell, R. & Hanson, D. (1964). Non-parametric upper confidence bounds for $\Pr\{Y < X\}$ and confidence limits for $\Pr\{Y < X\}$ when X and Y are normal. *Journal of the American Statistical Association* **59**, 906–924.
- [14] Govindarajulu, Z. (1967). Distribution free confidence bounds for $P(X < Y)$, *Annals of the Institute of Statistical Mathematics* **20**, 229–238.
- [15] Marshall, A. & Olkin, I. (1967). A multivariate exponential distribution, *Journal of the American Statistical Association* **62**, 30–44.
- [16] Enis, P. & Geisser, S. (1971). Estimation of the probability that $Y < X$, *Journal of the American Statistical Association* **66**, 162–168.

ASIT P. BASU

Reliability Optimization

Introduction

A reliability engineer's job is to make systems (hardware, software, human) as reliable as possible. In this article, we look at how systems can be made as reliable as possible, under the constraints that are imposed. We consider three approaches to reliability optimization.

1. Optimize the reliability of the components that make up the system. This can be done through the use of experimental design to find the parameters of the design or the parameters of the process that produce the components.
2. Redundancy can be used to increase system reliability (*see Reliability of Redundant Systems; Parallel, Series, and Series-Parallel Systems*).
3. Preventive maintenance can be done to extend the lifetimes of key components.

These three areas are the topics of the next sections. Other approaches to optimizing reliability can be found in Kuo *et al.* [1].

Optimizing Component Reliability

Meeker and Hamada [2] provide an overview of strategies for optimizing **component reliability**, including

1. life testing
2. **degradation** testing
3. designed experiments
4. monitoring field and warranty data.

We will focus on the use of designed experiments for achieving component reliability. Experimental design has been used for quite a while to choose the parameters of a product or process that optimize the resulting quality. See Box *et al.* [3] and Montgomery [4]. Until recently, most work has been done under the assumption that the output is a normally distributed random variable. One problem with applying these methods in reliability is that the output of a life test is a lifetime, and the normal distribution is usually a poor model for lifetimes. Another potential problem is the presence of censoring, that is the termination

of a life test before all units have failed. Hamada [5] describes likelihood-based methods for experiments in which the output is a lifetime and censoring is possible. Depending on the amount of censoring, the maximum-likelihood estimates of the factor effects might not exist. Hamada and Wu [6] present a Bayesian approach, for which factor effects always exist.

One example discussed in both Hamada [5] and Hamada and Wu [6], originally due to Taguchi [7], involves the lifetimes of a fluorescent lamp. The experiment involved five factors, which we will call A , B , C , D , and E . Each factor had two levels, denoted by “−” and “+”. Two lamps were made at each of eight different treatments (i.e., combinations of factors). The experimental design was a one-fourth replication of a 2^5 **factorial** design, denoted 2^{5-2} . These eight runs were chosen in such a way that all combinations of three variables gives a full 2^3 factorial design. The fourth and fifth variables, called D and E , are set according to the defining relationships $D = -AC$ and $E = -BC$. Under this assignment, the 2^{5-2} design has resolution III, meaning that main effects are confounded with two-factor interactions. The lamps were inspected periodically, so the exact lifetimes were unknown. Since all that was observed was that a light was burning at one time, t_1 , and not at the next inspection time, t_2 , it was known only that the failure occurred in the interval (t_1, t_2) . This is a special type of censoring called *interval censoring*. The last inspection was done at time 20, and 7 of the 16 lamps were still burning at this time. This is the usual situation of type I censoring. From the analyses in Hamada [5] and Hamada and Wu [6], the authors concluded that the significant factors, in order of significance, were D , B , E , and A . With this knowledge, it was possible to alter the process or product design so as to maximize the lifetime of the lamp.

Reliability Allocation

It may be possible to increase component reliabilities, but at a cost. Suppose we have an n -component system with component reliabilities, denoted $\mathbf{r} = (r_1, r_2, \dots, r_n)$. The system reliability is then a function of $\mathbf{r}$,

$$R_s = f(r_1, r_2, \dots, r_n) = f(\mathbf{r}) \quad (1)$$

Achieving a reliability of r_j for component j involves a cost $g(r_j)$. There are constraints on the total cost

2 Reliability Optimization

(B) and there may be lower and upper bounds (l_j and u_j) on each component's reliability. The constrained optimization problem is therefore

$$\begin{aligned} \text{Maximize : } & R_s = f(r_1, r_2, \dots, r_n) \\ \text{Subject to : } & \sum_{j=1}^n g_j(r_j) \leq B \\ & l_j \leq r_j \leq u_j, \quad j = 1, 2, \dots, n \end{aligned}$$

For example, consider an n -component series system with equal component reliabilities, for which

$$R_s = f(r_1, r_2, \dots, r_n) = r^n \quad (2)$$

Suppose that the reliability cost function for each system is

$$g_j(r_j) = \frac{1}{1 - r_j} \quad (3)$$

The total cost is therefore

$$\sum_{j=1}^n \frac{1}{1 - r_j} = \frac{n}{1 - r} \quad (4)$$

For this cost function, the cost goes to infinity as the reliability approaches 1 from the left; in other words, it becomes more and more expensive to increase reliability as the reliability gets close to 1. The optimization problem then becomes

$$\begin{aligned} \text{Maximize : } & r^n \\ \text{Subject to : } & \frac{n}{1 - r} \leq B \\ & 0 \leq r_j \leq 1, \quad j = 1, 2, \dots, n \end{aligned}$$

The solution to this problem is to take r to be as large as possible, while maintaining the constraint $n/(1 - r) \leq B$, that is $r = 1 - n/B$.

For a nontrivial example, consider the three-component system shown in Figure 1, where components 1a and 1b have identical reliability, r_1 , and component 2 has reliability r_2 .

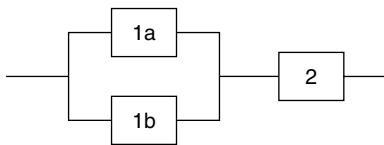


Figure 1 Structure diagram for series-parallel system

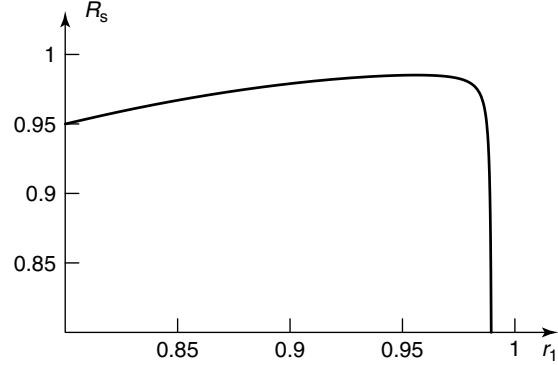


Figure 2 System reliability R_s as a function of r_1

The optimization problem becomes

$$\begin{aligned} \text{Maximize : } & [1 - (1 - r_1)^2] r_2 \\ \text{Subject to : } & \frac{1}{1 - r_1} + \frac{1}{1 - r_2} \leq B \\ & 0 \leq r_j \leq 1, \quad j = 1, 2 \end{aligned}$$

The highest reliability will be attained when the total cost is B , in which case

$$r_2 = \frac{B - 2 - (1 - B)r_1}{1 - B - Br_1} \quad (5)$$

Thus,

$$R_s = [1 - (1 - r_1)^2] \frac{B - 2 - (1 - B)r_1}{1 - B - Br_1} \quad (6)$$

a function of a single variable. Figure 2 shows R_s as a function of r_1 . A numerical procedure gave us the result that R_s attains a maximum when $r_1 = 0.956074$. For this value of r_1 , we have $r_2 = 0.987052$ giving a system reliability of $R_s = 0.985148$.

Only in special cases will we be able to reduce the problem to solving a single equation. Usually we will be left with a nonlinear constrained optimization problem.

Redundancy Allocation

Each component can be backed up by a number of redundant components. Let x_j denote the number of redundant parts for component j . The component

reliabilities are assumed to be fixed. The system reliability is then a function of the number of redundant parts for each component. For example, suppose the system is an n -component series system, having reliability $R_s = r_1 r_2 \cdots r_n$ if there is no redundancy. If component j has x_j redundant parts then the system is a series-parallel system, as shown for illustration in Figure 3 for $n = 4$, $x_1 = 2$, $x_2 = 1$, $x_3 = 2$, and $x_4 = 3$. With redundancy, the system reliability is

$$R_s = \prod_{j=1}^n \left[1 - (1 - r_j)^{x_j+1} \right] \quad (7)$$

Let $g_j(x_j)$ denote the cost of having x_j redundant parts for component j . There may be constraints on the number of redundant parts, for example physical size and space may pose a limit. The optimization problem is therefore

$$\text{Maximize : } R_s = \prod_{j=1}^n \left[1 - (1 - r_j)^{x_j+1} \right]$$

$$\text{Subject to : } \sum_{j=1}^n g_j(x_j) \leq B$$

$$L_j \leq x_j \leq U_j, \quad j = 1, 2, \dots, n$$

This is a nonlinear integer programming problem, because the variables x_j are necessarily integers.

For example, consider a three-component series system. Suppose that the cost functions are

$$g_j(x_j) = (x_j + 1)C_j \quad (8)$$

indicating that the cost is proportional to the number of components (operating and redundant), where $C_1 = 1$, $C_2 = 2$, and $C_3 = 3$. Suppose also that the component reliabilities are $r_1 = r_2 = 0.95$ and

$r_3 = 0.99$, and that the cost constraint is $B = 12$. The corresponding system reliabilities and costs are shown in Table 1 for a number of combinations of (x_1, x_2, x_3) . The entries below the line violate the constraint that the cost must not exceed 12. Among all the combinations whose cost does not exceed 12, the one with the highest reliability is $(x_1, x_2, x_3) = (2, 2, 0)$, giving a system reliability of 0.9898 and a cost of $B = 12$. Thus the optimum is to have the original and two spares for components 1 and 2, and to have the original and no spares for component 3. This makes intuitive sense, because component 3 already has the highest component reliability.

With only three components, the possibilities can be enumerated in a table like the one shown in Table 1. For larger problems, methods of nonlinear integer programming must be used to find the optimum redundancy.

Reliability and Redundancy Allocation

Suppose now that there are costs associated with increasing reliability and adding redundancy. Following the notation of the previous sections, we define

r_j := reliability of component j and

x_j := number of redundant parts for component j

Let $g_j(r_j, x_j)$ be the cost of achieving reliability r_j for each copy of component j , while having x_j redundant parts. We would then pose our optimization problem as follows:

$$\text{Maximize : } R_s = f(r_1, r_2, \dots, r_n)$$

$$\text{Subject to : } \sum_{j=1}^n g_j(r_j, x_j) \leq B$$

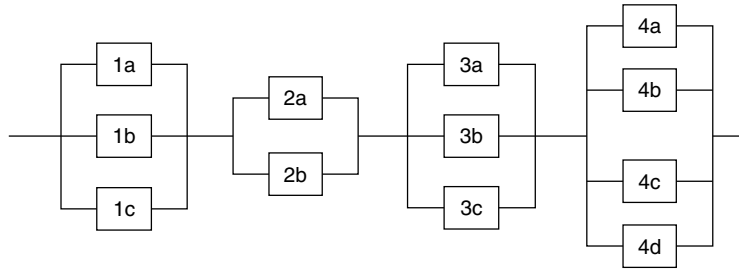


Figure 3 Structure diagram for series system with redundancy

4 Reliability Optimization

Table 1 Component redundancies along with system reliability and cost

x_1	x_2	x_3	R_s	Cost	x_1	x_2	x_3	R_s	Cost	x_1	x_2	x_3	R_s	Cost	x_1	x_2	x_3	R_s	Cost
0	0	0	0.8935	6	0	1	0	0.9381	8	0	2	0	0.9404	10	0	3	0	0.9405	12
0	0	1	0.9024	9	0	1	1	0.9475	11	0	2	1	0.9498	13	0	3	1	0.9499	15
0	0	2	0.9025	12	0	1	2	0.9476	14	0	2	2	0.9499	16	0	3	2	0.95	18
0	0	3	0.9025	15	0	1	3	0.9476	17	0	2	3	0.9499	19	0	3	3	0.95	21
1	0	0	0.9381	7	1	1	0	0.9851	9	1	2	0	0.9874	11	1	3	0	0.9875	13
1	0	1	0.9475	10	1	1	1	0.9949	12	1	2	1	0.9973	14	1	3	1	0.9974	16
1	0	2	0.9476	13	1	1	2	0.995	15	1	2	2	0.9974	17	1	3	2	0.9975	19
1	0	3	0.9476	16	1	1	3	0.995	18	1	2	3	0.9974	20	1	3	3	0.9975	22
2	0	0	0.9404	8	2	1	0	0.9874	10	2	2	0	0.9898	12	2	3	0	0.9899	14
2	0	1	0.9498	11	2	1	1	0.9973	13	2	2	1	0.99965	15	2	3	1	0.99977	17
2	0	2	0.9499	14	2	1	2	0.9974	16	2	2	2	0.99975	18	2	3	2	0.99987	20
2	0	3	0.9499	17	2	1	3	0.9974	19	2	2	3	0.99975	21	2	3	3	0.99987	23
3	0	0	0.9405	9	3	1	0	0.9875	11	3	2	0	0.9899	13	3	3	0	0.99	15
3	0	1	0.9499	12	3	1	1	0.9974	14	3	2	1	0.99977	16	3	3	1	0.99989	18
3	0	2	0.95	15	3	1	2	0.9975	17	3	2	2	0.99987	16	3	3	2	0.99999	21
3	0	3	0.95	18	3	1	3	0.9975	20	3	2	3	0.99987	22	3	3	3	0.99999	24

$$l_j \leq r_j \leq u_j, \quad j = 1, 2, \dots, n$$

$$L_j \leq x_j \leq U_j, \quad j = 1, 2, \dots, n$$

This is a nonlinear mixed integer programming problem. See Kuo *et al.* [1] for methods of solving such optimization problems.

Optimal Maintenance

For a repairable system, the question is often when to replace or overhaul an old system. Replacing (or renewing) the system too often is expensive, but an old deteriorating system may fail often, requiring many repairs. The objective is to choose the time between renewals to minimize the cost. There are many possible cost models for this scenario, but we discuss just a very simple one. Suppose that the failure process is a power law process (*see Repairable Systems: Statistical Inference*) with parameters β and θ . We assume that $\beta > 1$, otherwise the system is improving and it would never be optimal to replace an improving system. Let A denote the cost of replacing an old system with a brand-new one, and let B denote the cost of a repair. Some cost models impose a distribution on the value of B , but we will assume that it is a constant. If we consider the time interval $[0, T]$, where T is large, then the number of replacements will be $\lfloor T/a \rfloor$, occurring at times $a, 2a, 3a, \dots, \lfloor T/a \rfloor a$.

Over the interval $[(k-1)a, ka]$, $k = 1, 2, \dots, \lfloor T/a \rfloor$, there will be N_k failures/repairs, where

N_k has a Poisson distribution with mean $(a/\theta)^\beta$. (Since the times between replacements are constant and nonoverlapping, the N_k are independent and identically distributed (i.i.d.)) The number of failures between the last replacement at time $\lfloor T/a \rfloor a$ and time T has a Poisson distribution with mean $((T - \lfloor T/a \rfloor a)/\theta)^\beta$. The expected cost over the interval $[0, T]$ for the strategy of replacing the system every a time units is

$$C_T(a) = E(\text{Cost}[0, T]) = \left\lfloor \frac{T}{a} \right\rfloor A + \left\lfloor \frac{T}{a} \right\rfloor E(N_1) B + E(\text{number of failures in } [\lfloor T/a \rfloor a, T]) B \quad (9)$$

The long-run average cost is obtained by dividing by T and taking the limit as $T \rightarrow \infty$,

$$\begin{aligned} C_\infty(a) &= \lim_{T \rightarrow \infty} \frac{E(\text{Cost}[0, T])}{T} \\ &= \lim_{T \rightarrow \infty} \left\lfloor \frac{T}{a} \right\rfloor \frac{A}{T} + \lim_{T \rightarrow \infty} \left\lfloor \frac{T}{a} \right\rfloor \frac{B}{T} \left(\frac{a}{\theta} \right)^\beta \\ &\quad + \lim_{T \rightarrow \infty} E(\text{number of failures in } [\lfloor T/a \rfloor a, T]) \frac{B}{T} \\ &= \frac{A}{a} + \frac{B}{a} \left(\frac{a}{\theta} \right)^\beta = Aa^{-1} + \frac{B}{\theta^\beta} a^{\beta-1} \end{aligned} \quad (10)$$

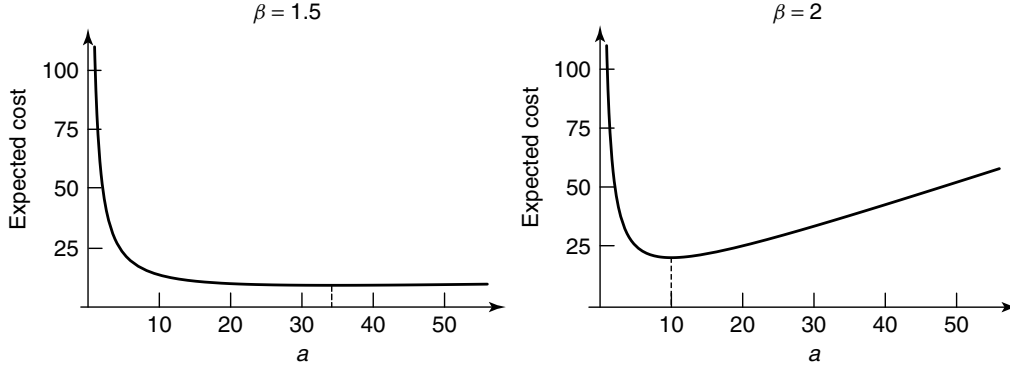


Figure 4 Cost functions $C_{\infty}(a)$ for $\beta = 1.5$ and $\beta = 2$

The minimum cost can be obtained by setting $C'_{\infty}(a) = 0$ and solving for a . This yields

$$a = \theta \left(\frac{A}{B(\beta - 1)} \right)^{1/\beta} \quad (11)$$

Figure 4 shows the cost functions $C_{\infty}(a)$ for $A = 100$, $B = 1$, $\theta = 1$ (i.e. the cost of replacement is 100 times the cost of a repair) and for two different values of β , namely $\beta = 1.5$ and $\beta = 2$. The optimal replacement times are

$$a = 1 \left(\frac{100}{1 \times (1.5 - 1)} \right)^{1/1.5} \approx 34.2 \quad (12)$$

and

$$a = 1 \left(\frac{100}{1 \times (2 - 1)} \right)^{1/2} = 10 \quad (13)$$

Figure 4 indicates that when $\beta = 1.5$ the optimal time between replacements is larger than when $\beta = 2$. This should be expected, because the system is deteriorating faster when $\beta = 2$, so more frequent replacements would be called for.

There are numerous other models for preventive maintenance. See, for example, Wang and Pham [8] and Blischke and Murthy [9].

References

- [1] Kuo, W., Prasad, V.R., Tillman, F.A. & Hwang, C.-L. (2001). *Optimal Reliability Design: Fundamentals and Applications*, Cambridge University Press, Cambridge.
- [2] Meeker, W.Q. & Hamada, M. (1995). Statistical tools for the rapid development & evaluation of high-reliability products, *IEEE Transactions on Reliability* **44**, 187–198.
- [3] Box, G.E.P., Hunter, W.G. & Hunter, J.S. (1978). *Statistics for Experimenters*, John Wiley & Sons, New York.
- [4] Montgomery, D.C. (2005). *Design and Analysis of Experiments*, 6th Edition, John Wiley & Sons, New York.
- [5] Hamada, M. (1995). Using statistically designed experiments to improve reliability and to achieve robust reliability, *IEEE Transactions on Reliability* **44**, 206–215.
- [6] Hamada, M. & Wu, C.F.J. (1995). Analysis of censored data from fractionated experiments: a Bayesian approach, *Journal of the American Statistical Association* **90**, 467–477.
- [7] Taguchi, G. (1980). *System of Experimental Design*, Unipub, New York.
- [8] Wang, H. & Pham, H. (2006). *Reliability and Optimal Maintenance*, Springer, New York.
- [9] Blischke, W.R. & Murthy, D.N.P. (2003). *Case Studies in Reliability and Maintenance*, Wiley, New York.

STEVEN E. RIGDON

Reliability Databases

Introduction

Reliability data consists of information relating to the performance of an item, or a collection of related items (referred to here as *systems*), such as the number of on-demand or time-dependent failures, the length of time taken to repair an item or the operating time between failures. A *reliability data bank* is an organized collection of processed reliability data based on the available raw data. The available raw data is the information that can be extracted from the primary sources. The nature of the primary sources will depend on the specific application, but commonly occurring primary data sources include records of maintenance work requests, measures of system output, and details of necessary repair work.

Recording system output and maintenance allows system owners to monitor performance. The objective of performance monitoring is to learn about the system with a view to predicting future system behavior. The system owner is concerned about the efficient operation of the system, and by improving his understanding of the frequency, and the causes, of system failure, he can act to improve operational efficiency in an informed way. The gathering of reliability data is a necessary step in learning about system performance. The reliability data bank contains the primary data, processed in such a way that the link to the raw data is clear [1].

Reliability databases (RDBs) are defined by Fragola [1], to be:

“A pre-selected, processed output produced and usually available in hard copy form designed to be useful for a variety of analyses whose general type and depths have been anticipated by the database design specification.”

This definition makes clear the dependence of the database on the analyses for which it is intended. There is, therefore, a requirement to consider the different aims of different users when designing an RDB. Cooke and Bedford [2] identify three groups of users of RDBs, and their different aims:

- maintenance engineers interested in measuring and optimizing maintenance performance;
- component designers interested in optimizing component performance;

- risk/reliability analysts wishing to predict reliability of complex systems.

The different aims of these users suggest they are likely to require different types of reliability data. The designer of the RDB should consider these different aims when formatting the database.

Published records of systematically collected reliability data date from 1959 with the *Martin Titan Handbook*. This contained generic failure rates for a large number of electronic components. Other notable early publications include *MIL-Handbook-217*, the *Failure Rate Data Bank*, and the *RADC Non-Electronic Reliability Notebook*. Fragola [1] provides an historical review of the evolution of RDBs. Since the 1980s, there have been attempts to standardize the methods of collecting and organizing raw data. This has resulted in a number of different standards being produced, particularly in relation to data gathering and exchange within the nuclear power industry (ISO6257, ISO7385).

There are a number of well-known RDBs currently in use. The GIDEP (Government-Industry Data Exchange Programme) data bank (maintained by the US military) contains information on the failure rate, failure mode, and replacement rate of items, systems, and subsystems based on field data, and test data. The Center for Chemical Process Safety/American Institute of Chemical Engineers (CCPS/AIChE) process equipment RDB contains reliability data relevant to the chemical and hydrocarbon processing industries. OREDA (Offshore RELiability DATA) is an organization sponsored by nine oil and gas companies that has compiled a data bank of reliability and maintenance data for offshore production equipment. The OREDA model has been accepted as an international standard (ISO14224) for database management in the offshore production industry.

The aforementioned databases contain generic data. This means the data is gathered about items used in a range of different systems and under different operating profiles. These data are then pooled, based on the specification of a range of engineering characteristics that allow the database designer to identify the data as similar in terms of, for example, physical design specification and operating conditions. These data pools may be treated as exchangeable data in some databases, enabling the use of Bayesian updating, but one should be wary of this owing to the fact that the updating

process will narrow the uncertainty range so that it no longer reflects the full range of behaviors seen within the pool. This provides average performance values, which are generally presented in terms of an expected failure rate with upper and lower bounds. RDBs can also be system-specific. System-specific databases concentrate on the failures occurring within a known system. System-specific RDBs should be subject to the same careful design as generic databases to ensure ease of interpretability in the data, although the data for system-specific databases are likely to exhibit greater homogeneity than would be expected in generic databases, thereby reducing the problems posed by grouping the data.

Taxonomies

Taxonomies provide an orderly classification of failure data. Taxonomies should be developed that use device attributes to meaningfully distinguish between the failure behavior of devices and allow investigation of the relationships between devices and their failure modes. For any given RDB, this will depend on the devices under consideration, the systems in which they are used, and the intended use of the database. However, general guidelines can be given on the development of failure and maintenance taxonomies.

Maintenance Taxonomy

Maintenance activities can be classified as either preventive maintenance (PM) or corrective maintenance (CM), which can be scheduled in different ways: calendar based, condition based, opportunity based, and emergency. **Maintenance policies** should be clearly defined, as the choice of policy influences the analysis of the failure data. PM leads to more censored data being recorded, whereas CM methods allow failure modes to be observed.

Preventive maintenance is defined as jobs that do not repair a fault or failure, but are part of regular servicing.

Corrective maintenance is defined as jobs to repair a defect, fault, or failure.

The scheduling types are as follows:

- Calendar based: planned maintenance, not based on observed **degradation**, but scheduled from

some measure of time. This scheduling type applies to PM, and does not involve repair work.

- Condition based: the component's state is observed to deviate from the manufacturer's intended state, but component functionality (as far as is required by the system) is maintained. Therefore, the component can be left until a more suitable maintenance opportunity arises. This form of scheduling involves CM.
- Opportunity based: maintenance is performed when a suitable opportunity occurs. For example, if the system has to be shut down for repairs to certain components, other components can also be maintained. There is an overlap between the concepts of condition-based and opportunity-based maintenance.
- Emergency: actions taken to repair a component that has lost functionality to the extent that the system is disabled. The system cannot function until the component has been repaired and postponing the repair is not feasible. The goal of other maintenance types is to prevent emergency action having to be taken.

Maintenance scheduling can influence the choice of maintenance actions.

Failure Taxonomy

Failure types can be classified as

- Degraded failure: the component is no longer in the state intended by the manufacturer, but is still fulfilling its system function. The component is in a noncritical degraded state.
- Critical failure: the component has ceased to function owing to critical degradation. A critically degraded component is repaired or renewed.
- Entrained functional unavailability: the component loses functionality owing to the failure of a supporting system. Repair is not necessarily required.
- Total failure: the component fails to meet its specifications to any degree.
- Nontotal critical failure: the component ceases to fulfill its function, but retains some residual functionality.

Operating Modes and Failure Mechanisms

The following three operating modes are generally considered:

- Continuous operation: the component is operating when the system is operating.
- Standby: the component is called to function on demand. Otherwise, it remains dormant.
- Alternating: the component rotates sojourns of continuous operation with one or more spares.

The alternating mode describes situations in which the components share workload according to a planned schedule, whereas standby components are generally used when the primary component can no longer function.

The operating mode of the component leads to a classification of component failures. Failures are commonly described as time related or demand related. Demand-related failures occur at the point of a component being called into service (from, for example, having been on standby). Failures occurring during continuous operation are defined to be time related. A component can have more than one operating mode, and can therefore be susceptible to both demand- and time-related failures.

Failure descriptions can be broken down into failure modes, failure mechanisms, and failure causes. Failure causes describe the reason a component fails in terms of flaws within the component, such as cracks or leaks. A failure mechanism is the physical process that leads to the failure cause. Classifying a fault as a failure cause or failure mechanism will often be a subjective process. The failure mode describes the way in which the component has failed, generally in terms of functionality. It is important to specify the failure when the failure consequences depend on how the component fails.

Structure of a Reliability Database

Example 1 A number of databases have been developed over a range of industries with the aim of providing a generic data source for common components. Fleming and Lydell [3] discuss the development of one such RDB – the PIPExp database, which is a source of pipe work reliability data, with a particular emphasis on piping used within nuclear power plants.

The development of this database was in direct response to an absence of collated reliability data on piping. Many RDBs focus on “active” components such as pumps and valves, but large parts of many industrial systems consist of “passive” components, such as piping. However, prior to 1990, most reliability studies related to piping were conducted for a specific application, and efforts were not made to exchange data, or to standardize **data collection** procedures.

Defining failure modes for passive components, such as piping, can be difficult. This is because piping failures are, in general, incipient failures, which result in loss of performance rather than loss of function. It is rare that a pipe suddenly loses structural integrity. The PIPExp database defines the failure modes as shown in Table 1.

These failure modes have been derived from the data validation and verification process to which the raw data has been subjected. The raw data for the PIPExp database consisted of a variety of source material, such as work order systems, regulatory reports, and data reconciliation with existing piping RDBs. The data has been taken from 1970 to the present. The data consists of records of 1750

Table 1 Failure modes within the PIPExp database

Failure mode	Description
Partly cracked	Crack with depth > 10% of wall thickness
Fully cracked	Through crack, but with no leakage
Wall thinning	Internal damage owing to cavitation/erosion
Small leak	Leak rate within a specified acceptable limit
Pinhole leak	As for small leak, but with different underlying damage mechanisms
Large leak	Leak rate exceeds specified limit, but can be safely managed
Severance	Full circumferential crack
Rupture	Large leak rate and major, sudden loss of structural integrity. Generally caused by a combination of internal degradation and overload

4 Reliability Databases

incidents spread over 22 million operating years. This data also allowed the identification of the following degradation mechanisms:

Corrosion
 Thermal fatigue
 Erosion cavitation
 Corrosion fatigue
 Erosion corrosion

 Stress corrosion cracking
 Vibrational fatigue
 Design dynamic loading
 Design and constructional defects

 Water hammer
 Over pressurization/overloading
 Human error
 Unreported/unknown

Point estimates for the frequency of pipe failure are given by the expression:

$$\lambda_{ijk} = \frac{n_{ijk}}{f_{ijk} N_{ij} T_{ij}} \quad (1)$$

where i is the pipe size index, j is the system index and k refers to the degradation mechanism, so n_{ijk} is the number of failure events for pipe size i in system j subject to degradation mechanism k , f_{ijk} is the fraction of components of size i in system j susceptible to degradation mechanism k , N_{ij} is the number of components belonging to the i, j subset and T_{ij} is the total operating time for those components. These estimates are presented with 5% and 95% confidence bounds, to illustrate the uncertainty estimates.

The analysis of the data is based on a Bayesian model using a Poisson likelihood for the number of observed events, and a “semi-informative” lognormal prior on λ_{ijk} . The stated confidence bounds are estimated from simulated realizations of the posterior distribution.

Analysis of RDB Data

Many different methods have been applied to the analysis of reliability data. A selection of these methods is presented here. The role of analysis is frequently to convert the available raw data into information used within other system or maintenance models. The method of analysis is determined by the

component’s failure behavior. The methods apply to the situation in which a component enters service and remains in service until it fails, with no other terminations in service.

Demand-Related Failures

Demand-related failures can be separated into *degradable* and *nondegradable* categories. A *degradable demand-related failure* is one in which the component is subject to degradation while in the standby state. Causes of degradation such as corrosion or oxidation occur, to some extent, irrespective of the operating mode of a component. Therefore, the analysis should consider this degradation.

Nondegradable demand-related failures, in which the component is considered not to have degraded over time can be modeled using a binomial distribution, in which each demand is considered to be single trial. The probability of failure per demand, p , is assumed to be independent and identical for each demand. Thus, the probability of observing f failures in d demands is given by

$$\begin{aligned} &\Pr(f \text{ failures in } d \text{ demands} | p) \\ &= \binom{d}{f} p^f (1 - p)^{d-f} \end{aligned} \quad (2)$$

The **maximum-likelihood estimation** (MLE) of p , given f failures in d demands is f/d , and the 90% confidence bounds are given by

p_l = largest p such that $\Pr(f \text{ or more failures in } d \text{ demands} | p) \leq 0.05$;
 p_u = least p such that $\Pr(f \text{ or fewer failures in } d \text{ demands} | p) \leq 0.05$.

For the *degradable* case, components are subject to maintenance. Consequently, knowledge of the maintenance procedures is required to determine reliability parameters from the data. This is illustrated using two common maintenance policies.

Under the *test and replace* policy, components are tested at a regular interval, I , and are replaced if they are found to have failed. Given uniform demand, the probability of failure on demand is given by

$$1 - \frac{1 - \exp(-\lambda I)}{\lambda I} \approx \frac{\lambda I}{2} \quad \text{for } \lambda I < 0.1 \quad (3)$$

The *fail and fix* policy allows components to fail and then performs repairs offline. During repair the

component is regarded as unavailable and therefore, any demands received would result in a demand failure. For components with a constant demand rate λ and repair rate μ , the equilibrium unavailability is given by $\frac{\lambda}{\lambda+\mu}$.

Time-Related Failures

Time-related failures are commonly modeled using exponential distributions. That is, the failure time of a component i , denoted by x_i , will have **cdf** $F(x) = 1 - \exp(-\lambda x)$. Given failure data from a collection of components, $x_1, \dots, x_n$, the MLE of λ is given by

$$\hat{\lambda} = \frac{n}{\sum_{i=1}^n x_i} \quad (4)$$

Confidence bounds for this estimate can be found by noting that $z = 2\lambda \cdot \sum_i x_i$ has a χ^2 distribution with $2n$ **degrees of freedom**. For these failures, the 90% confidence bounds are given by

$$\begin{aligned} \lambda_l &= \text{largest } \lambda \text{ such that } \Pr(\text{less than } \sum_i x_i | \lambda) \leq 5\% \\ \lambda_u &= \text{least } \lambda \text{ such that } \Pr(\text{more than } \sum_i x_i | \lambda) \leq 5\% \end{aligned}$$

If we consider a single component, with constant failure rate λ , the number of failures in given time, t , is a **Poisson process** with parameter λ and expected value λt . The MLE for λ is given by n/t .

Bayesian Models

Bayesian methods allow for the combining of different data types *via* hierarchical modeling. The imposition of a hierarchical structure allows more complex relationships to be taken into consideration, and in particular, offers a natural framework for modeling plant-specific and component-specific effects, thereby allowing data from both generic and plant-specific sources to be used in modeling.

A simple example of a one-stage Bayesian model is the Poisson–Gamma model for the rate of failure. The pattern of failures is assumed to follow a Poisson process with intensity λ . The intensity is given by the rate of occurrence of failures (ROCOF), and is itself a random quantity. We model λ as a Gamma distribution with parameters α and β , $\text{Ga}(\lambda|\alpha, \beta)$, the natural conjugate prior for Poisson variables.

If we then observe x failures in time T , our posterior for λ will be given by:

$$\begin{aligned} \Pr(\lambda|x, T) &\propto L(x, T|\lambda) \cdot \Pr(\lambda) \\ &\propto \text{Ga}(\lambda|\alpha + x, \beta + T) \end{aligned} \quad (5)$$

If the pattern of failures does follow a Poisson process with intensity $\tilde{\lambda}$, then $x/T \rightarrow \tilde{\lambda}$ as $T \rightarrow \infty$. Furthermore, our estimate is consistent.

Competing Risks

Competing-risk models are constructed by modeling each of the potential failure modes. The different failure modes are interpreted as competing to cause component failure. Assume we have k risks competing to cause component failure. These will be associated with lifetimes $X_1, \dots, X_k$. For a given realization of these variables, the value $\min\{x_1, \dots, x_k\}$ is the point at which the component fails. For such a realization, we would observe lifetime x_j and failure cause j .

Let $X_1 = X$ and $Z = \min(X_2, \dots, X_k)$. Under this construction, we are treating $X_1 = X$ as the failure mode we wish to learn about and all other failure modes as censoring variables. The naked failure rate, $r_X(t)$, is given by:

$$r_X(t) \equiv \frac{f_X(t)}{R_X(t)} = -\frac{1}{R_X(t)} \cdot \frac{dR_X}{dt} \quad (6)$$

where $R_X = 1 - F_X = \exp\left[-\int_0^t r_X(s) ds\right]$ is the subsurvival function. The naked failure rate describes the rate we could expect to observe for X , if it was never censored by Z , i.e. the rate having eliminated all other competing risks. However, when the other risks are not eliminated, the observed failure rate for X is defined as

$$\text{obr}_X(t) = \lim_{\Delta \rightarrow 0} \left[\frac{\Pr(x_j = X, X \in (t, t + \Delta) | x_j > t)}{\Delta} \right] \quad (7)$$

which coincides with $r_X(t)$ if the risks are independent.

Bounds of the value of R_X can be found [4] by noting that

$\Pr(X \leq t, X \leq Z) \leq \Pr(X \leq t) \leq \Pr(X \wedge Z \leq t)$
 $\Leftrightarrow 1 - F_X^*(t) \geq R_X(t) \geq R_X^*(t) + R_Z^*(t)$ and both $1 - F_X^*(t)$ and $R_X^*(t) + R_Z^*(t)$ are observable quantities. The Peterson bounds give an indication of the

range of possible marginal distributions. However, there can be several different marginal distributions compatible with observable data, however much data we can observe.

There are a number of viable competing-risk models available [2]. Each of these models specify an engineering situation that is being captured in the model. Therefore, model validation for these is not statistical, but engineering based.

Working with RDBs

Example 2 We now illustrate the application of a two-stage Bayesian model to the analysis of generic reliability data. This example follows the work of Bunea *et al.* [5]. Two-stage Bayesian models differ from the one-stage model discussed previously by introducing a further level of modeling to account for the uncertainty introduced by not knowing the “true” prior parameterization. In the analysis of generic reliability data, this type of modeling allows for the aleatory uncertainty associated with different data sources to be modeled, as well as the epistemic uncertainty.

The case considered by Bunea *et al.* [5] is that of the ZEDB data, which is a German-led RDB project containing events and operating times data for components at nuclear power plants in Switzerland, Germany, and the Netherlands. A two-stage hierarchical model is adopted to model the relationship between data from different plants. Data from a specific plant, i , takes the form of events, X_i , and operating/exposure times, T_i , and is modeled using a Poisson likelihood function, parameterized by λ_i :

$$l(X_i, T_i) = \frac{(\lambda_i T_i)^{X_i}}{X_i!} \exp\{-\lambda_i T_i\}, \quad \lambda_i > 0, \quad T_i > 0 \quad (8)$$

A prior distribution is assigned to each plant. The prior determines the value of λ_i and is conditioned on the hyperparameter, Q .

Q is the quantity through which the data between plants is assimilated, and the prior distribution for all i plants is conditioned on the same Q . The underlying probabilistic assumptions are (a) each plant's λ_j and its observable data X_j is independent of every other plant's λ_i given Q , and (b) each plant's observable data X_j is independent of Q and the other plants λ_i , X_i , given λ_j .

The analysis of the ZEDB data utilizes a lognormal prior for the distribution of λ_i :

$$f(\lambda|\mu, \sigma) = \frac{1}{\sigma\lambda\sqrt{2\pi}} \exp\left\{-\frac{(\ln\lambda - \mu)^2}{2\sigma^2}\right\} \quad (9)$$

which means our hyperparameters will be μ and σ . The hyperprior used in the analysis of the ZEDB data is chosen by applying Jeffrey's rule to the parameters μ and σ , resulting in a hyperprior of the form $f(\mu, \sigma) \propto 1/\sigma^2$.

Bunea *et al.* [5] explore the effects of varying the hyperpriors and the priors for this model, concluding that choosing a prior that is not conjugate is, in this case, beneficial to the analysis, as it allows for a truncation of the credible range of the hyperparameters.

Criticisms of RDBs

RDBs have faced a number of criticisms since their use has become more widespread. As the development of RDBs has evolved, some of the flaws found in early RDBs have been corrected. However, there are also a number of problems that have been compounded during development of successive databases. The problems concern both the construction of the database and the use of the data from these sources.

A major criticism of early RDBs was the way in which data was presented. Failure rates for components were often given as point estimates, with no attempt made to quantify the uncertainty surrounding this value. When estimating the failure rate from failure data, (i.e., the number of observed failures and the number of operating hours), the value obtained is an *estimate* of the typical component performance, over which we are uncertain. Recent RDBs have aimed to provide information about this uncertainty through the inclusion of confidence bounds, giving the reader a credible region in which the failure rate could lie.

However, the use of confidence bounds, or, more precisely, the way in which the confidence bounds are calculated, has also been criticized. Confidence bounds have frequently been calculated from a χ^2 distribution based on aggregated data from inhomogeneous populations. This aggregation, or pooling, of data is often conducted with the aim of collating all available data relating to a particular device or component, but overlooks the uncertainty introduced

by the different environmental and functional conditions of the subpopulations. For example, the same design of component may be more susceptible to certain types of failure depending on how it is used. The presentation of a single failure rate does not capture this information.

Attempts have been made to tackle this problem using scaling constants, called *K-factors* or π -factors. If a component is known to be operating under certain conditions, its failure rate is rescaled by the *K-factor* associated with these conditions. While *K-factors* allow component failure rates to be adjusted for different operating conditions, they do not deal with the underlying problem of incorrectly treating the uncertainty. Indeed, the use of *K-factors* compounds the problems of pooling, as inhomogeneous data is used to calculate the base failure rate for a component. The effect of ignoring the uncertainty introduced by different operating conditions is to reduce the calculated confidence bounds. The variance associated with a quantity will decrease as the sample size increases, if we assume we are sampling repeatedly from the same distribution. Pooling data makes this assumption, so it is important to assess prior to pooling that it is not an unreasonable assumption. Historically, RDBs have not always tested this assumption.

Recent databases have attempted to tackle this problem by considering the different amounts of observed variation when pooling data, using Fisher *F*-tests to assess homogeneity of samples, and increasingly, categorizing failures by failure mode. This development allows the separation of information pertaining to time- and demand-related failures and the inclusion of data on degraded and incipient failures. Categorizing the data in this way allows for greater clarity over what is meant by the failure rate data, and offers the possibility of learning about repair rates.

Another issue with generic RDBs is their limited use in troubleshooting existing designs. This has led to some critics calling for the rejection of generic data in favor of a “physics of failure” approach.

Physics of failure concentrates on modeling the potential modes of failure in terms of the underlying physical processes that cause them – such as stress, overloading, corrosion, and so on. The argument for the physics of failure approach is the limited relevance of generic data to systems operating under different conditions to those considered previously, and that physics of failure provides an approach in which the potential causes of failure are easier to identify, making appropriate maintenance easier to schedule. However, the use of generic data does allow for identification of potential failure modes, not accounted for, under physics of failure approaches. As the RDBs contain observed data, the failure modes they record include a range of “freak” failures that may not have been considered using the physics of failure approach. For this reason, the use of generic data sources and physics of failure methods can be viewed as complementary.

References

- [1] Fragola, J.R. (1996). Reliability and risk analysis data base development: an historical perspective, *Reliability Engineering and System Safety* **51**, 125–136.
- [2] Cooke, R. & Bedford, T. (2002). Reliability databases in perspective, *IEEE Transactions on Reliability* **51**, 294–310.
- [3] Fleming, K.N. & Lydell, B.Y. (2004). Database development and uncertainty treatment for estimating pipe failure rates and rupture frequencies, *Reliability Engineering and System Safety* **86**, 227–246.
- [4] Peterson, A.V. (1976). Bounds for a joint distribution function with fixed subdistribution functions: application to competing risks, *Proceedings of the National Academy of Science of the United States of America* **73**, 11–13.
- [5] Bunea, C., Charitos, T., Cooke, R.M. & Becker, G. (2005). Two-stage Bayesian models – application to ZEDB project, *Reliability Engineering and System Safety* **90**, 123–139.

TIM BEDFORD, GAVIN HARDMAN, LESLEY WALLS AND JOHN QUIGLEY

Fault Trees

Introduction

Fault trees are concerned with the representation of *failures*. As the name suggests, a tree structure is formed to describe the failures or faults for a system, from which analysis is carried out. The concept of the fault tree analysis (FTA) technique was developed back in 1961 by H. Watson at the Bell Telephone Laboratories. The analysis method, known as *kinetic tree theory*, was developed in 1970 by Vesely [1] and forms the basis of the technique still in use today. The method provides a visual description of the various causes of a specified system failure in terms of the failure of its components. Many different fault trees may be required to evaluate all the potential failures associated with a system.

To use the FTA technique, the components within the system under analysis must be independent. If the characteristics of the system are such that this assumption is not valid, then an alternative approach is required (refer to the section titled “Related Techniques”).

The fault tree has a specified treelike form, starting at the top with the system failure under analysis. The tree is built downward from this undesired event using deductive logic, a “what can cause this” approach, into its intermediate causes. The key to the visual nature of the approach is the descriptions used to document the events, as the logic is developed. The branches of the tree terminate with component failures (often referred to as basic events). Figure 1 illustrates a typical fault tree structure.

The basic elements of the fault tree are gates and events. Gates describe the fault logic between events. There are a number of different types of gates, each representing a causal relationship. The most common gate symbols are OR and AND (shown in Figure 1). The output of an OR gate, represented by the event described above the gate symbol, will occur if at least one of the input events (represented below the gate symbol) occurs, and for an AND gate then all input events must occur.

The other main element of the fault tree is the event. The most common events are intermediate and basic events (shown in Figure 1). All basic events have a distinct labeling or code on the tree

and are represented by circles. Data, in terms of failure and repair information, is required for these events. It is the probability of these basic event failures that can be used to calculate quantitative measures for the system, i.e., the probability of system failure.

Construction of the Fault Tree

To develop the inputs to each event the *immediate*, *necessary*, and *sufficient* causes need to be defined. This is often referred to as the *immediate cause effect*. Immediate refers to considering the next level of resolution or the next linked components in the system, and hence only taking small steps in the analysis. Necessary means making sure that all failures included are needed to cause the output event and that no redundant failures are included. Sufficient means that no failures are omitted from the input list. The immediate, necessary, and sufficient causes of the top event (TOP) then become subevents or intermediate events, where the same process is repeated until basic event inputs are reached on all branches.

When trying to construct a fault tree it may take a couple of attempts. After each attempt the logic of the diagram needs to be checked against your understanding of the system mechanics. The exact representation of the diagram may be different if drawn by different analysts, but will be equivalent if the logic is the same, namely, if the failure combinations are identical.

Qualitative Analysis

The qualitative phase of FTA is involved with determining the failure combinations that result in the system failure defined at the top of the tree. These failure combinations are known as *cut sets* or *minimal cut sets*. A *cut set* is a list of failure events such that if they occur, then so does the top event. A *minimal cut set* is a list of minimal, necessary, and sufficient conditions for the occurrence of the top event; basically, it is a cut set such that if any component failure event is removed from the combination it will no longer cause the system failure (see **Path Sets and Cut Sets in System Reliability Modeling**).

Establishing the minimal cut sets can then be used to form a logic expression for the top event. The

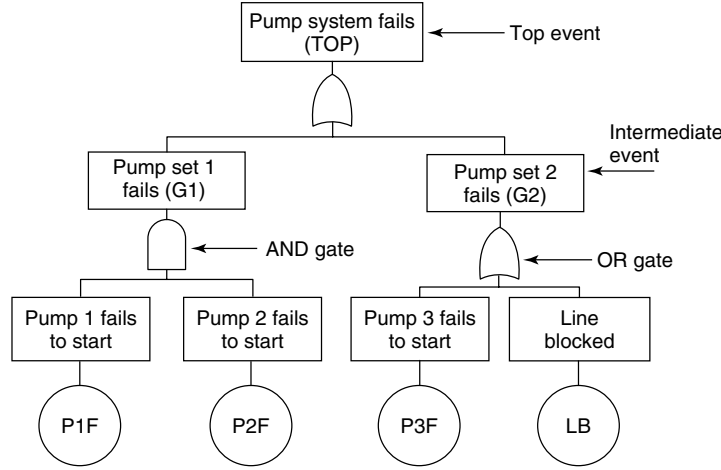


Figure 1 Example fault tree structure

minimal cut set expression for the top event (TOP) is such that $TOP = C1 + C2 + \dots + Cn$ where Ci are the minimal cut sets, with n in total, and “+” is the logical OR.

There are a number of methods that can be used to find these failure combinations, and the one most frequently used is the top-down approach (for other methods see reference [2]). For implementation, laws of Boolean algebra are required. The four laws needed are (a) the distributive law, (b) the idempotent law to remove repeated cut sets, (c) the idempotent law to remove repeated failure events, and (d) the absorption law to remove redundant combinations: Note a dot (.) refers to AND logic.

$$(A + B).(C + D) = A.C + A.D + B.C + B.D \text{ (a) (1)}$$

$$A + A = A \text{ (b) (2)}$$

$$A.A = A \text{ (c) (3)}$$

$$A + A.B = A \text{ (d) (4)}$$

The top-down approach incorporates four main steps:

- Step 1: Start at top of tree.
- Step 2: Express the top event in terms of its inputs using Boolean algebra.
- Step 3: Expand and simplify the expression using laws of Boolean algebra, if possible.
- Step 4: Substitute an expression for each gate working down through the tree and each time repeat step 3.

The fault tree in Figure 1 is used to demonstrate these steps. Step 1 involves starting at the top of the tree, i.e., at TOP. The inputs of TOP, step 2, are two intermediate events G1 and G2. The logical relationship between them is an OR gate, and therefore TOP can be expressed as $TOP = G1 + G2$. No simplification, using the laws of Boolean algebra, is necessary in this example for step 3. Step 4 in the process then considers G1 and G2 in turn, expressing these gates using Boolean algebra. G1 has inputs P1F and P2F. These are related using an AND gate, therefore the expression will involve the dot (.) symbol. Hence $G1 = P1F.P2F$. Substituting this expression to that for TOP gives $TOP = (P1F.P2F) + G2$. Next consider G2. Its inputs are P3F and LB, connected using an OR gate. Therefore, the expression for G2 is $P3F + LB$. Substituting this into the expression for TOP gives $TOP = P1F.P2F + P3F + LB$. No expansion or Boolean algebra simplification is required. Therefore the top event, pumping system failure, will occur if both Pump 1 and Pump 2 fail together, if Pump 3 fails on its own, or if the line is blocked.

Quantification

For a system to be deemed as having attained certain safety standards, a numerical analysis is required. There are a number of system performance measures that can be calculated using the fault tree analysis technique. These include calculating the system probability of failure, frequency of failure, **system**

downtime, and expected number of failures [2, 3]. The focus in this overview is on the approach for calculating the probability of system failure. As the state of components or the system cannot be predicted exactly, the use of **probability theory** is required to predict the likelihood of occurrence or nonoccurrence. In general terms, probability is a measure of the likelihood that a particular event will occur in any one experiment carried out in prescribed conditions. In terms of the reliability analysis of systems, the “particular event” refers to the component failure or system failure, with the “prescribed conditions” being the manner in which they function or are operated. When calculating the exact top event probability, the formula, referred to as the inclusion–exclusion expansion, equation (5), is used whereby successive consideration is given to the probability of failure of additional multiple minimal cut set combinations (an extension of the addition law of probability [4]). The first summation term refers to the addition of the probability of failure of each minimal cut set, the second summation to the addition of the probability of the failure of two minimal cut sets together, and so on.

$$\begin{aligned}
 &P(C_1 \text{ or } C_2 \text{ or } \dots \text{ or } C_n) \\
 &= \sum_{i=1}^n P(C_i) - \sum_{i=1}^{n-1} \sum_{j=i+1}^n P(C_i \text{ and } C_j) \\
 &\quad + \sum_{i=1}^{n-2} \sum_{j=i+1}^{n-1} \sum_{k=j+1}^n P(C_i \text{ and } C_j \text{ and } C_k) + \dots \\
 &\quad + (-1)^{n+1} P(C_1 \text{ and } C_2 \text{ and } \dots \text{ and } C_n) \quad (5)
 \end{aligned}$$

The steps to find the system failure probability can be listed as follows:

1. Find the top event logic expression.
2. Apply the inclusion–exclusion formula to obtain the probability expression.
3. Substitute in probabilities of failure for each component.

To illustrate, consider a system that has two minimal cut sets {A,B} and {B,C,E}.

1. The top event logic expression, using the minimal cut set expression is given by $TOP = A.B + B.C.E$.

2. The first term in the inclusion–exclusion formula is the summation of the probability of each single minimal cut set failure: $P(A.B) + P(B.C.E)$

The second term refers to the summation of the probability of two minimal cut set failures and is given by $P(C_1.C_2) = P(A.B.B.C.E)$, which can be simplified using the absorption law to $P(A.B.C.E)$.

Therefore the top event probability expression is given by

$$P(TOP) = P(A.B) + P(B.C.E) - P(A.B.C.E) \quad (6)$$

3. If it is assumed that the probability of failure of each component is 0.1, then the probability of the top event failure is as follows:

$$P(TOP) = 0.01 + 0.001 - 0.0001 = 0.0109 \quad (7)$$

As the number of minimal cut sets for a system failure increases, this has the effect of dramatically increasing the number of terms in the inclusion–exclusion formula. This can cause a problem even for fast modern digital computers; hence approximation methods have been derived.

The two simplest approximations are based on the principles of the inclusion–exclusion method; these are the lower and upper bounds for the top event probability. The lower bound can be calculated using the first two terms in the inclusion–exclusion principle, equation (8). The upper bound can be calculated by truncating the inclusion–exclusion formula after the first term, equation (9).

$$Q_{\text{LOWER}} = \sum_{i=1}^n P(C_i) - \sum_{i=2}^n \sum_{j=1}^{i-1} P(C_i \cap C_j) \quad (8)$$

$$Q_{\text{UPPER}} = \sum_{i=1}^n P(C_i) \quad (9)$$

For other approximations see references [2, 3].

Data

In the illustrative example, the probability of failure of each component was fixed at 0.1. The failure of a component, its unavailability or failure probability, defined as q_x , is dependent on how the component is maintained.

It is common practice to use a constant rate of failure; therefore, the exponential distribution for

times to failure is used. If the rate of failure is not constant, then the appropriate distribution is needed (see reference [2]). For a nonrepairable component the information commonly available is its rate of failure, referred to as λ . The unavailability for a nonrepairable component, x , with failure rate, λ , is given by equation (10). The “ t ” given in the exponential term of the equation refers to the length of operation. Any units of measurement can be used, i.e., hours or years, as long as all parameters use the same measurement.

For components that are repairable, then the extra piece of information known is the repair rate (ν). For these components the maintenance adopted is usually unscheduled, meaning that when the component fails it will be immediately repaired. To calculate the unavailability for this type of component equation (11) is used.

$$q_x = 1 - e^{-\lambda t} \quad (10)$$

$$q_x = \frac{\lambda}{\lambda + \nu} (1 - e^{-(\lambda + \nu)t}) \quad (11)$$

The other type of maintenance is scheduled maintenance, whereupon the component will be checked for failure at specified intervals, usually typical of standby or safety components. For this type of maintenance, an inspection interval (θ) is also available. Often an average unavailability is calculated for the component, using the simplest formula given in equation (12). If a failure occurs between two inspections, the precise time of failure is not known, hence the formula, equation (12), assumes that on average it occurs half way through the interval, hence the inclusion of $\theta/2$. More exact equations for this type of maintenance are available [2, 3].

$$q_x = \lambda \left(\frac{\theta}{2} + \tau \right) \quad (12)$$

Therefore, the probability of failure of each component in the system is evaluated using one of equations (10)–(12). This value is then used in calculating the system failure probability. Typical data available for components can be found in reference [5].

Advances in the Technique

Besides the standard qualitative and quantitative analysis that can be performed, the technique has developed to incorporate importance measures [2, 6], initiator and enabler theory [2], FTA on computers

with graphical user interfaces, automatic fault tree construction and binary decision diagrams [7]. Importance measures developed by researchers Birnbaum, Esary, Proschan, Fussel, and Vesely give a ranking for components in terms of their contribution to causing a system failure. Initiator and enabler theory (developed by Lambert and Dunglinson) gives a limited capability for the technique to deal with the order in which components fail (*see also Component Reliability Importance*). With the technique being so well developed, a number of commercial computer packages are widely available. The latest development in this area has been in terms of the analysis process and a more efficient mechanism, referred to as the binary decision diagram approach, has been developed to improve on the traditional kinetic tree theory.

Related Techniques

FTA provides one tool in an armory of tools that can be used to assess the risk and reliability of a system. Often the first step to understand the system and its components is to use the **FMEA** technique (failure mode and effects analysis, *see Failure Modes and Effects Analysis*), where critical system failure modes identified are further considered using more detailed approaches depending on the system characteristics. The fault tree approach is primarily suitable for independent systems. If the system under analysis involves a sequence of events, the technique of event tree analysis [2, 3] may be more suitable. If the system contains dependencies, the method of Markov analysis can be used to find the reliability of the system [2, 3]. For systems with even more complex characteristics, **Monte Carlo simulation** may be the way forward (*see Path Sets and Cut Sets in System Reliability Modeling*). It is key that the technique used to analyze the system is chosen depending on the characteristics of the systems.

References

- [1] Vesely, W.E. (1970). A time dependent methodology for fault tree evaluation, *Nuclear Engineering and Design* **13**, 337–360.
- [2] Andrews, J.D. & Moss, T.R. (2002). *Reliability and Risk Assessment*, 2nd Edition, Professional Engineering Publishing, London.

-
- [3] Rausand, M. & Hoyland, A. (2003). *System Reliability Theory: Models and Statistical Methods*, 2nd Edition, John Wiley & Sons, Chichester.
 - [4] Montgomery, D.C. (2003). *Applied Statistics and Probability for Engineers*, 3rd Edition, John Wiley & Sons, New York.
 - [5] Moss, T.R. (2005). *The Reliability Data Handbook*, Professional Engineering Publishing, London.
 - [6] Lambert, H.E. (1975). Measures of importance of events and cut sets in fault trees, *Reliability and Fault Tree Analysis: Theoretical and Applied Aspects of System Reliability and Safety Assessment*, SIAM, Philadelphia, pp. 77–100.
 - [7] Rauzy, A. (1996). A brief introduction to binary decision diagrams, *European Journal of Automation* **30**(8), 1033–1051.

Related Articles

Component Reliability Importance; Failure Modes and Effects Analysis; Path Sets and Cut Sets in System Reliability Modeling; Path Sets and Cut Sets in System Reliability Modeling.

LISA M. BARTLETT

System Reliability: Monte Carlo Estimation

Introduction

We discuss three situations where **Monte Carlo** simulation can be used to assess system reliability or availability.

First, consider an n -component coherent system (see **Coherent Systems**) having components whose reliabilities can be expressed in the vector $\mathbf{p} = (p_1, p_2, \dots, p_n)$. The system reliability is then a function of the component reliabilities, so we write

$$r(\mathbf{p}) = P(\text{system functions if components have reliabilities } \mathbf{p}) \quad (1)$$

For series, parallel, or **series-parallel systems**, it is straightforward to compute system reliability as a function of component reliabilities. (See **Parallel, Series, and Series-Parallel Systems**.) Many other system configurations, especially when the number of components is not large, lead to closed-form expressions for system reliability. However, if the number of components is large, or the system configuration is complex, closed-form expressions may be difficult to get and **Monte Carlo simulation** may be preferable.

Second, if component reliabilities are not known, but rather estimated on the basis of a number of pass/fail tests, then system reliability must be estimated from **component reliability**. Lower confidence limits for system reliability are difficult to obtain analytically. The **bootstrap** can be used to obtain lower confidence bounds [1]. Alternatively, if a Bayesian approach is taken, we can use Monte Carlo simulation to draw samples from the posterior distribution of the system reliability.

Third, there are many measures of system performance that can be estimated only through simulation. We consider, for example, the case of a one-component system with one stand-by unit and one repair server. This situation can be modeled with **queues** (see **Queues in Reliability**). The availability $A(t)$ is defined to be the probability that the system is up (operating) at time t . Availability can be evaluated analytically only in the simplest cases; otherwise, a Monte Carlo simulation is required.

Reliability for Complex System

When a system's configuration diagram is complex, it may be difficult to compute system reliability even if all component reliabilities are known. Consider, for example, the "bridge" system in Figure 1. (This system is not that complicated and, in fact, the system reliability could be computed directly, but it will due for illustration.)

For this system, the minimal cut sets, that is, minimal sets of components whose failure guarantees the failure of the system, are $\{1, 2\}$, $\{4, 5, 6, 7\}$, $\{1, 3, 6, 7\}$, and $\{2, 3, 4, 5\}$. If we let $X_i = 1$ if component i is working, and 0 otherwise, then the structure function for this system is

$$\begin{aligned} \phi(x_1, x_2, \dots, x_7) = & [1 - (1 - x_1)(1 - x_2)] \\ & \times [1 - (1 - x_4)(1 - x_5)(1 - x_6)(1 - x_7)] \\ & \times [1 - (1 - x_1)(1 - x_3)(1 - x_6)(1 - x_7)] \\ & \times [1 - (1 - x_2)(1 - x_3)(1 - x_4)(1 - x_5)] \end{aligned} \quad (2)$$

If we let $r_i = P(X_i = 1)$, $i = 1, 2, \dots, 7$ denote the component reliabilities, we can use Monte Carlo simulation to estimate the system reliability by simulating the performance of each component, obtaining simulated values for $x_1, x_2, \dots, x_7$. For illustration, suppose that the component reliabilities are

$$r_1 = r_2 = 0.9 \quad (3)$$

$$r_3 = 0.99 \quad (4)$$

$$r_4 = r_5 = r_6 = r_7 = 0.95 \quad (5)$$

We ran 100 000 simulations, and 98 986 times the system was successful. Our estimate for the system reliability is therefore

$$\hat{r} = \frac{98\,986}{100\,000} = 0.98986 \quad (6)$$

In this example, most simulations resulted in all components working, in which case the system is sure to work. If we could focus attention on a smaller set of possible values, for example those outcomes with at least one failure, then we might eliminate much of the work and therefore be able to run more simulations in the same amount of time. This is the concept of *restricted-sampling Monte Carlo*, developed by Kumamoto *et al.* [2], and described in Henley and Kumamoto [3]. If there are n components, let S

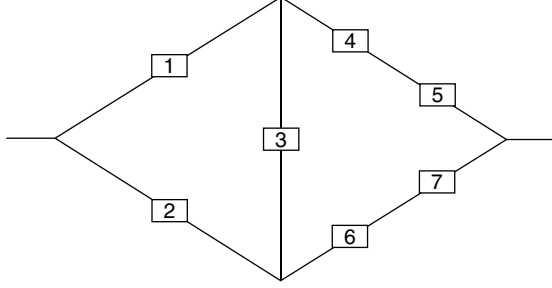


Figure 1 Structure diagram of a bridge system

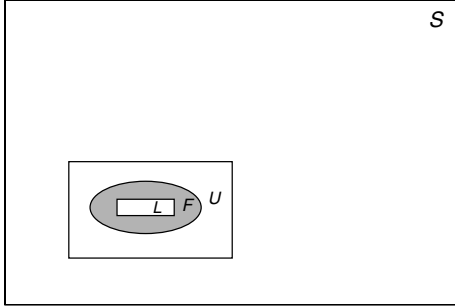


Figure 2 Restricted-sampling Monte Carlo

denote the set of all 2^n possible outcomes for the system. The key idea is to find sets U and L , which are subsets of all possible outcomes, in such a way that these sets, respectively, contain and are contained in the set F . The set U consisting of all outcomes where at least one component has failed and the set L consisting of the outcome where all components fail will always work, although more efficient choices may be made in particular cases. We must be able to compute P_U and P_L , the system reliabilities for outcomes U and L . Figure 2 illustrates the idea.

If we could then restrict our simulated values to be contained in the set $U \setminus L$, we could avoid the vast number of simulations where all components worked and hence the system worked. If we found that M outcomes out of the N simulations fell in the set $F - L$, then we would have

$$\frac{M}{N} = \text{estimate of } \frac{P_F - P_L}{P_U - P_L} \quad (7)$$

Since we know P_U and P_L , we can obtain

$$\text{estimate of } P_F = P_L + \frac{M}{N} (P_U - P_L) \quad (8)$$

Other methods of reducing computation time are discussed in [2] and [3].

Estimating System Reliability from Component Data

Suppose now that the component reliabilities are not known, and are estimated from testing. We assume that a pass/fail testing is done and that for component i , the number of passing components is y_i out of n_i that were tested. Since, given the component reliability r_i for component i we know that y_i has a binomial distribution with parameters n_i and r_i . The maximum likelihood estimate (MLE) of the system reliability is thus

$$\hat{r}_i = \frac{y_i}{n_i} \quad (9)$$

We could easily obtain (assuming that the system is not too complex) a point estimate for system reliability by substituting the estimates of the component reliabilities into the system reliability function. However, it is a remarkably difficult problem to find a lower confidence bound on the system reliability, even for simple systems like a three-component series system. Martz and Duran [4] and Leemis [1] have investigated this problem. Leemis has proposed using the bootstrap to find a lower confidence bound for system reliability. The bootstrap involves repeatedly resampling (with replacement) from the sample. He suggests some modifications for the case where a component passes all tests ($y_i = n_i$).

A fully Bayesian approach is suggested in Martz and Waller [5]. The likelihood for the data from component i is

$$L(r_i | y_i) = c r_i^{y_i} (1 - r_i)^{n_i - y_i} \quad (10)$$

For the binomial distribution, the natural conjugate prior distribution is the beta (denoted $\text{BETA}(\alpha_i, \beta_i)$), which has probability density function (PDF)

$$p(r_i) = c r_i^{\alpha_i - 1} (1 - r_i)^{\beta_i - 1}, \quad 0 < r_i < 1 \quad (11)$$

The posterior is then

$$p(r_i | y_i) = c r_i^{y_i + \alpha_i - 1} (1 - r_i)^{n_i - y_i + \beta_i - 1}, \quad 0 < r_i < 1 \quad (12)$$

which is the BETA ($y_i + \alpha_i, n_i - y_i + \beta_i$) distribution. If for example, we had a three-component series system, the system reliability would be

$$\text{System reliability} = r_1 r_2 r_3 \quad (13)$$

The posterior distribution of the quantity $r = r_1 r_2 r_3$ is rather messy, and it may be preferable to take Monte Carlo samples from it.

Suppose, for the sake of illustration, that the components for the three-component series system are believed, *a priori*, to be rather reliable with, say, a reliability of around 0.9. Thus we choose for a prior a distribution that has mean 0.9, but still has reasonable spread. The mean of the BETA(α, β) is $\alpha / (\alpha + \beta)$. Thus any combination of α and β satisfying $\alpha / (\alpha + \beta) = 0.9$ would work, but the larger α and β are, the smaller the variance. The choice of $\alpha = 18$, $\beta = 2$ gives the prior shown in Figure 3.

Suppose also that the components have undergone testing, with the results as shown in Table 1.

We ran 100 000 simulations from the posterior distribution of the system reliability, which for the series system is $r = r_1 r_2 r_3$. A histogram, which should approximate the posterior distribution, is shown in Figure 4. The lower **percentiles** for the posterior distribution are $p_{0.01} = 0.6836$, $p_{0.05} = 0.7362$, and

$p_{0.10} = 0.7620$. These numbers would serve as lower Bayesian confidence limits for the system reliability.

For more complicated systems, the method described here would be similar, except the system reliability function would be more complicated. For example, if the system were a 3-out-of-4 system, then the system reliability would be

$$r = 1 - \left[\prod_{i=1}^4 (1 - r_i) + \sum_{i=1}^4 r_i \prod_{\substack{j=1 \\ j \neq i}}^4 (1 - r_j) \right] \quad (14)$$

The procedure for estimating system reliability would be similar: collect data on each component, compute the posterior distribution for each component, simulate from each posterior, and compute the expression in equation (14).

For complicated models, especially hierarchical models (i.e., models where there are **prior distributions** on the parameters of the prior distribution), the posterior distribution cannot be found directly. In such cases, Markov Chain Monte Carlo (MCMC) methods can be used to draw samples from the posterior distribution. See Gilks *et al.* [6].

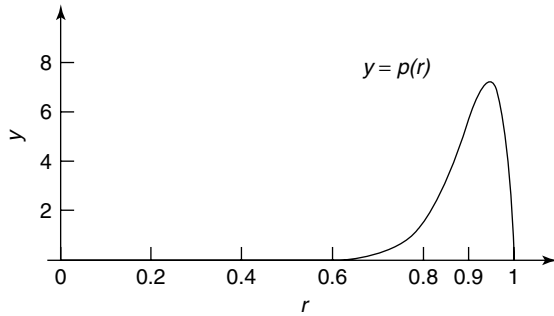


Figure 3 Prior distribution for component reliability

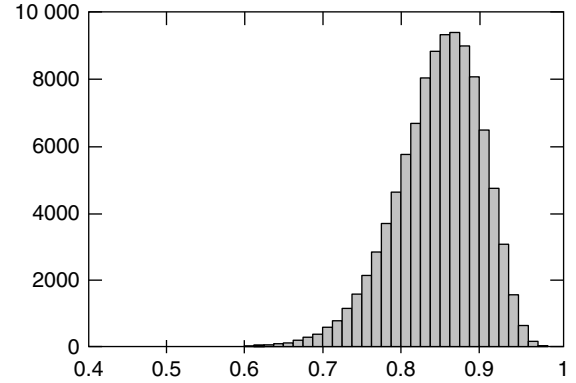


Figure 4 Monte Carlo estimate of posterior distribution

Table 1 Prior distributions, test data, and posterior distribution for example

i	Prior parameters		Test data		Posterior distribution	Posterior mean
	α_i	β_i	n_i	y_i		
1	18	2	100	97	BETA(115, 5)	0.95833
2	18	2	25	24	BETA(43, 3)	0.93478
3	18	2	10	10	BETA(28, 2)	0.93333

System Availability

Measures of system effectiveness, such as availability, are often complicated to evaluate unless the simplest assumptions are made. Even though many measures of system effectiveness, such as availability, can be evaluated in the context of queuing models (*see Queues in Reliability*), even these queuing models become too complicated to be evaluated analytically.

Consider for instance a one-component system with one component available as a standby. Switching over is instantaneous, but there is only one repair person, so if both the components are in a failed state, then the system is down. We make the following assumptions about the failure and repair times.

1. Times to failure for system 1 are independent and identically distributed (i.i.d.) and have cumulative distribution function (cdf) $F_1(x)$.
2. Times to failure for system 2 are i.i.d. and have cdf $F_2(x)$.
3. Repair times for system 1 are i.i.d. and have cdf $G_1(x)$.
4. Repair times for system 2 are i.i.d. and have cdf $G_2(x)$.
5. All failure times and repair times are independent.

The assumption 5 about the independence of failure and repair times can be weakened, but for now we will consider this case. Under these assumptions, each system can be in one of the four states:

W – working
 S – operable, but in standby mode
 R – failed and being repaired
 Q – failed and queued for repair.

Although it might seem like there would be $4 \times 4 = 16$ possible states for the system (since the overall system is comprised of two systems), there are actually only six possible states. They are

State 1: WS
 State 2: RW
 State 3: SW
 State 4: RQ
 State 5: WR
 State 6: QR

The possible state transitions are shown in Figure 5.

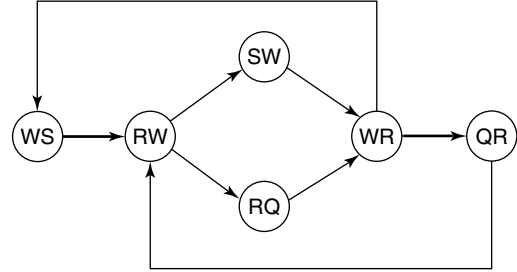


Figure 5 State diagram for one component plus standby system

We will let $X_1, X_2, \dots$ and $Y_1, Y_2, \dots$ denote the sequences of failure times for systems 1 and 2, respectively, and we will let $R_1, R_2, \dots$ and $S_1, S_2, \dots$ denote the corresponding repair times for systems 1 and 2, respectively. If the system begins in state 1 (WS), the failure of system 1 causes an immediate switch to system 2 and the commencement of repair for system 1. This is a transition to state 2. The next state depends on which happens first: the repair of system 1 or the failure of system 2. In the former case, the overall system moves to state 3 (SW), and in the latter case, the overall system moves to state 4 (RQ). In terms of the random variables defined above, the overall system transitions from state 2 to 3 if $R_1 < Y_1$ and from state 2 to 4 if $Y_1 > R_1$. Either way, the next state must be state 5 (WR). The rest of the transitions are similar.

Monte Carlo simulation can then be done by randomly generating the random variables: $X_1, X_2, \dots, Y_1, Y_2, \dots, R_1, R_2, \dots$ and $S_1, S_2, \dots$. We must then keep track of the sequence of states and the actual event times (where an event occurs whenever there is a transition from one state to another). Note that the overall system is down in states 4 and 6, and the system is up otherwise. We could estimate the long run unavailability by counting the “down” time and dividing it by the total time. For illustration, we did this under several different scenarios, as described in Table 2.

The simulations shown in Table 2 were done in Matlab. There are several packages that specialize in simulation which can be used to create similar models. For example, in Arena, the “Preempt” module can be used to switch machines after a failure, while the “Tally” can report the total downtime. Other

Table 2 Estimation of average downtime or unavailability for various failure repair distributions

Distribution of X_i	Distribution of Y_i	Distribution of R_i	Distribution of S_i	Average downtime
Exponential(1)	Exponential(1)	Exponential(0.5)	Exponential(0.5)	14.18%
Exponential(1)	Exponential(1)	Exponential(0.4)	Exponential(0.4)	10.13%
Exponential(1)	Exponential(1)	Exponential(0.3)	Exponential(0.3)	6.39%
Exponential(1)	Exponential(1)	Exponential(0.1)	Exponential(0.5)	8.01%
Gamma(1,4)	Gamma(1,4)	Gamma(0.5,3)	Gamma(0.5,3)	1.86%
Gamma(1,4)	Gamma(1,4)	Gamma(0.3,3)	Gamma(0.3,3)	0.31%
Gamma(3,4)	Gamma(3,4)	Gamma(1,3)	Gamma(1,3)	0.43%
Gamma(6,4)	Gamma(3,4)	Gamma(1,3)	Gamma(1,3)	0.15%
Weibull(6,2)	Weibull(6,2)	Gamma(1,3)	Gamma(1,3)	7.00%
Weibull(6,2)	Weibull(7,2)	Gamma(0.5,3)	Gamma(0.5,3)	1.13%

packages, such as ProModel and Flexsim, need similar modifications to create the model.

Convergence of Monte Carlo Methods

Typically a large number, say N , simulations are run and the quantity of interest is estimated from the result of these N simulations. In most cases, including all of those discussed here, the estimated value converges to the true value as $N \rightarrow \infty$. Convergence is often slow, however, with the **standard error** being proportional to $1/\sqrt{N}$.

References

- [1] Leemis, L. (2006). Lower system reliability bounds from binary failure data using bootstrapping, *Journal of Quality Technology* **38**, 1–12.
- [2] Kumamoto, H., Tanaka, K. & Inoue, K. (1977). Efficient evaluation of system reliability by Monte Carlo method, *IEEE Transactions on Reliability* **R-26**, 311–315.

- [3] Henley, E.J. & Kumamoto, H. (1981). *Reliability Engineering and Risk Assessment*, Prentice Hall, Englewood Cliffs.
- [4] Martz, H.F. & Duran, B.S. (1985). A comparison of three methods for calculating lower confidence limits on system reliability using binomial component data, *IEEE Transactions on Reliability* **R-34**(5), 311–315.
- [5] Martz, H. & Waller, R. (1982). *Bayesian Reliability Analysis*, John Wiley & Sons, New York.
- [6] Gilks, W.R., Richardson, S. & Spiegelhalter, D.J. (1996). *Markov Chain Monte Carlo in Practice*, Chapman & Hall, London.

Related Articles

Monte Carlo Methods; Monte Carlo Methods, Univariate and Multivariate.

CHRISTOPHER J. RIGDON AND
STEVEN E. RIGDON

Repairable Systems: Statistical Inference

Introduction

A repairable system is a system which, upon experiencing failure, is brought back to working condition by any corrective measure other than replacement of the whole system. Examples of **repairable systems** abound in the literature. Some of the commonly cited datasets include the aircraft air conditioning data first presented in Proschan [1], the bus motor failure data analyzed by Davis [2], the **software system failures** examined in Musa [3], and the failures of load-haul-dump (LHD) machines used in underground mines in Sweden and presented in Blischke and Murthy [4, Table 2.7].

In modeling repairable systems, the common assumption of independent and identically distributed (i.i.d.) observations will typically not hold. The failure distribution of a repaired system is, in general, different from that of a new system since the failure distribution of a repaired system depends on the age at failure as well as the type of repair performed. The two types of repair most commonly found in the literature are perfect repair and minimal repair. A perfect repair restores a system to the good-as-new state. Assuming that a perfect repair is performed at each failure results in times between failures being i.i.d. Hence the renewal process is an appropriate model to employ. In particular, if the common distribution of the times between failures is taken to be the exponential distribution, then one gets the homogeneous **Poisson process** (HPP). On the other hand, a minimal repair brings a system back to its working state just prior to failure. If it is assumed that a minimal repair is undertaken at each failure, then the nonhomogeneous Poisson process (NHPP) is the model of choice.

In addition to the renewal process and the NHPP, various models for repairable systems have been proposed and continue to be proposed in the literature. Dorado *et al.* [5] proposed a general class of **imperfect repair** models that includes the ones introduced by Brown and Proschan [6], Block *et al.* [7], and Kijima [8]. A class of models using the concept of virtual ages was proposed by Last and Szekli [9]. More recently, Lindqvist *et al.* [10] examined

heterogeneous trend renewal processes, while Peña and Hollander [11] discussed dynamic models for recurrent events that are applicable in both the repairable systems scenario as well as in medical settings. Both papers presented very general models that include most of the existing models as special cases.

In this article, the main interest is in statistical inference procedures for repairable systems. In particular, the issues of estimation and goodness-of-fit testing will be addressed. Since it is virtually impossible to present inference procedures for all models, we will focus on the HPP, NHPP, and a generalized age-dependent minimal repair model. The age-dependent minimal repair model, a model that subsumes both the HPP and NHPP, was originally proposed by Block *et al.* [7].

This paper is organized as follows. In the section titled “Inference Procedures for Homogeneous and Nonhomogeneous Poisson Processes”, we present estimation procedures for the HPP and the NHPP. In the section titled “**Age-Dependent Minimal Repair Models**”, we describe the model of interest and formulate it in terms of the counting process framework. We discuss the estimation of unknown parameters in the section titled “Point Estimation” and address goodness-of-fit testing in the section titled “Goodness-of-Fit Testing”. Finally, we present some possible extensions as well as ongoing research in the section titled “Concluding Remarks”.

Inference Procedures for Homogeneous and Nonhomogeneous Poisson Processes

In this section, we tackle the estimation problem for the two basic models for repairable systems. The method of maximum likelihood will be used to find estimators. We will consider both the failure-truncated and the time-truncated scenarios. The observed data are said to be failure-truncated if the testing of the repairable system ceases after a prespecified number of failures. On the other hand, the data are said to be time-truncated if the testing stops at a prespecified time. We will focus on the case where we observe several failures in a single repairable system.

Homogeneous Poisson Process

Let $T_1, T_2, \dots$ denote the time of the i th failure

2 Repairable Systems: Statistical Inference

and let $X_1, X_2, \dots$ denote the times between failures. The failure times can then be represented as $T_1 = X_1, T_2 = X_1 + X_2, \dots, T_k = X_1 + X_2 + \dots, X_k, \dots$. A perfect repair being undertaken at each failure results in i.i.d. interfailure times. Suppose we assume that the times between failures follow an $EXP(\theta)$ distribution, that is, the probability density function of X_i is given by $f(x) = (1/\theta) \exp(-x/\theta), x > 0, \theta > 0$. We then have an HPP with **intensity function** $\lambda(t) = 1/\theta, t > 0$. Note that the HPP models the situation where the system neither deteriorates nor improves over time.

Let us now derive the maximum-likelihood estimate for θ . If the system is tested until n failures are obtained, then the observed failure times are $0 < t_1 < t_2 < \dots < t_n$. The likelihood and log-likelihood functions are given by

$$L(\theta|t) = \left(\frac{1}{\theta}\right)^n e^{-t_n/\theta} \quad \text{and} \\ \ell(\theta|t) = -n \ln \theta - \frac{t_n}{\theta} \quad (1)$$

respectively. Taking the derivative of the log-likelihood function with respect to θ and setting the derivative equal to zero results in $\hat{\theta} = t_n/n$. Consequently, by the invariance property of maximum-likelihood estimators, the maximum-likelihood estimate of $\lambda(t_n)$ is $\hat{\lambda}(t_n) = 1/\hat{\theta} = n/t_n$.

If we want to come up with a $(1 - \alpha)100\%$ **confidence interval** (CI) for θ , then we need to examine the distribution of T_n . Since $T_n = X_1 + X_2 + \dots + X_n$, where X_i 's are i.i.d. $EXP(\theta)$, then it follows that T_n has a gamma distribution with scale parameter θ and shape parameter n . Thus, the random variable $2T_n/\theta$ has a χ^2 distribution with $2n$ **degrees of freedom**. If we let $\chi_p^2(u)$ denote the 100 p th percentile of a χ^2 distribution with u degrees of freedom, we obtain

$$\mathbf{P} \left[\chi_{\alpha/2}^2(2n) < \frac{2T_n}{\theta} < \chi_{1-\alpha/2}^2(2n) \right] = 1 - \alpha \quad (2)$$

Hence, a $(1 - \alpha)100\%$ CI for θ is $(2t_n/\chi_{1-\alpha/2}^2(2n), 2t_n/\chi_{\alpha/2}^2(2n))$.

Example 1 For the purpose of illustration, let us take the failure times for Aircraft 3 in the aircraft air conditioning data set first examined in [1]. The

original data were given in terms of the times between failures. Suppose we observe the aircraft until $n = 20$ failures have occurred. The observed failure times are

90 100 160 346 407 456 470 494 550 570
649 733 777 836 865 983 1008 1164 1474 1550

The estimated mean time between failures is then $\hat{\theta} = 1550/20 = 77.5$. In addition, a 90% CI for θ is given by

$$\left(\frac{2t_n}{\chi_{0.95}^2(40)}, \frac{2t_n}{\chi_{0.05}^2(40)} \right) = \left(\frac{2(1550)}{55.7585}, \frac{2(1550)}{26.5093} \right) \\ = (55.5969, 116.9401) \quad (3)$$

We can therefore conclude, with 90% confidence, that the mean time between failures is somewhere between 55.5969 and 116.9401.

Let us now consider the situation where the system is tested until a prespecified time t resulting in observed data $t_1 < t_2 < \dots < t_N < t$. Note that the number of observed failures, denoted by N , is a random variable which follows a Poisson distribution with mean t/θ . Given that $N = n$, the likelihood and log-likelihood functions are given by

$$L(\theta|n, t) = \frac{(t/\theta)^n \exp^{-t/\theta}}{n!} \quad \text{and} \\ \ell(\theta|n, t) = n \ln t - n \ln \theta - t/\theta - \ln n! \quad (4)$$

respectively. Taking the derivative of the log-likelihood function with respect to θ and equating it to zero yields the maximum-likelihood estimator $\hat{\theta} = t/N$. To come up with a $(1 - \alpha)100\%$ CI for θ , we use the result of Pearson and Hartley [12] involving a $(1 - \alpha)100\%$ CI for t/θ . They proved that if $N = n$ failures are observed, then a $(1 - \alpha)100\%$ CI for t/θ is $(1/2\chi_{\alpha/2}^2(2n), 1/2\chi_{1-\alpha/2}^2(2n + 2))$. Consequently, a $(1 - \alpha)100\%$ CI for θ is given by

$$\left(\frac{2t}{\chi_{1-\alpha/2}^2(2c + 2)}, \frac{2t}{\chi_{\alpha/2}^2(2c)} \right)$$

Example 2 Consider the same data set examined in Example 1 but assume that we observed the aircraft until $t = 1575$. Note that with this particular choice

of t , we end up with exactly the same set of failure times as in Example 1. Our point estimate of θ , however, changes to $\hat{\theta} = 1575/20 = 78.75$. It then follows that a 90% CI for θ is

$$\left(\frac{2t}{\chi_{0.95}^2(42)}, \frac{2t}{\chi_{0.05}^2(40)} \right) = \left(\frac{2(1575)}{58.1240}, \frac{2(1575)}{26.5093} \right) \\ = (54.1945, 118.8262) \quad (5)$$

Thus we conclude, with 90% confidence, that the mean time between failures lie between 54.1945 and 118.8262. It is worthwhile to note that the estimates change depending on how the observed data were generated.

Nonhomogeneous Poisson Process

Let us now examine the scenario where a minimal repair is performed at each failure thereby resulting in the nonhomogeneous Poisson process as the appropriate model. In contrast to the HPP, the NHPP is capable of modeling systems which wear out or improve over time. In particular, if we let the intensity function of the NHPP be of the form

$$\lambda(t; \gamma, \theta) = \frac{\gamma}{\theta} \left(\frac{t}{\theta} \right)^{\gamma-1}, \quad \theta > 0 \quad (6)$$

then we obtain the power law process, a widely used model in the repairable systems setting owing to its flexibility. For instance, a deteriorating system can be modeled by taking $\gamma > 1$, while an improving system can be modeled by assuming a value of $\gamma < 1$. Note that taking $\gamma = 1$ allows us to recover an HPP with intensity $1/\theta$.

Suppose the repairable system is observed until n failures occur. Then the observed failure times are $t_1, t_2, \dots, t_n$, where n is fixed. The joint **probability density function** of the failures times $T_1, T_2, \dots, T_n$ for a failure-truncated NHPP with intensity function $\lambda(t)$ is

$$f(t_1, t_2, \dots, t_n) = \left[\prod_{i=1}^n \lambda(t_i) \right] \exp \left[- \int_0^{t_n} \lambda(s) ds \right], \\ 0 < t_1 < t_2 < \dots < t_n \quad (7)$$

Consequently, the likelihood function for the power law process can be expressed as

$$L(\gamma, \theta | \mathbf{t}) = \frac{\gamma^n}{\theta^{n\gamma}} \left[\prod_{i=1}^n t_i \right]^{\gamma-1} \exp \left[- \left(\frac{t_n}{\theta} \right)^\gamma \right], \\ 0 < t_1 < t_2 < \dots < t_n \quad (8)$$

with corresponding log-likelihood function given by

$$\ell(\gamma, \theta | \mathbf{t}) = n \ln \gamma - n\gamma \ln \theta \\ + (\gamma - 1) \sum_{i=1}^n \ln t_i - \left(\frac{t_n}{\theta} \right)^\gamma \quad (9)$$

Taking the partial derivatives of the log-likelihood function with respect to γ and θ and equating them to zero results in

$$\frac{n}{\hat{\gamma}} - n \ln \hat{\theta} + \sum_{i=1}^n \ln t_i - \left(\frac{t_n}{\hat{\theta}} \right)^{\hat{\gamma}} \ln \left(\frac{t_n}{\hat{\theta}} \right) = 0 \quad \text{and} \\ \frac{\hat{\gamma} t_n^{\hat{\gamma}}}{\hat{\theta}^{\hat{\gamma}+1}} - \frac{n \hat{\alpha}}{\hat{\theta}} = 0 \quad (10)$$

respectively. Solving the system of equations simultaneously yields the maximum-likelihood estimates

$$\hat{\gamma} = \frac{n}{\sum_{i=1}^n \ln \left(\frac{t_n}{t_i} \right)} \quad \text{and} \quad \hat{\theta} = \frac{t_n}{n^{1/\hat{\gamma}}} \quad (11)$$

To develop an interval estimate for γ , we use the fact that $2n\gamma/\hat{\gamma}$ follows a χ^2 distribution with $2(n-1)$ degrees of freedom (see [13, Theorem 31]). Consequently, we can write the probability statement

$$\mathbf{P} \left[\chi_{\alpha/2}^2(2(n-1)) \leq \frac{2n\gamma}{\hat{\gamma}} \leq \chi_{1-\alpha/2}^2(2(n-1)) \right] \\ = 1 - \alpha \quad (12)$$

A few algebraic manipulations yield a $(1-\alpha)100\%$ CI for γ given by

$$\left(\frac{\hat{\gamma}}{2n} \chi_{\alpha/2}^2(2(n-1)), \frac{\hat{\gamma}}{2n} \chi_{1-\alpha/2}^2(2(n-1)) \right) \quad (13)$$

4 Repairable Systems: Statistical Inference

The parameter θ has no simple interpretation and hence, CIs for θ are usually not of interest.

Example 3 Let us consider an excerpt from the software system failures dataset originally presented in [3]. The data consist of failure times, measured in execution time in CPU seconds, of a real-time command and control software system. Suppose we observed the system until $n = 20$ failures were realized resulting in the following failure times

3 33 146 227 342 351 353 444 556 571
709 759 836 860 968 1056 1726 1846 1872 1986

The point estimates of γ and θ are therefore

$$\begin{aligned}\hat{\gamma} &= \frac{20}{29.9260} = 0.668315 \quad \text{and} \\ \hat{\theta} &= \frac{1986}{20^{1/0.668315}} = 22.451613\end{aligned}\quad (14)$$

It should not be surprising that the estimated value of γ is less than one. An examination of the data reveals that the times between failures seem to be increasing which is an indication of a system that is improving over time. A 95% CI for γ is then given by

$$\begin{aligned}&\left(\frac{\hat{\gamma}}{2n}\chi_{0.025}^2(38), \frac{\hat{\gamma}}{2n}\chi_{0.975}^2(38)\right) \\ &= \left(\frac{0.668315(22.8785)}{40}, \frac{0.668315(56.8955)}{40}\right) \\ &= (0.382251, 0.950603)\end{aligned}\quad (15)$$

Hence, using a method that captures the true value of γ 95% of the time, we estimate that γ lies between 0.382251 and 0.950603.

Let us now examine the time-truncated case so that our observed values are $0 < t_1 < t_2 < \dots < t_N < t$. Keeping in mind that the number of observed failures N is a random variable, the likelihood and log-likelihood functions, given $N = n$ and $\mathbf{t}$, are

$$\begin{aligned}L(\gamma, \theta | n, \mathbf{t}) &= \frac{\gamma^n}{\theta^{n\gamma}} \left(\prod_{i=1}^n t_i\right)^{\gamma-1} \exp\left[-\left(\frac{t}{\theta}\right)^\gamma\right], \\ n &\geq 1, 0 < t_1 < \dots < t_n < t\end{aligned}\quad (16)$$

and

$$\begin{aligned}\ell(\gamma, \theta | n, \mathbf{t}) &= n \ln \gamma - n\gamma \ln \theta \\ &+ (\gamma - 1) \sum_{i=1}^n \ln t_i - \left(\frac{t}{\theta}\right)^\gamma\end{aligned}\quad (17)$$

respectively. Taking the partial derivatives of the log-likelihood function with respect to γ and θ , equating them to zero, and solving the resulting system of equations yields the maximum-likelihood estimators

$$\hat{\gamma} = \frac{N}{\sum_{i=1}^N \ln\left(\frac{t}{T_i}\right)} \quad \text{and} \quad \hat{\theta} = \frac{t}{N^{1/\hat{\gamma}}}\quad (18)$$

An interval estimate for γ can be obtained by using the result that, conditioned on $N = n$, the quantity $2n\gamma/\hat{\gamma}$ follows a χ^2 distribution with $2n$ degrees of freedom (see [13, Theorem 32]). Thus, given $N = n$, a $(1 - \alpha)100\%$ CI for γ is

$$\left(\frac{\hat{\gamma}}{2n}\chi_{\alpha/2}^2(2n), \frac{\hat{\gamma}}{2n}\chi_{1-\alpha/2}^2(2n)\right)\quad (19)$$

Note the difference in degrees of freedom between the failure-truncated and time-truncated cases. Just as in the failure-truncated situation, there is not much interest in interval estimation for the parameter θ .

Example 4 Let us revisit the data set used in Example 3. However, we now assume that we observed the system until $t = 2000$ CPU seconds. This particular value of t was chosen to illustrate the difference in estimates between the failure-truncated and time-truncated cases. The observed failure times are exactly the same as those presented in Example 3. The resulting point estimates of γ and θ turn out to be

$$\begin{aligned}\hat{\gamma} &= \frac{20}{30.0655} = 0.665192 \quad \text{and} \\ \hat{\theta} &= \frac{2000}{20^{1/0.665192}} = 22.139031\end{aligned}\quad (20)$$

Consequently, a 95% CI for γ is

$$\left(\frac{\hat{\gamma}}{2n}\chi_{0.025}^2(40), \frac{\hat{\gamma}}{2n}\chi_{0.975}^2(40)\right)$$

$$\begin{aligned}
 &= \left(\frac{0.665192(24.4330)}{40}, \frac{0.665192(59.3417)}{40} \right) \\
 &= (0.406316, 0.986841) \quad (21)
 \end{aligned}$$

Hence, using a method that captures the true value of γ 95% of the time, we estimate that γ lies between 0.406316 and 0.986841. Note that even though we used the same set of failure times for Examples 3 and 4, we obtained different estimates depending on the type of truncation that was used.

Age-Dependent Minimal Repair Models

In practice, the type of repair that is undertaken at each failure need not be always perfect or always minimal. In this section, we formalize the age-dependent minimal repair model, hereon referred to as the *BBS model*. Suppose that at time $t = 0$, a component or system with time-to-first-failure density function f is put on test. Upon failure at time t , one of two types of repair will be undertaken: a perfect repair with probability $p(t)$ or a minimal repair with probability $1 - p(t)$. Recall that a perfect repair restores the effective age of a system to 0, while a minimal repair returns the system's effective age to what it was just before failure. The BBS model admits as special cases some of the more common models in repairable system literature. For instance, if the probability of perfect repair does not depend on time, then the imperfect repair model proposed by Brown and Proschan [6] is obtained. On the other hand, taking $p(t) = 0$ recovers the NHPP model, while assuming $p(t) = 1$ results in the i.i.d. model or the renewal model.

In most instances, the conditions under which the systems operate affect the distribution of future lifetimes. The effects of operating or environmental conditions are typically taken into account through the use of covariates. To allow for even more generality in the BBS model, we incorporate covariates via a proportional hazards model. Hence, we model the system's hazard rate function as

$$\lambda(t; \mathbf{X}) = \lambda_0(t) \psi(\mathbf{X}(t), \boldsymbol{\beta}) \quad (22)$$

where $\lambda_0(t)$ is called the *baseline hazard rate function* and $\psi(\mathbf{X}(t), \boldsymbol{\beta})$ is a function that takes into

account the effects of various factors that impact system failure through the covariates $\mathbf{X}(t)$ and regression parameters $\boldsymbol{\beta}$, which characterize the effect of the covariates. Note that we are allowing the possibility that the covariates may be time dependent. Some of the proposed forms for $\psi(\mathbf{X}(t), \boldsymbol{\beta})$ include exponential, logistic, linear, and inverse linear. In this article, we will use the **Cox proportional hazards** model [14], which postulates that the system's hazard rate function is of the form

$$\lambda(t; \mathbf{X}) = \lambda_0(t) \exp[\boldsymbol{\beta}' \mathbf{X}(t)] \quad (23)$$

We will use counting processes as the main tool in formalizing the model and developing the statistical inference procedures. The formulation of the BBS model in terms of the counting process framework is attributed to Hollander *et al.* [15].

All random entities are defined on the filtered probability space $(\Omega, \mathbf{F} = (\mathcal{F}_t : t \geq 0), \mathcal{P})$, where the filtration is typically the natural filtration. In addition, the upper endpoint of the observation period is denoted by $\tau > 0$, which may be fixed or random. Denote by $W_0 \equiv 0 < W_1 < W_2 < \dots$ the successive failure times of a system, and let $U_1, U_2, \dots$ be a sequence of i.i.d. Uniform[0,1] random variables that are independent of the failure times. Since a perfect repair returns the system to the good-as-new state, it suffices to observe a system only until the time of its first perfect repair. Hence, the sequence $(W_1, W_2, \dots, W_\nu)$, where $\nu = \inf\{k \in \{1, 2, \dots\} : U_k < p(W_k)\}$, is an epoch of the BBS model. Consider observing n independent BBS epochs which may be generated by observing n independent systems, each until the time of its first perfect repair, or observing a single system up to the time of the n th perfect repair. In the latter case, the operating time reverts to zero after each perfect repair. The sample data can then be represented by $\{W_{jk} : 1 \leq j \leq n, 1 \leq k \leq \nu_j\}$. Define the stochastic processes $N = \{N_1(t), N_2(t), \dots, N_n(t)\}$ and $Y = \{Y_1(t), Y_2(t), \dots, Y_n(t)\}$, where

$$\begin{aligned}
 N_j(t) &= \sum_{k=1}^{\infty} I\{W_{jk} \leq t \wedge W_{j\nu_j}\} \quad \text{and} \\
 Y_j(t) &= I\{W_{j\nu_j} \geq t\} \quad (24)
 \end{aligned}$$

The process $N_j(t)$ basically counts the number of minimal repairs observed until time t for the j th

system, while $Y_j(t)$ is an indicator of whether the j th system is still “at risk” at time t .

With respect to the filtration $\mathbf{F}$, the cumulative intensity function is the compensator function associated with the counting process N . Hence, if the j th system has a possibly time-dependent covariate process $X_j(\cdot)$, then the compensator of N is $A = \{(A_1(t; \xi, \beta), \dots, A_n(t; \xi, \beta))\}$ with

$$A_j(t; \xi, \beta) = \int_0^t Y_j(s) \lambda(s; \xi) \exp\{\beta' X_j(s)\} ds \quad (25)$$

where $\lambda(s; \xi)$ is some baseline hazard function, β is a $q \times 1$ vector of regression coefficients, and $X_1(s), \dots, X_n(s)$ are $q \times 1$ vectors of locally bounded predictable covariate processes. In this article, we will assume a parametric specification for the baseline hazard rate function. In order to guarantee that the waiting time to the first perfect repair is almost surely finite, we assume that

$$\int_0^\infty p(t) \lambda(t; \xi) \exp\{\beta' X_j(t)\} dt = \infty \quad (26)$$

Block *et al.* [7] showed that if condition (26) holds, then the waiting time to the first perfect repair has hazard rate function $\lambda^*(t; \xi) = p(t) \lambda(t; \xi) \exp\{\beta' X_j(t)\}$.

Point Estimation

Let us first address the issue of estimating the unknown parameters. For the estimation problem, we will use the method of maximum likelihood. Under the generalized BBS setup, the likelihood process is given by

$$L(t; \xi, \beta) = \left[\prod_{s \leq t} \prod_{j=1}^n \left[Y_j(s) \lambda(s; \xi) \times \exp\{\beta' X_j(s)\} \right]^{dN_j(s)} \right] \times \exp \left[- \sum_{j=1}^n A_j(t) \right] \quad (27)$$

where Π denotes the product integral. The product integral Π can be viewed as a continuous version of the ordinary product $\prod$, in much the same way that the integral $\int$ is a continuous version of the sum $\sum$. We refer the reader to Gill and Johansen [16] or Andersen *et al.* [17] for further details about the product integral. Consequently, the log-likelihood process is

$$\begin{aligned} \ell(t; \xi, \beta) &= \int_0^t \sum_{j=1}^n \log [Y_j(s) \lambda(s; \xi) \exp\{\beta' X_j(s)\}] \\ &\quad \times dN_j(s) - \sum_{j=1}^n A_j(t) \\ &= \sum_{j=1}^n \left[\int_0^t \log [Y_j(s) \lambda(s; \xi) \exp\{\beta' X_j(s)\}] \right. \\ &\quad \times dN_j(s) - \int_0^t Y_j(s) \lambda(s; \xi) \\ &\quad \times \exp\{\beta' X_j(s)\} ds \left. \right] \end{aligned} \quad (28)$$

The maximum-likelihood estimators of the parameters ξ and β can be obtained by direct maximization of equations (27) or (28). Depending on the form of the integrand in the likelihood or log-likelihood equations, it may be possible to get a closed-form expression for the estimators by applying the principles of calculus which involves taking the partial derivatives with respect to each parameter and then setting them to zero. However, in the event that the integrand is fairly involved, then iterative methods such as the Newton–Raphson method, or an **expectation-maximization (EM) algorithm** (*cf.*, Dempster *et al.* [18]) are some of the more common approaches. An example of an EM approach can be found in Stocker and Peña [19].

Example 5 Let us revisit the power law process that was discussed in the section titled “Nonhomogeneous Poisson Process”. Note that this model can be recovered from the BBS model by taking the no-covariate case with probability of perfect repair $p(t) = 0$ and assuming that the intensity function is of the form $\lambda(t; \gamma, \theta) = (\gamma/\theta) (t/\theta)^{\gamma-1}$. In contrast to Example 3 where we observed a single repairable system, suppose we observe n such independent

repairable systems. Let t_{ji} denote the i th failure on the j th system. Take the failure-truncated case where the j th system is observed until the k_j th failure so that the observed failure times for the j th system are $0 < t_{j1} < t_{j2} < \dots < t_{jk_j}$.

Noting that the process N_j jumps only at the failure times and the size of each jump is one, then evaluating the log-likelihood function given in equation (28) at the end of the observation period yields

$$\begin{aligned} \ell(\gamma, \alpha) &= \sum_{j=1}^n \left[\sum_{i=1}^{k_j} \log \left[\frac{\alpha}{\gamma^\alpha} t_{ji}^{\alpha-1} \right] \right. \\ &\quad \left. - \int_0^{t_{jk_j}} \frac{\alpha}{\gamma^\alpha} s^{\alpha-1} ds \right] \\ &= \log \alpha \sum_{j=1}^n k_j - \alpha \log \gamma \sum_{j=1}^n k_j + (\alpha - 1) \\ &\quad \times \sum_{j=1}^n \sum_{i=1}^{k_j} \log t_{ji} - \sum_{j=1}^n \left(\frac{t_{jk_j}}{\alpha} \right)^\beta \end{aligned} \quad (29)$$

Taking the partial derivatives with respect to α and γ and equating them to zero results in the well-known estimators

$$\begin{aligned} \hat{\gamma} &= \left(\frac{\sum_{j=1}^n t_{jk_j}}{\sum_{j=1}^n k_j} \right)^{1/\hat{\alpha}} \quad \text{and} \\ \hat{\alpha} &= \frac{\sum_{j=1}^n k_j}{\sum_{j=1}^n \left(\frac{t_{jk_j}}{\hat{\gamma}} \right)^{\hat{\alpha}} \log t_{jk_j} - \sum_{j=1}^n \sum_{i=1}^{k_j} \log t_{ji}} \end{aligned} \quad (30)$$

An examination of the above expressions reveals that an iterative algorithm needs to be employed in order to obtain the estimates. More detailed discussion concerning estimation issues, including Bayesian estimation, for the power law process is given in Rigdon and Basu [13].

Goodness-of-Fit Testing

In some instances, the problem on hand is not concerned with estimating unknown parameters, but rather with determining whether a claim about a parameter and/or distribution is valid in light of the observed data. The interest therefore shifts to **hypothesis testing**. In this section, we will discuss a particular form of hypothesis testing known as *goodness-of-fit testing*. In particular, we will focus on validating the assumption concerning the baseline hazard function in the generalized BBS model. This problem is often encountered in the engineering and operations research settings, where one is typically interested in predicting the reliability of systems, developing **maintenance policies**, and providing spare parts.

There is an abundance of goodness-of-fit procedures that can be utilized. In this article, however, we will explore the family of hazard-based smooth goodness-of-fit tests introduced by Peña [20, 21] in the i.i.d setting and extended by Agustin and Peña [22, 23] and Agustin [24] to the repairable systems scenario. The main idea in developing smooth goodness-of-fit tests is the embedding of the hypothesized hazard rate function in a larger parametric family, and subsequently forming score tests. Members of the larger parametric family are generated from smooth transformations of the members of the hypothesized class. The idea of embedding was first proposed by Neyman [25], who used density functions instead of hazard rate functions. An excellent reference for the density-based smooth goodness-of-fit tests is the book by Rayner and Best [26].

Let us first consider the simple hypothesis setting that involves testing $H_0 : \lambda(\cdot) = \lambda_0(\cdot)$ versus $H_1 : \lambda(\cdot) \neq \lambda_0(\cdot)$, where $\lambda_0(\cdot)$ is a fully specified function. Consider the embedding class for the hazard rate functions given by

$$\mathcal{A}_k = \{\lambda_k(\cdot; \theta) = \lambda_0(\cdot) \exp[\theta' \Psi(\cdot)] : \theta \in \mathbb{R}^k\} \quad (31)$$

where k is some prespecified positive integer called the *smoothing order* and $\Psi(\cdot) = (\Psi_1(\cdot), \dots, \Psi_k(\cdot))'$ is some function or process. By choosing appropriate forms for $\Psi(\cdot)$, a wide range of test statistics can be obtained. Some possibilities for $\Psi(\cdot)$ are polynomial, trigonometric, or wavelet functions. By setting $\theta = (\theta_1, \dots, \theta_k)' = \mathbf{0}$ in equation (31), the hypothesized class $\mathcal{C}$ is recovered. Thus, the hypotheses of interest

8 Repairable Systems: Statistical Inference

can be restated as

$$H_0^* : \theta = \mathbf{0}, \beta \in \mathfrak{N}^m \text{ versus } H_1^* : \theta \neq \mathbf{0}, \beta \in \mathfrak{N}^m \quad (32)$$

The family of tests is then derived by forming the score test for $H_0^* : \theta = \mathbf{0}$ and determining its distributional properties. Agustin and Peña [22] established that under H_0^* and assuming certain regularity conditions hold, the estimated score process

$$\frac{1}{\sqrt{n}} \hat{U}(\tau; \mathbf{0}, \hat{\beta}) = \frac{1}{\sqrt{n}} \sum_{j=1}^n \int_0^\tau \Psi(s) [dN_j(s) - Y_j(s) \times \lambda_0(s) \exp[\hat{\beta}' X_j(s)] ds] \quad (33)$$

converges in distribution to a normal random variable with mean $\mathbf{0}$ and covariance matrix $\Sigma(\tau)$. The expression for $\Sigma(\tau)$ is quite involved so we refer the reader to [23]. The estimator $\hat{\beta}$ is the partial-likelihood maximum-likelihood estimator (PLMLE). Consequently, an asymptotic α -level goodness-of-fit test for $H_0 : \lambda(\cdot) = \lambda_0(\cdot)$ rejects H_0 if

$$S(\tau; \hat{\beta}) \equiv \frac{1}{n} \hat{U}(\tau; \mathbf{0})' \hat{\Sigma}^-(\tau) \hat{U}(\tau; \mathbf{0}) \geq \chi_{k^*, \alpha}^2 \quad (34)$$

where $\hat{\Sigma}^-(\cdot)$ is a generalized inverse of the covariance matrix estimate $\hat{\Sigma}(\cdot)$ and $\chi_{k^*, \alpha}^2$ is the $(1 - \alpha)$ 100th percentile of the χ^2 distribution with degrees of freedom $k^* = \text{rank}[\hat{\Sigma}(\tau)]$ (cf., Li and Doss [27, Theorem 2]).

As mentioned earlier, varying the form of the smoothing process Ψ generates a rich class of tests. In order to come up with closed-form expressions, let us consider the no-covariate setting, that is, we take $X = \mathbf{0}$. Under this framework, three possible consistent estimators of $\Sigma(\tau)$ are

$$\hat{\Sigma}(\tau) = \frac{1}{n} \sum_{j=1}^n \int_0^\tau \Psi(t) \Psi(t)' Y_j(t) \lambda_0(t) dt \quad (35)$$

$$\tilde{\Sigma}(\tau) = \frac{1}{n} \sum_{j=1}^n \int_0^\tau \Psi(t) \Psi(t)' dN_j(t) \quad (36)$$

and

$$\begin{aligned} \hat{\Sigma}(\tau) &= \frac{1}{2n} \sum_{j=1}^n \int_0^\tau \Psi(t) \Psi(t)' \\ &\quad \times [dN_j(t) + Y_j(t) \lambda_0(t) dt] \\ &= \frac{\hat{\Sigma}(\tau) + \tilde{\Sigma}(\tau)}{2} \end{aligned} \quad (37)$$

Monte Carlo simulation studies have shown that using the estimator $\hat{\Sigma}(\tau)$ results in tests with achieved levels that are more consistent with the prespecified levels.

If we let $k = 1$ and $\Psi(t) = 1$, then we obtain the test statistic

$$Q(\tau) \equiv \frac{1}{\sqrt{n}} U(\tau) = \frac{1}{\sqrt{n}} \sum_{j=1}^n [N_j(\tau) - R_{jv_j}^\tau] \quad (38)$$

where $R_{jv_j}^\tau = \Lambda_0(\tau \wedge W_{jv_j})$ and $\Lambda_0(t) = \int_0^t \lambda_0(s) ds$ is the cumulative hazard function corresponding to λ_0 . Using $\hat{\Sigma}(\tau)$ as the estimator of the variance function results in a test with rejection region

$$S(\tau) = \frac{\left[\sum_{j=1}^n (N_j(\tau) - R_{jv_j}^\tau) \right]^2}{\sum_{j=1}^n R_{jv_j}^\tau} \geq \chi_{1; \alpha}^2 \quad (39)$$

This test can be viewed as a generalization of the Pearson-type 1-degree-of-freedom test proposed by Akritas [28].

On the other hand, choosing $\Psi(t) = \exp[-\Lambda_0(t)]$ and using $\hat{\Sigma}(\tau)$ as the variance estimator yields the test statistic

$$S(\tau) \equiv \frac{2 \left\{ \sum_{j=1}^n \left[\sum_{i=1}^{N_j(\tau)} \exp\{-R_{ji}\} - (1 - \exp\{-R_{jv_j}^\tau\}) \right] \right\}^2}{\sum_{j=1}^n (1 - \exp\{-2R_{jv_j}^\tau\})} \quad (40)$$

where $R_{ji} = \Lambda_0(X_{ji})$ are the generalized **residuals**. Hence, H_0 is rejected when $S(\tau) \geq \chi_{1; \alpha}^2$. This test can be regarded as a generalization of Moses' test (cf.,

Hollander and Proschan [29], Gatsonis *et al.* [30], Horowitz and Neumann [31]).

Another possibility is to let the process $\Psi(t)$ be a stochastic process. In particular, by taking $\Psi(t) = B(t)/B(\tau) - 1/2$, where $B(s) = \sum_{i=1}^n \int_0^s Y_i(u) \lambda_0(u) du$, we obtain the test statistic

$$Q(\tau) = \frac{1}{\sqrt{n}} \times \left[\frac{\sum_{j=1}^{N_*(\tau)} \sum_{l=1}^j \left[\sum_{i=1}^n I\{R_{i v_i}^\tau \geq R_{(l)}\} \right]}{\sum_{i=1}^n R_{i v_i}^\tau} - \frac{N_*(\tau)}{2} \right] \quad (41)$$

and the estimated variance function

$$\hat{\Sigma}(\tau) = \frac{1}{12n} \sum_{i=1}^n R_{i v_i}^\tau \quad (42)$$

where $N_*(\tau) = \sum_{i=1}^n N_i(\tau)$ and $R_{(1)}, \dots, R_{(N_*(\tau))}$ are the ordered generalized residuals of the aggregated processes. The test, which rejects H_0 if $S(\tau) \equiv Q^2(\tau)/\hat{\Sigma}(\tau) \geq \chi_{1,\alpha}^2$, results in a generalization of the spacings test of Barlow *et al.* [32].

One can even take the combination of the three previous smoothing processes to get

$$\Psi(t) = \left[1, \exp\{-\Lambda_0(t)\}, \frac{B(t)}{B(\tau)} - \frac{1}{2} \right]' \quad (43)$$

Simulation results revealed that the combined test is sensitive to a wide range of alternatives and thus shows promise as an omnibus test. We refer the reader to Agustin [24] for the details in this setting. In addition, the polynomial form $\Psi(t) = [1, \Lambda_0(t), \dots, \Lambda_0(t)^{k-1}]'$ was examined in Agustin and Peña [22].

Example 6 Consider the well-known dataset presented in Proschan [1, Table 1] which contains the times between failures of the air conditioning system of 13 Boeing jet planes. Following Presnell *et al.* [33], we will take a major overhaul, indicated by

**, as the time of the first perfect repair. In cases where a major overhaul was not performed, we will take the last observed failure time as the time of the first perfect repair. Since all failure times in between will be regarded as minimal repair times, we assume that τ is large enough so we are able to observe all failures until the time of the first perfect repair. Presnell *et al.* [33] tested the minimal repair assumption and found no evidence against it. Suppose we want to test the hypothesis that $H_0 : \lambda(t) = 1/94.34$, that is, the distribution of the times to first failure of the systems is exponential with mean 94.34 hours. The test statistics and corresponding p values for the generalized Moses' test and the generalized spacings test are 6.21 ($p = 0.0127$) and 5.64 ($p = 0.0176$), respectively. Hence, both tests rejected the exponential assumption at the 5% level of significance.

Let us now turn our attention to the composite hypothesis scenario, where interest is on testing whether the baseline hazard rate function belongs to some parametric family. Hence, the hypotheses of interest are $H_0 : \lambda(\cdot) \in \mathcal{C} \equiv \{\lambda_0(\cdot; \xi) : \xi \in \Xi\}$ versus $H_1 : \lambda_0(\cdot) \notin \mathcal{C}$, with the functional form of $\lambda_0(\cdot)$ assumed to be known, except for the value of ξ .

The embedding class for the hazard rate functions in the composite hypothesis case is given by

$$\mathcal{A}_k = \{\lambda_k(\cdot; \theta, \xi, \beta) = \lambda_0(\cdot; \xi) \exp[\theta' \Psi(\cdot; \xi)] : \theta \in \mathfrak{R}^k\} \quad (44)$$

Since setting $\theta = \mathbf{0}$ recovers the hypothesized class $\mathcal{C}$, we can express the original hypotheses as

$$\begin{aligned} H_0^* : (\theta, \xi, \beta) &\in \{\mathbf{0}\} \times \Xi \times \mathcal{B} \quad \text{versus} \\ H_1^* : (\theta, \xi, \beta) &\in \mathfrak{R}^k \setminus \{\mathbf{0}\} \times \Xi \times \mathcal{B} \end{aligned} \quad (45)$$

The relevant score process for θ is therefore

$$\begin{aligned} U(t; \theta, \xi, \beta) &= \sum_{j=1}^n \int_0^t \Psi(s; \xi) \\ &\quad \times [dN_j(s) - Y_j(s) \lambda_0(s; \xi) \\ &\quad \times \exp[\beta' X_j(s)] \exp[\theta' \Psi(s; \xi)] ds] \end{aligned} \quad (46)$$

The challenge in this setting is the need to contend with two nuisance parameters. For the Cox proportional hazards model, it is common to estimate the

vector of regression coefficients by the PLMLE (cf., Keiding *et al.* [34]). The PLMLE of β , denoted by $\hat{\beta}$, is obtained by solving the equation

$$\sum_{j=1}^n \int_0^\tau [X_j(s) - E(s; \beta)] dN_j(s) = \mathbf{0} \quad (47)$$

where

$$S_{(m)}(t; \beta) = \frac{1}{n} \sum_{j=1}^n X_j(t)^{\otimes m} Y_j(t) \exp[\beta' X_j(t)],$$

$$m = 0, 1, 2 \quad (48)$$

$$E(t; \beta) = \frac{S_{(1)}(t; \beta)}{S_{(0)}(t; \beta)} \quad (49)$$

and for a vector $\mathbf{a}$, $\mathbf{a}^{\otimes 0} = 1$, $\mathbf{a}^{\otimes 1} = \mathbf{a}$, and $\mathbf{a}^{\otimes 2} = \mathbf{a}\mathbf{a}'$. The estimator of ξ , denoted by $\hat{\xi} \equiv \hat{\xi}(\hat{\beta})$, is the solution of

$$\begin{aligned} & \sum_{j=1}^n \int_0^\tau \left(\frac{\partial}{\partial \xi} \log \lambda_0(s; \xi) \right) \\ & \times \left[dN_j(s) - Y_j(s) \lambda_0(s; \xi) \exp[\hat{\beta}' X_j(s)] ds \right] \\ & = \mathbf{0} \end{aligned} \quad (50)$$

Assuming that certain regularity conditions hold, Agustin and Peña [23] showed that an asymptotic α -level goodness-of-fit test for $H_0 : \lambda(\cdot) \in \mathcal{C}$ rejects H_0 whenever

$$\frac{1}{n} \hat{U}(\hat{\xi}, \hat{\beta})' \hat{\Gamma}^-(\hat{\xi}, \hat{\beta}) \hat{U}(\hat{\xi}, \hat{\beta}) \geq \chi_{k^*, \alpha}^2 \quad (51)$$

where $\hat{U}(\hat{\xi}, \hat{\beta})$ is the estimated score process, $\hat{\Gamma}^-(\hat{\xi}, \hat{\beta})$ is a generalized inverse of the estimated limiting covariance matrix of $\hat{U}/\sqrt{n}$ under H_0 , and $\chi_{k^*, \alpha}^2$ is the $(1 - \alpha)100$ th percentile of the central χ^2 distribution with $k^* = \text{rank}[\hat{\Gamma}]$ degrees of freedom.

A key element of this family of tests is the process $\Psi(\cdot)$ which characterizes the family of alternatives that a specific test will be able to detect powerfully. We now explore the use of the specification $\Psi(t) = [\Lambda_0(t; \xi), \Lambda_0^2(t; \xi), \dots, \Lambda_0^k(t; \xi)]'$, where $k \in \{1, 2, \dots\}$ is a fixed smoothing order. With this particular choice of Ψ , we obtain tests with relatively simple forms, which turn out to be functions of the model's

generalized residuals. In the simpler case, where the covariates are time independent, the resulting estimated score statistic is

$$Q(\tau) = \frac{1}{\sqrt{n}} \sum_{j=1}^n \left[\sum_{i=1}^{N_j(\tau)} R_{ji}^\ell - \exp\{\hat{\beta}' X_j\} \frac{(R_{jv_j}^\tau)^{\ell+1}}{\ell+1} \right]_{\ell=1, \dots, k} \quad (52)$$

with $R_{ji} = \Lambda_0(W_{ji}; \hat{\xi})$, $(i = 1, \dots, v_j, j = 1, \dots, n)$ and $R_{jv_j}^\tau = \Lambda_0(\tau \wedge W_{jv_j}; \hat{\xi})$. For smoothing order k , an asymptotic α -level test has rejection region

$$S_k \equiv Q'(\tau) \hat{\Gamma}^-(\tau) Q(\tau) \geq \chi_{k^*, \alpha}^2 \quad (53)$$

where k^* is the rank of $\hat{\Gamma}(\tau)$. The components of the estimated covariance matrix are presented in Agustin and Peña [23]. Simulation results show that these tests are powerful omnibus tests.

Example 7 Consider the LHD machines dataset presented in Blischke and Murthy [4, Table 2.7]. LHD machines were used for moving ore and rock in underground mines in Sweden and a vital subsystem of this machine is the hydraulic system. Over a 2-year period, data consisting of time between successive failures of the hydraulic systems of six LHD machines were recorded. For the purpose of this example, we consider all intermediate repairs as minimal repair and the final repair as the first perfect repair. Since the machines are classified as old, medium age, and relatively new, we will use relative age of the machine as a covariate. Owing to the nature of the classification, we need two binary variables, X_1 and X_2 . Old machines are coded as $X_1 = 1$ and $X_2 = 0$; medium age ones as $X_1 = 0$ and $X_2 = 1$, and relatively new ones as $X_1 = 0$ and $X_2 = 0$. We are interested in testing whether an exponential distribution is a plausible model for the time to first failure of the systems. The relevant hypotheses are therefore, $H_0 : \lambda(\cdot) \in \mathcal{C} \equiv \{\lambda_0(\cdot; \xi) = \xi : \xi \in (0, \infty)\}$ versus $H_1 : \lambda(\cdot) \notin \mathcal{C}$. Incorporating the covariates into the model via the Cox proportional hazards model and using the partial likelihood to derive the estimates of β , we obtain $\hat{\beta} = (0.0541, -0.1028)$. Using these estimates in equation (50) results in $\hat{\xi} = 0.0077$. Suppose we use the polynomial specification with smoothing order $k = 4$. Simulation results in Agustin and Peña [23] showed that $k \in \{3, 4\}$ are promising

omnibus tests. The resulting test statistic and its corresponding p value are $S_4 = 9.9291$ and $p = 0.0416$, respectively. Hence, we conclude that the exponential distribution is not an appropriate model for the time to first failure of the hydraulic systems of the LHD machines.

Concluding Remarks

This article focused on inference procedures for the homogeneous and nonhomogeneous Poisson processes as well as a generalized BBS model. In particular, the maximum-likelihood principle was used to obtain estimators of unknown parameters. Other methods of estimation such as the Bayesian approach are certainly applicable and we refer the reader to the book by Rigdon and Basu [13] for details. For a more general class of models that take frailty into account, parametric estimation procedures are proposed in Stocker and Peña [19], while nonparametric estimation issues are tackled in Peña *et al.* [35]. The issue of efficiently determining the smoothing order for the goodness-of-fit problem is still unresolved. Ledwina [36] and Kallenberg and Ledwina [37] used the Schwartz information criterion as a basis for a data-driven smooth goodness-of-fit test. Their tests, however, are formulated using density functions instead of hazard rate functions. Research is still ongoing for a data-driven version of the hazard-based smooth goodness-of-fit tests.

References

- [1] Proschan, F. (1963). Theoretical explanation of observed decreasing failure rate, *Technometrics* **5**, 375–383.
- [2] Davis, D. (1952). An analysis of some failure data, *Journal of the American Statistical Association* **47**, 113–150.
- [3] Musa, J. (1979). *Software Reliability Data; Data Analysis Center for Software*, Rome Air Development Center, New York.
- [4] Blischke, W. & Murthy, P. (2000). *Reliability: Modeling, Prediction and Optimization*, Wiley-Interscience, New York.
- [5] Dorado, C., Hollander, M. & Sethuraman, J. (1997). Nonparametric estimation for a general repair model, *Annals of Statistics* **25**, 1140–1160.
- [6] Brown, M. & Proschan, F. (1983). Imperfect repair, *Journal of Applied Probability* **20**, 851–859.
- [7] Block, H., Borges, W. & Savits, T. (1985). Age-dependent minimal repair, *Journal of Applied Probability* **22**, 370–385.
- [8] Kijima, M. (1989). Some results for repairable systems with general repair, *Journal of Applied Probability* **26**, 89–102.
- [9] Last, G. & Szekli, R. (1998). Asymptotic and monotonicity properties of some repairable systems, *Advances in Applied Probability* **30**, 1089–1110.
- [10] Lindqvist, B., Elvebakk, G. & Heggland, K. (2003). The trend-renewal process for statistical analysis of repairable systems, *Technometrics* **45**, 31–44.
- [11] Peña, E. & Hollander, M. (2004). Mathematical reliability: an expository perspective, in *Models for Recurrent Events in Reliability and Survival Analysis*, R. Soyer, T. Mazzuchi & N. Singpurwalla, eds, Kluwer Academic Publishers, Chapter 6, pp. 105–123.
- [12] Pearson, E. & Hartley, H. (1966). *Biometrika Tables for Statisticians*, 3rd Edition, Cambridge University Press, Cambridge, Vol. 1.
- [13] Rigdon, S. & Basu, A. (2000). *Statistical Methods for the Reliability of Repairable Systems*, John Wiley & Sons, New York.
- [14] Cox, D. (1972). Regression models and life-tables (with discussion), *Journal of the Royal Statistical Society. Series B* **34**, 187–220.
- [15] Hollander, M., Presnell, B. & Sethuraman, J. (1992). Nonparametric methods for imperfect repair models, *Annals of Statistics* **20**, 879–896.
- [16] Gill, R. & Johansen, S. (1990). A survey of product-integration with a view toward application in survival analysis, *Annals of Statistics* **18**, 1501–1555.
- [17] Andersen, P., Borgan, O., Gill, R. & Keiding, N. (1993). *Statistical Models Based on Counting Processes*, Springer-Verlag, New York.
- [18] Dempster, A., Laird, N. & Rubin, D. (1977). Maximum likelihood from incomplete data via the EM algorithm, *Journal of the Royal Statistical Society. Series B* **39**, 1–38.
- [19] Stocker, R. & Peña, E. (2007). A general class of parametric models for recurrent event data, *Technometrics*, **49**, 210–221.
- [20] Peña, E. (1998). Smooth goodness-of-fit tests for the baseline hazard in Cox's proportional hazards model, *Journal of the American Statistical Association* **93**, 673–692.
- [21] Peña, E. (1998). Smooth goodness-of-fit tests for composite hypothesis in hazard-based models, *Annals of Statistics* **26**, 1935–1971.
- [22] Agustin, Z. & Peña, E. (2001). Goodness-of-fit of the distribution of time-to-first-occurrence in recurrent event models, *Lifetime Data Analysis* **7**, 289–306.
- [23] Agustin, Z. & Peña, E. (2005). A basis approach to goodness-of-fit testing in recurrent event models, *Journal of Statistical Planning and Inference* **133**, 285–303.
- [24] Agustin, Z. (2002). Generalized hazard-based goodness-of-fit tests for a repairable system, *Journal of Statistical Computation and Simulation* **72**, 519–531.
- [25] Neyman, J. (1937). Smooth test for goodness of fit., *Skandinavisk Aktuarietidskrift* **20**, 149–199.

- [26] Rayner, J. & Best, D. (1989). *Smooth Tests of Goodness of Fit*, Oxford, New York.
- [27] Li, G. & Doss, H. (1993). Generalized Pearson-Fisher chi-square goodness-of-fit tests, with applications to models with life history data, *Annals of Statistics* **21**, 772–779.
- [28] Akritas, M. (1988). Pearson-type goodness-of-fit tests: the univariate case, *Journal of the American Statistical Association* **83**, 222–230.
- [29] Hollander, M. & Proschan, F. (1979). Testing to determine the underlying distribution using randomly censored data, *Biometrics* **35**, 393–401.
- [30] Gatsonis, C., Hsieh, H. & Korwar, R. (1985). Simple nonparametric tests for a known standard survival based on censored data, *Communications in Statistics-Theory and Methods* **14**, 2137–2162.
- [31] Horowitz, J. & Neumann, G. (1992). A generalized moments specification test of the proportional hazards model, *Journal of the American Statistical Association* **87**, 234–240.
- [32] Barlow, R., Bartholomew, D., Bremner, J. & Brunk, H. (1972). *Statistical Inference Under Order Restrictions*, John Wiley & Sons, New York.
- [33] Presnell, B., Hollander, M. & Sethuraman, J. (1994). Testing the minimal repair assumption in an imperfect repair model, *Journal of the American Statistical Association* **89**, 289–297.
- [34] Keiding, N., Andersen, C. & Fledelius, P. (1998). The Cox regression model for claims data in non-life insurance, *Astin Bulletin* **28**, 95–118.
- [35] Peña, E., Slate, E. & Gonzalez, J. (2007). Semiparametric inference for a general class of models for recurrent events, *Journal of Statistical Planning and Inference*, **137**, 1727–1747.
- [36] Ledwina, T. (1994). Data-driven version of Neyman's smooth test-of-fit, *Journal of the American Statistical Association* **89**, 1000–1005.
- [37] Kallenberg, W. & Ledwina, T. (1995). Consistency and Monte Carlo simulation of a data driven version of smooth goodness-of-fit tests, *Annals of Statistics* **23**, 1594–1608.

Related Articles

Age-Dependent Minimal Repair and Maintenance; Analysis of Recurrent Events from Repairable Systems; Intensity Functions for Nonhomogeneous Poisson Processes; Poisson Processes; Repairable Systems Reliability; Repairable Systems: Bayesian Analysis.

MA. ZENIA N. AGUSTIN

System Reliability: Computational Algebra Methods

Introduction

Monte Carlo and discrete-event simulation use the “experimental” approach to solving complicated modeling problems in reliability and survival analysis. This is typically a fallback approach taken when analytic solutions to a problem do not exist due to mathematical intractability. This chapter outlines the use of computer algebra systems for solving reliability-modeling problems. Computer algebra systems provide exact solutions to these problems, as opposed to the approximate solutions provided by simulation.

There are many choices of base language that could be used in this context. We use the Maple-based APPL (A Probability Programming Language) developed by Glen *et al.* [1] due to our familiarity with the language. Another reasonable choice would have been the MathStatca program developed by Rose and Smith [2].

We begin with a brief introduction to the data structures used in APPL. We then proceed to a series of examples of reliability-modeling problems that can be solved using the language.

Data Structures

One of the initial issues associated with the development of a computer language for solving probability problems concerns how to represent a random variable. APPL uses a “list-of-lists” data structure to represent a random variable, where the first sublist gives the functional form that represents the distribution of the random variable (typically given by a **probability density function**, survivor function, or hazard function), the second sublist gives the support of the random variable, and the third sublist indicates whether the random variable is discrete or continuous and the form of the distribution representation in the first sublist.

As a simple introductory example, consider a continuous random variable X that has a “triangular”

distribution with a minimum of 0, a mode of 1, and a maximum of 3, with probability density function

$$f(x) = \begin{cases} 2x/3 & 0 \leq x < 1 \\ 1 - x/3 & 1 \leq x < 3 \end{cases} \quad (1)$$

This random variable is instantiated in the list-of-lists format with the Maple command

```
> X := [[x -> 2 * x/3, x -> 1 - x/3],  
[0, 1, 3], ["Continuous", "PDF"]];
```

The three integers in the second sublist denote the endpoints of the piecewise support of X . Likewise, a discrete random variable X with probability density function

$$f(x) = \begin{cases} 0.6 & x = 1 \\ 0.3 & x = 2 \\ 0.1 & x = 3 \end{cases} \quad (2)$$

is instantiated with the Maple command

```
> X := [[0.6, 0.3, 0.1], [1, 2, 3],  
["Discrete", "PDF"]];
```

The data structure described above is needed for arbitrarily distributed random variables. Common distributions, such as the binomial and normal distributions, have shortcuts for defining the random variable of interest. The names of the functions for these distributions begin with an uppercase letter and end with an appended RV, for random variable. For example, the APPL command

```
> X := TriangularRV(0, 1, 3);
```

is a shortened and equivalent statement to the earlier instantiation of the distribution of X . Similarly, the statements

```
> Y := BinomialRV(10, 2 / 3);  
> Z := NormalRV(mu, sigma);
```

define a random variable Y having the binomial distribution with parameters $n = 10$ and $p = 2/3$, and a random variable Z having a normal distribution with a symbolic mean parameter μ and symbolic standard deviation parameter σ , respectively.

Once the data structure for defining the distribution of a random variable had been established and common distributions preloaded, algorithms were developed that allowed the manipulation of these random variables. A simple example is to find the expected value of a particular random variable. In order to find the mean, variance, **skewness**, and **kurtosis** of the random variable Y defined above for example, execute the APPL statements

2 System Reliability: Computational Algebra Methods

```
> Mean(Y);
> Variance(Y);
> Skewness(Y);
> Kurtosis(Y);
```

which return the mean, variance, skewness, and kurtosis of the binomial random variable as

$$\frac{20}{3}, \frac{20}{9}, -\frac{\sqrt{5}}{10}, \frac{57}{20}$$

The next section considers a series of examples that use algorithms that have been developed in APPL to manipulate random variables.

Examples

This section contains the following examples of the use of a computer algebra system for reliability analysis: (a) analyzing a k -out-of- n system; (b) finding the mean time to system failure of a parallel system with symbolic parameters; (c) finding a conditional percentile of a system with differing supports for each of its components; (d) finding the distribution of the sum of independent and identically distributed random variables in order to determine the probability of mission success; (e) finding a lower confidence bound for the system reliability; and (f) finding the exact distribution of a goodness-of-fit statistic.

System Survivor Functions

Consider a 4-out-of-5 system that operates when four or five of the components are functioning. Let each component have an independent and identically distributed Weibull time to failure with survivor function

$$S(t) = e^{-(3t)^2}, \quad t > 0 \quad (3)$$

Find

1. the hazard function of the time to system failure;
2. the mean time to system failure;
3. the variance of the time to system failure;
4. the probability that the system survives to time 0.25.

The APPL code to calculate these quantities is

```
X := WeibullRV(3, 2);
Y := OrderStat(X, 5, 2);
HF(Y);
```

```
Mean(Y);
Variance(Y);
SF(Y, 1 / 4);
```

The first statement defines the random variable X , which is associated with the lifetimes of the components. Since this system has five components, and the system failure occurs upon the second failure, the system lifetime is the second-order statistic (see Evans *et al.* [3], for more details) of five independent Weibull random variables, denoted by Y . The HF function returns the hazard function of the system time to failure, which is

$$h(t) = \frac{360t(e^{-9t^2} - 1)}{4e^{-9t^2} - 5}, \quad t > 0 \quad (4)$$

The mean and variance of the system time to failure are

$$E[T] = \frac{5\sqrt{\pi}}{12} - \frac{2\sqrt{5}\pi}{15} \cong 0.2101 \quad \text{and} \\ V[T] = \frac{1}{20} - \frac{21\pi}{80} - \frac{\sqrt{5}\pi}{9} \cong 0.005867 \quad (5)$$

The probability that the system survives to time 0.25 is determined by the SF function, which yields $S(0.25) = 5e^{-9/4} - 4e^{-45/16} \cong 0.2868$.

Symbolic Parameters

Consider an n -component parallel system consisting of components with independent and identically distributed exponential(λ) lifetimes. Find the expected time to system failure.

Since the reliability function of a parallel system of n identical components, each with reliability p , is

$$r(p) = 1 - (1 - p)^n \quad (6)$$

and the survivor function of an exponential random variable with failure rate λ is

$$S(t) = e^{-\lambda t}, \quad t > 0 \quad (7)$$

the system survivor function is

$$S(t) = 1 - (1 - e^{-\lambda t})^n, \quad t > 0 \quad (8)$$

The mean time to system failure is

$$E[T] = \int_0^\infty (1 - (1 - e^{-\lambda t})^n) dt \quad (9)$$

This quantity can be calculated, for instance when $n = 5$, with the APPL statements

```
X := ExponentialRV(lambda);
T := MaximumIID(X, 5);
Mean(T);
```

which return a mean time to system failure of $E[T] = 137/60\lambda$. This problem can also be solved symbolically using the following Maple statements, avoiding the use of APPL entirely:

```
assume(lambda > 0);
n := 5;
s := 1 - (1 - exp(-lambda*t)) ^ n;
int(s, t = 0 .. infinity);
```

which also return a mean time to system failure of $E[T] = 137/60\lambda$. The advantage afforded by APPL is that the mathematical details associated with the derivation can be bypassed.

Lifetimes with Varying Support

Consider the three-component system with component 1 in parallel with the series arrangement of components 2 and 3, i.e., a three-component system with minimal path sets $\{1\}$ and $\{2, 3\}$. Let T_1 , T_2 , and T_3 be the independent lifetimes of the three components, where

$$T_1 \sim U(2, 10), \quad T_2 \sim U(4, 14), \quad T_3 \sim U(6, 18) \quad (10)$$

Find the 95th percentile of the distribution of the failure time of a system that has survived to time 8.

Although this particular problem is conceptually straightforward, keeping track of the various supports is tedious. The APPL code to solve this problem is given below.

```
T1 := UniformRV(2, 10);
T2 := UniformRV(4, 14);
T3 := UniformRV(6, 18);
T := Maximum(T1, Minimum(T2, T3));
TC := Truncate(T, 8, infinity);
CriticalPoint(TC, 0.95);
```

This code returns the 95th percentile of the distribution of the remaining system lifetime as

$$\frac{32 - \sqrt{31}}{2} \cong 13.2161 \quad (11)$$

One way to verify this solution is *via Monte Carlo simulation*, which is also possible in APPL and Maple, although not computationally efficient in either language. The Maple code below generates 10 000 system lifetimes and estimates the quantity of interest. Maple's rand function generates a random 12-digit integer.

```
variates := [];
for i from 1 to 10000 do
    t1 := 2 + 8 * rand()/1000000000000;
    t2 := 4 + 10 * rand()/1000000000000;
    t3 := 6 + 12 * rand()/1000000000000;
    t := max(t1, min(t2, t3));
    if (t > 8) then
        variates := [op(variates), t];
    fi;
od;
variates := sort(variates);
n := nops(variates);
print(evalf(variates[floor(0.95*n)]));
```

The output of several runs of this code segment hovers around the theoretical value, which verifies the analytic solution provided by APPL.

Cold Standby Systems

Let X_1, X_2, X_3 be independent gamma distributed random variables that represent the lifetime of a unit (X_1) and the lifetimes of two backup cold standby units (X_2 and X_3) taken on a mission of duration 50. If $X_1 \sim \text{gamma}(1/2, 3)$, $X_2 \sim \text{gamma}(1/4, 5)$, and $X_3 \sim \text{gamma}(1/6, 7)$, find the probability that the mission is successful.

The two standard ways to address a problem of this type are the central limit theorem and Monte Carlo simulation. Since there are only three random variables being added here, the central limit theorem will not be particularly accurate. Monte Carlo simulation, on the other hand, requires custom programming and the result is typically stated as an interval around the true value. Also, the number of replications required is a quadratic function of the desired accuracy: each additional digit of accuracy to construct such an interval requires a 100-fold increase in the number of Monte Carlo replications.

Alternatively, the APPL code required to solve this problem exactly is

4 System Reliability: Computational Algebra Methods

```
> X := GammaRV(1 / 2, 3);
> Y := GammaRV(1 / 4, 5);
> Z := GammaRV(1 / 6, 7);
> W := Convolution(X, Y);
> T := Convolution(W, Z);
> SF(T, 50);
```

which yields the probability of a successful mission as

$$\frac{246\,009\,667}{256}e^{-25/3} - 61\,758\,360e^{-25/2} + \frac{2199}{256}e^{-25} \cong 0.8371 \quad (12)$$

Bootstrapping

Bootstrapping is a popular statistical technique that involves **resampling** data in order to draw inferences. Bootstrapping involves drawing B “**bootstrap**” samples of n observations without replacement from the n data values collected. Resampling introduces error (because of the finite value of B) that, in some instances, can be eliminated by using APPL. The advice in the literature makes the choice of B nontrivial for a novice to bootstrapping, e.g., “ B in the range of 50 to 200 usually makes $\hat{se}_{boot}$ a good **standard error** estimator, even for estimators like the median” (Efron and Tibshirani [4, p. 14]) or “ B should be ≥ 500 or 1000 in order to make the variability of [the estimated 95th percentile] acceptably low” for estimating 95th **percentiles** (Efron and Tibshirani [4, p. 275]). Because exact values are always preferred to approximations, there is no point to adding the resampling error to the sampling error inherent in the data. In addition, a novice to bootstrapping can easily confuse the **sample size** n with number of bootstrap samples B . Eliminating the resampling of the dataset B times simplifies the conceptualization of the bootstrap process.

Consider the problem of using bootstrapping to determine a 95% lower confidence bound on the system reliability for a series system of three independent components using the failure data (y_i, n_i) , where

Table 1 Failure data for a three-component series system

Component number	$i = 1$	$i = 2$	$i = 3$
Number passing (y_i)	21	27	82
Number on test (n_i)	23	28	84

y_i is the number of components of type i that pass the test and n_i is the number of components of type i on test, $i = 1, 2, 3$ (see Leemis [5]). The number of components tested and the number of passes for each type of component for the example are given in Table 1. The point estimate for the system reliability is

$$\frac{21}{23} \cdot \frac{27}{28} \cdot \frac{82}{84} = \frac{1107}{1288} \cong 0.8595 \quad (13)$$

To calculate an approximate 95% lower bound for the system reliability, the bootstrapping procedure resamples B of the systems, calculates the system reliability, and then returns the 0.05 B th-ordered system reliability. More specifically, for the three-component system of interest, the bootstrapping algorithm follows these steps:

- For the first component, the dataset (21 ones and 2 zeros) is sampled with replacement 23 times.
- These values are summed and divided by 23 yielding a reliability estimate for the first component.
- The previous two steps are repeated for components 2 and 3.
- The product of the reliability estimates for the three components are multiplied (because the components are arranged in series and their failures are independent) to give a system reliability estimate.

The above procedure is repeated B times. The B system reliability estimates are then sorted. Finally, the 0.05 B th-ordered system reliability, a lower 95% bootstrap **confidence interval** bound on the system reliability, is returned.

As in the previous example, a finite choice for B will result in resampling error, which can be avoided by using APPL. The APPL statements given below utilize the `Product` and `Transform` procedures. This alternative approach to determining a lower 95% bootstrap confidence interval bound on the system reliability is equivalent to using an infinite value for B .

```
> X1 := BinomialRV(23, 21 / 23);
> X1 := Transform(X1, [[x -> x / 23],
[-infinity, infinity]]);
> X2 := BinomialRV(28, 27 / 28);
> X2 := Transform(X2, [[x -> x / 28],
[-infinity, infinity]]);
```

```

> X3 := BinomialRV(84, 82 / 84);
> X3 := Transform(X3, [[x -> x / 84],
[-infinity, infinity]]);
> Temp := Product(X1, X2);
> T := Product(Temp, X3);

```

The binomial distribution is appropriate since the resampling from the dataset is performed with replacement. Because the support of X_i contains $n_i + 1$ values, for $i = 1, 2, 3$, there are a possible $24 \cdot 29 \cdot 85 = 59\,160$ potential mass values for the random variable T determined by the `Product` procedure. Of these, only 6633 are distinct because the `Product` procedure combines repeated values. The lower 95% bootstrap confidence interval bound is the 0.05 fractile of the distribution of T , which is $6723/9016 \cong 0.7457$. This value has been verified by Monte Carlo simulation.

Kolmogorov–Smirnov Test Statistic

After fitting a parametric distribution to component or system life data, the model adequacy should be assessed. The Kolmogorov–Smirnov test statistic D_n is used for testing goodness of fit, and is most appropriate when data are sampled randomly from a continuous population. The Kolmogorov–Smirnov test statistic D_n is the largest vertical distance between an empirical cumulative distribution function $F_n(x)$ and a fitted or hypothesized cumulative distribution function $\hat{F}(x)$ for all values of x , i.e.,

$$D_n = \sup_x \{|F_n(x) - \hat{F}(x)|\} \quad (14)$$

The empirical cumulative distribution function $F_n(x)$ is the proportion of the observations $X_1, X_2, \dots, X_n$ that are less than or equal to x and is defined as

$$F_n(x) = \frac{I(x)}{n} \quad (15)$$

where n is the size of the random sample and $I(x)$ is the number of X_i 's less than or equal to x . In other words, $F_n(x)$ is the proportion of the sample that is less than or equal to x .

Unfortunately, the distribution of the test statistic under the null hypothesis is mathematically intractable. An algorithm has been developed to find the distribution of the test statistic in the “all-parameters-known” case (Drew *et al.* [6]). Although computationally intensive for large sample sizes, the

algorithm returns the cumulative distribution function of the Kolmogorov–Smirnov test statistic. This algorithm was implemented in the APPL procedure `KSRV` as illustrated in the example below.

Find the probability density function and the variance of the Kolmogorov–Smirnov test statistic in the all-parameters-known case for a sample of size $n = 3$. The APPL code to solve this problem is

```

> X := KSRV(3);
> Variance(X);

```

which returns the cumulative distribution function as

$$F(x) = \begin{cases} 48x^3 - 24x^2 + 4x - 2/9 & 1/6 \leq x < 1/3 \\ -12x^3 + 14x^2 - 8x/3 & 1/3 \leq x < 1/2 \\ -4x^3 + 2x^2 + 10x/3 - 1 & 1/2 \leq x < 2/3 \\ 2x^3 - 6x^2 + 6x - 1 & 2/3 \leq x < 1 \end{cases} \quad (16)$$

and the variance as exactly $38\,827/2\,099\,520$, or approximately 0.01849. The lower limit of the support of X corresponds to the case of the hypothesized cumulative distribution function bisecting the risers of the empirical cumulative distribution function. The `KSRV` function works for all values of its argument, limited by CPU and memory requirements associated with the algorithm.

Conclusion

There are a variety of problems associated with the calculation of the reliability of a system of components that can be effectively addressed with a computer algebra system. These provide an exact solution to analytically tractable systems analysis problems, providing an alternative to Monte Carlo simulation, discrete-event simulation, and approximate methods.

Acknowledgments

The author is grateful for the APPL algorithm development and programming work of John Drew, Diane Evans, and Andy Glen.

References

- [1] Glen, A., Evans, D. & Leemis, L. (2001). APPL: A probability programming language, *The American Statistician* **55**(2), 156–166.
- [2] Rose, C. & Smith, M.D. (2002). *Mathematical Statistics and Mathematics*, Springer-Verlag.

- [3] Evans, D., Leemis, L. & Drew, J. (2006). The Distribution of order statistics for discrete random variables with applications to bootstrapping, *INFORMS Journal on Computing* **18**(1), 19–30.
- [4] Efron, B. & Tibshirani, R. (1993). *An Introduction to the Bootstrap*, Chapman & Hall, New York.
- [5] Leemis, L. (2006). Lower confidence bounds for system reliability from binary failure data using bootstrapping, *Journal of Quality Technology* **38**(1), 2–13.
- [6] Drew, J. Glen, A. & Leemis, L. (2000). Computing the cumulative distribution function of the Kolmogorov–Smirnov Statistic, *Computational Statistics and Data Analysis* **34**(1), 1–15.

LAWRENCE M. LEEMIS

Genetic Algorithms in Reliability

Introduction

Reliability is an important research area attracting many researchers due to its critical importance in various systems. In fact, reliability plays a vital role in society. People everywhere depend continually on complex technological systems such as washing machines, telephones, televisions, cars, planes, and so on, and every time we use any of these systems, we expect that they perform the task they are intended for. This general idea that something should be fit for purpose is more properly defined as reliability.

For engineering purposes, reliability, $R(t)$, is defined as the probability that a system will perform its intended function during a specified period under stated conditions. Mathematically, this can be expressed as

$$R(t) = \int_t^{\infty} f(x) dx \quad (1)$$

where $f(t)$ is the **probability density function** for the failure time.

In recent years, metaheuristics have been selected and successfully applied to solve several reliability-optimization problems, as described by Gen and Cheng [1]. Genetic algorithms (GAs) are one of the most widely used metaheuristics and they have proved to be powerful tools for effectively solving complex engineering optimization problems.

Genetic Algorithms

GAs, originally developed by Holland [2], are non-deterministic stochastic search/optimization methods that simulate the process of natural evolution to solve problems with a complex solution space. GAs remain the most recognized form of evolutionary computation algorithms. In GA terminology, a solution vector $\mathbf{x} \in \mathbf{X}$ is called an *individual* or a *chromosome*. Chromosomes are made of discrete units called *genes*. Each gene controls one or more features of the chromosome. In the original implementation of GA by Holland [2], genes are assumed to be binary digits.

In later implementations, more varied gene types have been introduced. Normally, a chromosome corresponds to a unique solution $\mathbf{x}$ in the solution space. This requires a mapping mechanism between the solution space and the chromosomes. This mapping is called an *encoding*. In fact, GA works on the *encoding* of a problem, not on the problem itself [3].

In its general form, a GA has an initial population of individuals that is generated either randomly or heuristically. At every generation, the individuals in the current population are decoded and evaluated according to some predefined quality criterion, referred to as the *fitness function*. Creation of new members is done by *crossover* and *mutation* operations. The performance of the algorithm is mainly influenced by these two operators. The primary purpose of the crossover operator is to get genetic material from the previous generation to the subsequent generation, while the main purpose of the mutation operator is to introduce a certain amount of randomness into the search.

The effectiveness of the *crossover* operator dictates the rate of convergence; while the *mutation* operator prevents the algorithm from prematurely converging to a local optimum. In the selection procedure, individuals are chosen according to their fitness value. Individuals with high fitness values have better chances of reproducing, while low-fitness ones are more likely to disappear. The procedure is terminated either when the search process stagnates or when a predefined number of generations is reached.

Since GAs are advanced search mechanisms, they offer the capability of exploring large and complex problem spaces. However, it is important to not forget that GAs are stochastic iterative processes and they are not guaranteed to converge to the global optimal solution. Hence, the termination condition may be specified as some fixed maximum number of generations or as the attainment of an acceptable fitness level. Figure 1 shows the flowchart of a single-objective GA.

A very remarkable point of the GAs is that although they use the idea of randomness when performing a search, it must be clearly understood that GAs are not simply random search algorithms. Random search algorithms can be inherently inefficient due to the directionless nature of the search. The GAs are not directionless. They use knowledge from previous generations of chromosomes (i.e., solution encodings) to construct new chromosomes that will

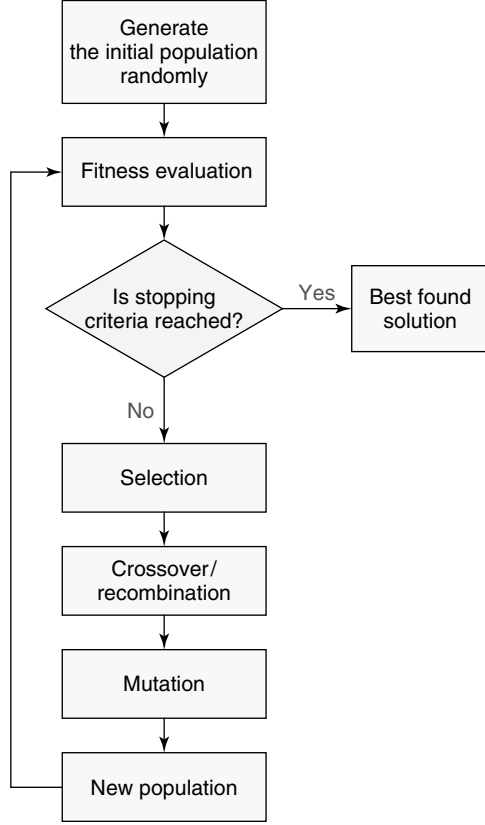


Figure 1 Single-objective genetic algorithm

approach the optimal solution. Thus, GAs are a form of a randomized search and the way that the parent solutions are chosen and combined comprise a stochastic process; see Lisnianski and Levitin [4]. A comprehensive description of GA theory and its application in engineering can be found in Goldberg [5] and Gen and Chen [1].

Redundancy Allocation Problem

There are many system design engineering projects that require the use of **redundancy** to meet high reliability specifications. To make systems more reliable, redundant units are usually allocated, such that the overall reliability is maximized. However, the use of redundancy also adds more cost, weight, volume, and so on, to the system. There also generally exist system-level constraints and the problem is then to select the design configuration that maximizes some

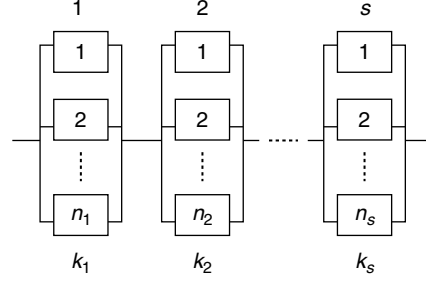


Figure 2 Series-parallel redundancy system

stated objective functions. This is known as the redundancy allocation problem (RAP). Solving the RAP has been shown to be NP hard by Chern [6] and different mathematical optimization approaches have been used to determine optimal or good solutions to this problem [7–12].

The RAP is a system design optimization problem and probably the most-studied design configuration of the RAP is the series-parallel system, such as depicted in Figure 2. This system has a total of s subsystems arranged in series. For each subsystem, there are n_i functionally equivalent components arranged in parallel. Each component has potentially different levels of cost, weight, reliability, and other characteristics. The n_i components are to be selected from m_i available component types, where multiple copies of each type can be selected. k_i represents the minimum number of components required for subsystem i to function and is specified for each subsystem.

Different problem formulations of the RAP have been presented in the literature. For instance, one well-known example is the following formulation in which the objective is to maximize the system reliability given restrictions on the system cost, C , and the system weight, W :

$$\max \prod_{i=1}^s R_i(\mathbf{x}_i | k_i) \quad (2)$$

subject to:

$$\sum_{i=1}^s C_i(\mathbf{x}_i) \leq C \quad (3)$$

$$\sum_{i=1}^s W_i(\mathbf{x}_i) \leq W \quad (4)$$

$$k_i \leq \sum_{j=1}^{m_i} x_{ij} \leq n_{\max,i} \quad \text{for } i = 1, 2, \dots, s \quad (5)$$

$$x_{ij} \in \{0, 1, 2, \dots\}$$

The problem presented in equations (2–5) is to determine which design alternative to select with the specified level of **component reliability** to achieve the maximum reliability given the overall restrictions on system cost of C and weight of W . In the formulation, x_{ij} is the number of the j th component selection used for subsystem i , and $C_i(\mathbf{x}_i)$ and $W_i(\mathbf{x}_i)$ are the cost and weight for subsystem i .

There are other alternative methods that have been used to solve the RAP. However, metaheuristics seem to be among the preferred techniques. In general, metaheuristics [13–16] can be defined as high-level strategies for exploring search spaces by using different methods. Many researchers have made use of metaheuristics to find solutions for different versions of the RAP [12, 17–20]. However, this article is devoted to the use of GAs in reliability, which is one of the most widely used metaheuristics.

Single-Objective Optimization of Redundant Systems

Solving the RAP as a single-objective optimization problem involves either the maximization of the system design subject to limits on resource consumption or the minimization of the consumption of one resource subject to minimum requirements of system reliability and other resource constraints.

Gen *et al.* [21] and Yokota *et al.* [22] proposed a simple GA to solve reliability-optimization problems. They considered the reliability-optimization problem of a redundancy system with several failure modes originally presented by Tillman [23]. In 1995, Painton and Campbell [24] solved a reliability-optimization problem related to the design of a personal computer. The functional block diagram of the personal computer to be optimized has a series–parallel configuration and a GA was developed to identify the best system configuration by selecting the best combination of components to yield the highest fifth-percentile mean time between failures (MTBF) while remaining within a budget constraint.

In Coit and Smith [25], a GA is presented to analyze series–**parallel systems**, in which each subsystem is **k -out-of- n :G** to determine the **optimal**

design configuration given constraints on reliability, cost, and/or weight, and in [26] they presented a penalty-guided GA to solve the RAP. A different approach was proposed for the solution of the RAP using a combination of neural networks and GAs [27]. Later, Coit and Smith [28] presented a GA to solve the RAP when the objective is to maximize a lower percentile of the system time-to-failure distribution. In 2000, an approach to solve the optimal allocation of a plant design under conflicting safety and economic constraints based on the coupling of **Monte Carlo simulation** and GAs was proposed by Cantoni *et al.* [29].

An extensive overview of some of the methods that have been developed over the last years for solving various system reliability-optimization problems can be found in Kuo and Prasad [30]. A more recent discussion on the use of GAs for the solution of reliability-optimization problems is presented in Marseguerra *et al.* [31].

Genetic Algorithm Operators for Reliability Optimization. An effective GA for RAP depends on problem-specific solution encoding, fitness functions, and complementary crossover and mutation operators. The effectiveness of the crossover operator dictates the rate of convergence, while the mutation operator prevents the algorithm from prematurely converging to a local optimum. The number of children and mutants produced in each generation are tunable parameters that are generally held constant during a specific GA run. There are different possible encodings and operators for RAP; however, those used by Coit and Smith [25] are described in more detail.

Traditionally solution encodings have been a binary string [2]. For combinatorial optimization, an encoding using integer values can potentially be more efficient. For RAP [25], the m_i components are indexed in descending order in accordance with their reliability (i.e., 1 representing the most reliable, etc.). The solution encoding is a vector representation with $s \times n_{\max}$ positions. Each of the s subsystems is represented by $n_{\max}$ positions with the components, which form a particular solution listed according to their reliability index. An index of $m_i + 1$ is assigned to a position where an additional component was not used (i.e., $n_i < n_{\max,i}$). The subsystem representations are then placed adjacent to each other to complete the vector representation.

4 Genetic Algorithms in Reliability

As an example, consider a system with $s = 3$, $m_1 = 5$, $m_2 = 4$, $m_3 = 5$, and $n_{\max,i}$ predetermined to be 5. The following example,

$$\mathbf{v}_1 = (1 \ 1 \ 6 \ 6 \ 6 | 2 \ 2 \ 3 \ 5 \ 5 | 4 \ 6 \ 6 \ 6 \ 6) \quad (6)$$

represents a prospective solution with two of the most reliable components used in parallel for the first subsystem, two of the second most reliable and one of the third most reliable components used in parallel for the second subsystem, and one of the fourth most reliable components used for the third subsystem.

Each solution encoding (called a *genotype*) is associated with a solution to the original problem (called a *phenotype*). For each solution, a *fitness function* is determined. This may be the original objective function adjusted by penalty functions to enforce constraint feasibility.

The crossover breeding operator provides a thorough search of areas of the sample space already demonstrated to produce good solutions. Parents are selected (e.g., roulette wheel, tournament selection) based on their fitness function. There are different types of crossover operators including single-point, double-point, and uniform crossover. As an example [25], uniform crossover retains identical portions of the parents, and then randomly selects, with equal probability, from the two parents for those parts which differ. For example, consider the two following parent vectors ($\mathbf{v}_1$, $\mathbf{v}_2$) and resulting child vector ($\mathbf{v}_c$),

$$\left. \begin{array}{ll} \mathbf{v}_1 = (1 \ 1 \ 6 \ 7 \ 7 \ 3 \ 7 \ 7 \ 7 \ 7) \\ \mathbf{v}_2 = (2 \ 3 \ 6 \ 6 \ 7 \ 2 \ 2 \ 7 \ 7 \ 7) \\ \hline \mathbf{v}_c = (\text{x} \ \text{x} \ 6 \ \text{x} \ 7 \ \text{x} \ \text{x} \ 7 \ 7 \ 7) & - \text{identical indices} \\ \mathbf{v}_c = (2 \ 1 \ 6 \ 7 \ 7 \ 3 \ 2 \ 7 \ 7 \ 7) & - \text{random selection of nonmatching indices} \\ \mathbf{v}_c = (1 \ 2 \ 6 \ 7 \ 7 \ 2 \ 3 \ 7 \ 7 \ 7) & - \text{final arrangement} \end{array} \right\} \quad (7)$$

The mutation operator performs random perturbations to selected solutions. A predetermined number of mutations within a generation is set for each GA trial. Each value within the solution vector (which was randomly selected to be mutated) is changed with probability equal to the mutation rate.

The new solutions can either form a new population, or they can be combined with the previous population and undergo a culling operation. Often elitism is used to retain the solutions with the best fitness functions.

Multiple-Objective Optimization of Redundant Systems

Most actual system engineering optimization problems are multiobjective in nature. This means that these problems have designs with several criteria or design objectives. Most multiple-objective problems (MOPs) have no unique solution that can simultaneously optimize all objectives. Instead, MOPs have a set of potential (nondominated) solutions, the so-called Pareto-optimal set. Existing methods for the solution of MOPs involve either the consolidation of all objectives into a single-objective function, or the determination of a Pareto-optimal set.

Busacca *et al.* [32] proposed a multiobjective GA approach that was applied to a design problem with the aim to identify the optimal system configuration and components with respect to reliability and cost objectives. In 2006, Tian and Zuo [33] applied physical programming to a RAP for multistate systems. They sought to maximize system performance utility while minimizing system cost and system weight simultaneously. In their approach, the multiple objectives considered are aggregated into a single-objective function and a GA is used to solve the proposed physical programming optimization model.

In Taboada and Coit [34] and Taboada *et al.* [35], the RAP was formulated as a multiobjective problem with the system reliability to be maximized, and cost and weight of the system to be minimized as shown

in equations (8 and 9). The Pareto-optimal set was initially obtained using the fast elitist nondominated sorting genetic algorithm (NSGA)-II [36]. Then, the decision-making stage was performed with the aid of two pruning methods to reduce the size of the Pareto-optimal set and obtain a smaller representation of the multiobjective design space.

$$\max \left[\prod_{i=1}^s R_i(\mathbf{x}_i) \right], \min \left[\sum_{i=1}^s \sum_{j=1}^{m_i} c_{ij} x_{ij} \right],$$

$$\min \left[\sum_{i=1}^s \sum_{j=1}^{m_i} w_{ij} x_{ij} \right] \quad (8)$$

Subject to

$$1 \leq \sum_{j=1}^{m_i} x_{ij} \leq n_{\max,i} \quad \text{for } i = 1, 2, \dots, s \quad (9)$$

$$x_{ij} \in \{0, 1, 2, \dots\}$$

Multistate Optimization of Redundant Systems

As we earlier discussed, reliability is defined as the probability that a device is able to perform its intended functions satisfactorily under specified conditions for a specified period. However, traditional reliability assumes that a system and its components can only be in either a working or a failed state. This condition has facilitated the development of a robust and extensive theory to analyze system performance. However, in some cases traditional reliability theory fails to represent the true model of the system. Failure to acknowledge this situation can represent a major deficiency since systems can have a range of intermediate states that are not accounted for in the reliability estimation.

To describe the satisfactory performance of a device, it may be necessary to use more than two levels of satisfaction, e.g., excellent, average, and poor. Multistate reliability has been proposed as a complementary theory to cope with the problem of analyzing systems where traditional reliability theory and models become insufficient. Then, in a multistate system, both the system and its components are allowed to experience more than two possible states, e.g., completely working, partially working or partially failed, and completely failed.

The universal moment generating function (UMGF) can be used to determine the availability of a redundant multistate system, as in Ushakov [37, 38]. In 1998, Levitin *et al.* [39] used a GA for solving the multistate RAP, where the system and its components have a range of performance levels. **System availability** is represented by a multistate availability function, which extends the binary state availability. They evaluated the multistate system availability based on the UMGF. Later, Levitin *et al.* [40] addressed the multistage expansion

problem for multistate series-parallel systems. In this problem, the system-study period is divided into several stages. A more comprehensive explanation completely devoted to multistate system reliability analysis and optimization can be found in Lisnianski and Levitin [4].

Summary and Other Important Related Topics

We have presented an overview of some of the most relevant research that has been done on the use of GAs in **reliability optimization** of redundant systems. As we have described, GAs are powerful tools for effectively solving reliability-optimization problems. However, there are many other applications of GAs in reliability-related problems that were not addressed in this article and are worthy of comment, such as the use of GAs on reliability optimization with alternative design, reliability optimization with time-dependent reliability, reliability optimization with interval coefficients, bicriteria reliability optimization, and reliability optimization with fuzzy goals and fuzzy constraints. For those who are more interested on the topic, reference [41] presents a state-of-the-art survey that covers several GA-based approaches for various reliability-optimization problems.

References

- [1] Gen, M. & Cheng, R. (1997). *Genetic Algorithms and Engineering Design* H.R. Parsaei ed, John Wiley & Sons.
- [2] Holland, J. (1975). *Adaptation in Natural and Artificial Systems*, Univeristy of Michigan Press, Ann Arbor.
- [3] Konak, A., Coit, D. & Smith, A. (2006). Multi-objective optimization using genetic algorithms: a tutorial, *Reliability Engineering and System Safety* **91**(9), 992–1007.
- [4] Lisnianski, A. & Levitin, G. (2003). *Multi-State System Reliability. Assessment, Optimization and Applications, Series on Quality, Reliability and Engineering Statistics, Vol. 6*, M. Xie, T., Bendell & A.P. Basu eds, World Scientific.
- [5] Goldberg, D. (1989). *Genetic Algorithms in Search, Optimization and Machine Learning*, Addison-Wesley, Reading.
- [6] Chern, M.S. (1992). On the computational complexity of reliability redundancy allocation in a series system, *Operations Research Letters* **11**, 309–315.
- [7] Bellman, R.E. & Dreyfus, E. (1958). Dynamic programming and reliability of multicomponent devices, *Operations Research* **6**, 200–206.

- [8] Fyffe, D.E., Hines, W.W. & Lee, N.K. (1968). System reliability allocation and a computational algorithm, *IEEE Transactions on Reliability* **17**, 64–69.
- [9] Sakawa, M. (1978). Multiobjective optimization by the surrogate worth trade-off method, *IEEE Transactions on Reliability* **27**, 311–314.
- [10] Bulfin, R.L. & Liu, C.Y. (1985). Optimal allocation of redundant components for large systems, *IEEE Transactions on Reliability* **34**, 241–247.
- [11] Dhingra, A.K. (1992). Optimal apportionment of reliability and redundancy in series systems under multiple objectives, *IEEE Transactions on Reliability* **41**(4), 576–582.
- [12] Kulturel-Konak, S., Smith, A. & Coit, D.W. (2003). Efficiently solving the redundancy allocation problem using tabu search, *IEEE Transactions* **35**(6), 515–526.
- [13] Glover, F. (1989). Tabu search part I, *ORSA Journal on Computing* **1**, 190–206.
- [14] Glover, F. (1990). Tabu search part II, *ORSA Journal on Computing* **2**, 4–32.
- [15] Dorigo, M., Di Caro, G. & Gambardella, L.M. (1999). Ant algorithms for discrete optimization, *Artificial Life* **5**(2), 137–172.
- [16] Blum, C. & Roli, A. (2003). Metaheuristics in combinatorial optimization: overview and conceptual comparison, *ACM Computing Surveys* **35**(3), 268–308.
- [17] Majety, S.R., Venkatasubramanian, S. & Smith, A. (1996). Optimal reliability allocation in series-parallel systems from components' discrete cost reliability data sets: a nested simulated annealing approach, in *Proceedings of the Fifth International Industrial Engineering Research Conference*, Minneapolis, May pp. 435–440.
- [18] Yun-Chia, L. & Smith, A.E. (2004). An ant colony optimization algorithm for the redundancy allocation problem (RAP), *IEEE Transactions on Reliability* **53**(3), 417–423.
- [19] Nahas, N. & Nourelfath, M. (2005). Ant system for reliability optimization of a series system with multiple-choice and budget constraints, *Reliability Engineering and System Safety* **87**, 1–12.
- [20] Muñoz Zavala, A., Villa Diharce, E. & Hernandez Aguirre, A. (2005). Particle evolutionary swarm for design reliability optimization, in *EMO*, C.A. Coello Coello, A. Hernández Aguirre & E. Zitzler, eds, Springer-Verlag Berlin, Heidelberg, pp. 856–869.
- [21] Gen, M., Ida, K. & Taguchi, T. (1993). Reliability optimization problems: a novel genetic algorithm approach, Technical report, ISE93-5, Ashikawa Institute of Technology, Ashikawa.
- [22] Yokota, T., Gen, M. & Ida, K. (1995). System reliability optimization problems with several failure modes by genetic algorithm, *Japanese Journal of Fuzzy Theory and Systems* **7**(1), 117–185.
- [23] Tillman, F. (1969). Optimization by integer programming of constrained reliability problems with several modes of failure, *IEEE Transactions on Reliability* **18**, 47–53.
- [24] Painton, L. & Campbell, J. (1995). Genetic algorithms in optimization of system reliability, *IEEE Transactions on Reliability* **44**(2), 172–178.
- [25] Coit, D.W. & Smith, A. (1996a). Reliability optimization of series-parallel systems using a genetic algorithm, *IEEE Transactions on Reliability* **45**(2), 254–260.
- [26] Coit, D.W. & Smith, A. (1996b). Penalty guided genetic search for reliability design optimization, *Computers and Industrial Engineering* **30**(4), 895–904.
- [27] Coit, D.W. & Smith, A. (1996c). Solving the redundancy allocation problem using a combined neural network/genetic algorithm approach, *Computers and Operations Research* **23**(6), 515–526.
- [28] Coit, D.W. & Smith, A. (1998). Redundancy allocation to maximize a lower percentile of the system time-to-failure distribution, *IEEE Transactions on Reliability* **47**(1), 79–87.
- [29] Cantoni, M., Marseguerra, M. & Zio, E. (2000). Genetic algorithms and Monte Carlo simulation for optimal plant design, *Reliability Engineering and System Safety* **68**, 29–38.
- [30] Kuo, W. & Prasad, V.R. (2000). An annotated overview of system-reliability optimization, *IEEE Transactions on Reliability* **49**(3), 323–330.
- [31] Marseguerra, M., Zio, E. & Martorelli, S. (2006). Basics of genetic algorithms optimization for RAMS applications, *Reliability Engineering and System Safety* **91**, 977–991.
- [32] Busacca, P.G., Marseguerra, M. & Zio, E. (2001). Multi-objective optimization by genetic algorithms: application to safety systems, *Reliability Engineering and System Safety* **72**, 59–74.
- [33] Tian, Z. & Zuo, M.J. (2006). Redundancy allocation for multi-state systems using physising programming and genetic algorithms, *Reliability Engineering and System Safety* **91**, 1049–1056.
- [34] Taboada, H.A. & Coit, D.W. (2006). Data clustering of solutions for multiple objective system reliability optimization problems, *Quality Technology & Quantitative Management Journal* **4**(2), 35–54.
- [35] Taboada, H., Baheranwala, F., Coit, D.W. & Wattanapongsakorn, N. (2006). Practical solutions of multi-objective optimization: an application to system reliability design problems, *Reliability Engineering and System Safety* **92**(3), 314–322.
- [36] Deb, K., Agrawal, S., Pratap, A. & Meyarivan, T. (2002). A fast and elitist multi-objective genetic algorithm: NSGA-II, *IEEE Transactions on Evolutionary Computation* **6**(2), 182–197.
- [37] Ushakov, I.A. (1986). A universal generating function, *Soviet Journal of Computer and Systems Sciences* **24**, 37–49.
- [38] Ushakov, I.A. (1987). Optimal standby problems and a universal generating function, *Soviet Journal of Computer and Systems Sciences* **25**(4), 79–82.
- [39] Levitin, G., Lisnianski, A., Ben-Haim, H. & Elmakis, D. (1998). Redundancy optimization for series-parallel

-
- multi-state systems, *IEEE Transactions on Reliability* **47**(2), 165–172.
- [40] Levitin, G. (2000). Multi-state series-parallel system expansion-scheduling subject to availability constraints, *IEEE Transactions on Reliability* **49**(1), 71–79.
- [41] Yen, M. & Yun, Y.S. (2006). Soft computing approach for reliability optimization: state-of-the-art survey, *Reliability Engineering and System Safety* **91**, 1008–1026.

DAVID W. COIT AND HEIDI A. TABOADA

Reliability, Safety, and Risk Management

Safety and Reliability

A close association between safety and reliability has existed since the earliest times, Cox and Tait [1]. Throughout the history of engineering, reliability improvement, arising as a direct consequence of the analysis of failure, has been a central feature in the development of safe systems (Smith [2]). Until the advent of modern scientific theories, technical progress was made by a sophisticated process of “test and correct”. Engineers and designers learned not only by their own mistakes but also from other people’s misfortunes. This process was quite successful and a significant progress was made, as evidenced by the construction of twelfth- and thirteenth-century buildings throughout Europe and the rest of the developed world. However, the relationship between safety and reliability was intensified as industrial processes proliferated. During the Industrial Revolution (Hobsbawm [3], Ashton [4]) new sources of power, using water or steam, not only gave great opportunities for the rapid development of manufacturing technology, but also presented safety challenges.

At this time scientific developments were strongly deterministic. The approach to safety and reliability during this time was along similar deterministic lines. Structures were designed to a predetermined “load factor” or *factor of safety*, this being the ratio of the load predicted to cause collapse of the structure to the normal operating load. The factor was made large enough to ensure that the structure operated safely even if significant corrosion was present. Fixing the magnitude of the load factor was a relatively arbitrary process, and led at times to considerable argument among practitioners.

Toward the end of the nineteenth century, science had begun to make use of statistical and probabilistic techniques (for example, in gas kinetics). This probabilistic approach entered the safety and reliability field on a large scale as attempts were made to operate electronic and other delicate equipment under battle conditions during the Second World War. Application of the well-tried factor of safety was no longer able to provide a solution, and under the harsh operating

conditions encountered on board ship or in combat aircraft, reliability (or rather, unreliability) became a major problem. Typical airborne radar, for instance, would do well to operate continuously for 1 hour without failure (Wickens [5]). In these circumstances it was necessary to study the causes and effects of component breakdown in order to improve reliability.

At the same time it could be argued that the improvements in **component reliability** and developments in economic prosperity have lead many practitioners to consider a sociotechnical systems approach and include the “human” factor in the process of safety improvement (see Figure 1). The most obvious impetus for this interest has been a growing public concern over the great costs of human error: the Tenerife air disaster in 1977, Three Mile Island 2 years later, Bhopal in 1984 with its considerable loss of lives, and Chernobyl in 1986 with its implications for the public image of the nuclear-power-generating industry worldwide.

An additional spur to the developments in our understanding of human error has come from theoretical and methodological developments in cognitive psychology. It has become increasingly apparent that to provide an adequate picture of control processes psychologists must explain not only correct performance, but also more predictable varieties of human error.

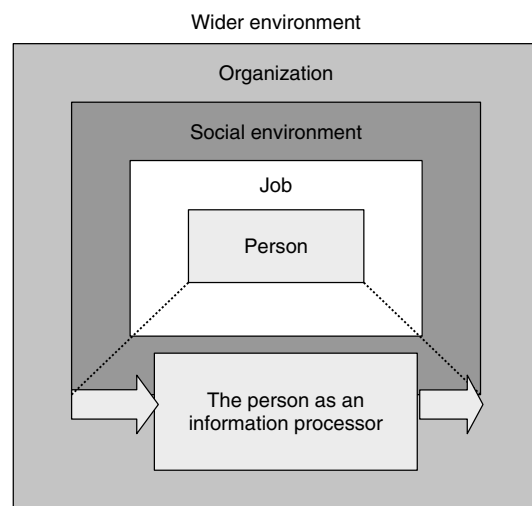


Figure 1 People-related issues – a broader view [[6], © Elsevier, 1998.]

The Human Factor

The study of human factors in systems reliability, through the application of psychology, ergonomics, and human factors engineering, has grown considerably over recent decades. Early researchers in this area complained of the seeming lack of interest in human error by their fellow psychologists, *crammed as psychological writings are, and most needs to be, with allusions to errors in an incidental manner, they hardly ever arrive at considering these profoundly, or even systematically*, Spearman [7].

Human factors have been implicated in the etiology of many accidents within high-reliability organizations. The work of human factors engineers and ergonomists has developed a greater understanding of human task performance and the interactions between humans and complex systems. The type and degree of human participation, especially in “high” consequence areas, has been a matter of increasing concern.

A direct manifestation of the current link between safety and reliability is illustrated through the safety regulation and the requirements for submission of safety case (for further details see www.hse.gov.uk).

The Systems Approach

The systems approach to safety management is based on the applications of general systems theory; the approach adopted by the author integrates engineering, management procedures, and human factors. The approach provides a framework for analyzing, modeling, and managing work and work functions. There are three important characteristics of “systems” that can be applied to a simple safety systems model. First, the components are connected in an organized way and changes in one component affect the others. Second, the behavior of the system changes if any one component is excluded from the system. Third, the organized assembly of components performs some function. Safety systems also have an adaptive response to their wider environments. Figure 2 outlines an objective setting model of the safety of system X.

Risk Management

Risk management is a technique utilized by organizations and public bodies to increase safety and

reliability, and minimize losses. The process involves the identification, evaluation, and control of risks. Risk identification can be achieved using a number of techniques.

Risk evaluation encompasses the measurement and assessment of risk. Implicit in the process is the need for sound decision making on the nature of potential sociotechnical systems and their predicted reliability. Decisions on the acceptability of risk are dependent on a number of factors, which include social, economic, political, and legislative concerns.

The final stage in the management of risk is risk control. Risk control strategies can be classified into four categories: risk avoidance, retention, transfer, and reduction. Risk avoidance involves a conscious decision on the part of the organization to avoid a particular risk by discontinuing the operation that is producing the risk. Risk retention may occur with knowledge (a deliberate decision to retain the risk by, for example, self financing) or without knowledge (occurs when risks have not been identified). Risk transfer is a conscious transfer of risk, i.e., to another organization *via* insurance. Finally risk reduction is the management of systems to reduce risks. A number of techniques, concepts, and strategies in relation to technology, management systems, and human factors are described in the following section.

Analysis Techniques and the Control of the System

Human Reliability Techniques

Human reliability should be an important consideration at all stages of systems design and its later implementation and management. Human reliability can be enhanced by appropriate selection, training, and continuing development of all personnel. Additional enhancement may be obtained from heightened awareness, in appropriate staff, of the design features of plant and equipment that may lead to operational or maintenance errors. Attempts to quantify human reliability have been incorporated into systems thinking since the 1950s – although the majority of the work has taken place in high-reliability organizations, it has been related to the probabilities of human error for critical functions and particularly in emergency situations (Humphreys [8]). The most common measure of human reliability is the human error probability

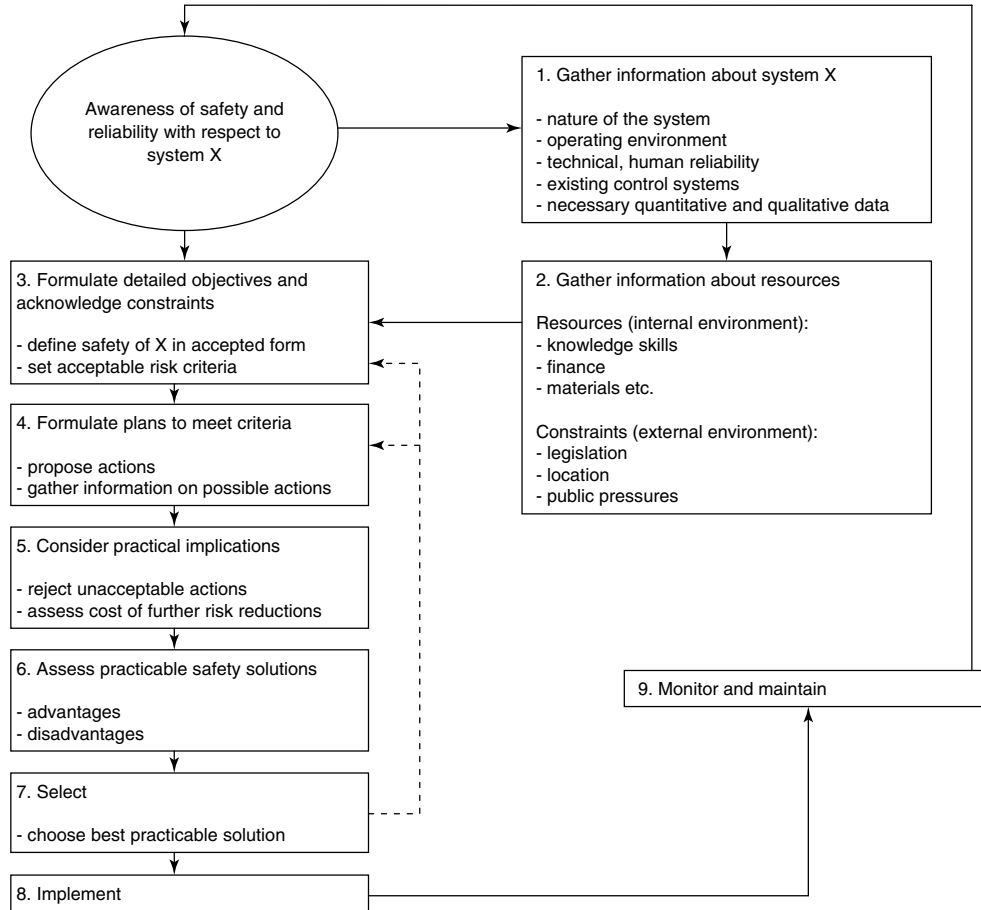


Figure 2 An objective setting model of the safety of system X [© Elsevier, 1998.]

(HEP), which is the probability of an error occurring during a specified task. HEP is estimated from the ratio of errors committed to the total number of opportunities for error as follows:

$$HEP = \frac{\text{Number of human errors}}{\text{Total number of opportunities for error}} \quad (1)$$

The successful performance probability of a task (or the task reliability) can generally be expressed by $1 - HEP$. Table 1 lists the common human reliability assessment techniques.

The majority of human reliability assessment techniques are based in part on behavioral psychology. They have been derived from empirical models using statistical inference (but are not

always adequately validated) and are dependent on expert judgments, which can be subject to **bias**. One of the most commonly used human reliability assessment techniques is the technique for human error rate prediction (Swain and Guttman [9]); the technique can be broken down into five distinct phases:

1. a systems description, including goal definition and functions;
2. a job and task analysis by personnel to identify likely error situations;
3. an estimation of the likelihood of each potential error as well as of it being undetected (taking account of the performance-shaping factors);
4. an estimate of the consequences of any undetected error; and

Table 1 Human reliability assessment techniques

HRA ^(a) technique
APJ ^(b)
PC ^(c)
TESEO ^(d)
THERP ^(e)
HEART ^(f)
IDA ^(g)
SLIM ^(h)
HCR ⁽ⁱ⁾

^(a) HRA: human reliability assessment

^(b) APJ: absolute probability judgment

^(c) PC: paired comparison

^(d) TESEO: *tecnica empirica stima operatori*

^(e) THERP: technique for human error rate prediction

^(f) HEART: technique for human assessment and reduction technique

^(g) IDA: influence diagram approach

^(h) SLIM: success likelihood index method

⁽ⁱ⁾ HCR: human cognitive reliability method

5. suggested and evaluated changes to the system in order to increase success probability.

Hazard Analysis Techniques

The process of hazard analysis is an overarching technique that allows the systems designer and operators to ensure compliance with safety regulations. Hazard analysis will help analysts understand the hazards resulting from failure better and requires an understanding of the physical and sometimes the chemical properties of the materials involved. Moreover, hazard analysis requires an understanding of the drivers of human behaviors. An important example of the bottom-up approach is the failure modes and effects analysis (FMEA). This method is of particular use when examining the performance of relatively simple components, for determining which types of failure are of danger and which are of safety, and finally for calculating overall failure rates to the two states for the complete component.

Fault tree analysis (or HAZard ANalysis (HAZAN)) provides a powerful and commonly used example of a top-down procedure. The fault tree analysis begins with a hazardous outcome, known as the *top event*. The process formulates a logic diagram indicating the sequence of events leading to the top event; component probabilities are

then used to calculate the probability of the top event.

Performance monitoring and evaluation is a fundamental part of risk control. Measurement of performance involves comparisons with standards. Performance criteria should thus be established for plant, personnel, and management activities such as selection and training for the effectiveness of the overall system. Performance monitoring is therefore not restricted to the recording of accidents and incidents, which are the traditional parameters of “safety” but it is also an “active” process. This philosophy of safety measurement is concerned with actively seeking information on systems with a view to promoting positive actions.

Strategies for the reduction and management of errors have been developed by practitioners and researchers alike. Four potential error-handling strategies that draw upon some of the human factors issues considered in the previous sections include

1. intelligent decision support systems, i.e., safety-parameter display systems;
2. memory aids for maintenance personnel, i.e., portable interactive maintenance auxiliary (Reason [10]);
3. safety by design, i.e., designing out errors at the design phase of development; and
4. error management, i.e., promoting the positive and mitigating the negative effects of training errors in a systematic fashion (Frese [11]).

Summary

This paper has presented a brief exposition on the connection between safety and reliability. The paper has presented a systems approach to safety, reliability, and risk management that demonstrates the fundamental importance of good “people” management; this approach encourages organizations to consider the “human” factor and to develop systems which not only minimize accidents and incidents, but also make better safety inevitable.

References

- [1] Cox, S. & Tait, R. (1998). *Safety, Reliability and Risk Management: An Integrated Approach*, 2nd Edition, Butterworth-Heinemann, Oxford.

-
- [2] Smith, D. (2005). *Reliability, Maintainability and Risk: Practical Methods for Engineers Including Reliability Centred Maintenance and Safety-Related Systems*, Elsevier, Burlington.
 - [3] Hobsbawm, J. (1994). *The Age of Extremes: The Short Twentieth Century, 1914–1991*, Abacus, London.
 - [4] Ashton, T. (1997). *The Industrial Revolution (1760–1830)*, Oxford Paperbacks, Oxford.
 - [5] Wickens, C. (1992). *Engineering Psychology and Human Performance*, 2nd Edition, HarperCollins, New York.
 - [6] Cox, S. & Cox, T. (1996). *Safety Systems and People*, Butterworth-Heinemann, Oxford.
 - [7] Spearman, C. (1928). The origin of error, *Journal of General Psychology* **1**, 29.
 - [8] Humphreys, P. (1988). *Human Reliability Assessors Guide*, UKAEA, Culcheth.
 - [9] Swain, A.D. & Guttman, H.E. (1983). *Handbook of Human Reliability Analysis with Emphasis on Nuclear Power Plant Applications*, NUREG/CR-1278, Sandia National Laboratories.
 - [10] Reason, J. (1990). *Human Error*, Cambridge University Press, New York.
 - [11] Frese, M. (1987). A theory of control and complexity: Implications for software design and integration of computer system into the work place, in M. Frese, E. Ulich & W. Dzida, eds, *Psychological Issues of Human-Computer Interaction at the Work Place*, North Holland, Amsterdam, pp. 313–337.

SUE COX

Block Replacement

Introduction and Historical Perspective

The scientific interest in maintenance management originated only a few decades ago. The basis for the scientific support of maintenance is found in reliability engineering. Books by Barlow and Proschan [1] and Jorgenson and McCall [2] indicated the start of scientific interest in maintenance problems. In the early stages, maintenance was almost exclusively restricted to correcting failures. Later in the 1960s **preventive maintenance** programs were set up with the intention of avoiding failures. A main element in these programs is the preventive replacement of components. Although today there are many condition-based techniques to predict a required replacement, time-based preventive replacements (or hard-time replacements) are still common practice in many systems, like aircrafts, because of safety reasons. The timing of preventive replacements is therefore one of the key problems in this respect. The development of some simple but insightful models describing aging phenomena of technical components, as well as the choice of appropriate actions to cope with these phenomena, showed that a scientific approach can be useful. In the 1970s and early 1980s, the basic models were extended in several directions.

During the last 15 years substantial progress has been made in developing quantitative decision support systems for maintenance management of complex systems. Sophisticated systems are in operation nowadays in the oil industry (both for maintenance of refineries as well as offshore installations), road and railways maintenance, and electric power generation. They include the block-replacement model in one way or another.

Block Replacement

Block replacement is the preventive replacement of aging items after fixed intervals of time. Failures of items in between are covered by individual failure replacements and do not change the timing of the block replacement. Alternatives to block replacement are age replacement, where the replacement timing is based on the age of each component, and opportunity maintenance, where the timing of the maintenance

depends on the occurrence of opportunities for doing maintenance. For example, if a block replacement is scheduled on the first of May and the component involved fails and is replaced in February, the preventive replacement will still be done in May. The age-replacement policy, however, would reschedule the replacement. The block-maintenance policy is usually based on calendar time, but it can also be based on run-hours, kilometers driven, or number of takeoffs and landings; in this contribution we only treat the time base and assume that the other cases can be dealt with by transformations. We will next discuss the following issues relating to block replacement, *viz* characterizing the optimal block time, computation of the optimal block-replacement interval, comparison with other policies, and extensions.

Modeling

Consider a component which is subject to failures. Let the random variable X denote the time to failure, and assume that a cumulative distribution function (**cdf**) $F(\cdot)$ and probability distribution function (**pdf**) $f(\cdot)$ of X are available after statistical (data) analysis. Next assume that all components used for replacement have independent and identically distributed (i.i.d.) failure times. Let c_p, c_f be the cost of a preventive, corrective replacement respectively and let T be the time interval for block replacement. As each block replacement is a renewal of the system, the renewal reward theorem can be applied to derive a formula for the long-run average cost. A key quantity in this respect is the expected number of failures in the interval $[0, T]$. This is usually referred to as the *renewal function*, here indicated by $M(T)$. Now the long-run average costs $g(T)$ of a block-replacement policy with interval T are given by

$$g(T) = \frac{c_p + c_f M(T)}{T} \quad (1)$$

The first question is under which conditions there exists a finite optimum replacement interval, or in other words, whether preventive replacement is cost effective. The derivative of $g(T)$ is given by

$$\begin{aligned} g'(T) &= \frac{1}{T^2} \{c_f m(T)T - (c_p + c_f M(T))\} \\ &= \frac{1}{T} \{c_f m(T) - g(T)\} \end{aligned} \quad (2)$$

2 Block Replacement

where $m(T)$ indicates the renewal density $M'(T)$. To see whether $g'(T)$ has a zero point, we need to know more about the behavior of the renewal function. It is well known that under mild conditions we have the following asymptotic behavior of the renewal function

$$\lim_{T \rightarrow \infty} \left(M(T) - \frac{T}{E(X)} \right) = \frac{1}{2}(c_X^2 - 1) \quad (3)$$

where c_X^2 denotes the squared coefficient variation of X , i.e. $\text{Var}(X)/(E(X))^2$. Next remark that $\lim_{T \rightarrow \infty} m(T) = \frac{1}{E(X)}$ and that $\lim_{T \rightarrow \infty} g(T) = \frac{c_f}{E(X)}$. This does not indicate whether there exists a zero point of $g'(T)$ and another condition needs to be made. It appears that requiring

$$\lim_{T \rightarrow \infty} \left(M(T) - \frac{T}{E(X)} \right) > \frac{c_p}{c_f} \quad (4)$$

or, equivalently

$$\frac{c_p}{c_f} < \frac{1}{2}(c_X^2 - 1) \quad (5)$$

is sufficient for the existence of a finite optimum. For instance, if we assume that X follows a Weibull distribution, with cdf $F(T) = 1 - e^{-(T/\lambda)^\beta}$, then the requirement comes down to the failure rate being increasing (which is the case if the shape parameter $\beta > 1$) and the cost of preventive replacement be sufficiently lower than that of failure replacement. There may be multiple minima, but the first minimum is likely to be the overall minimum. Numerical evaluation of block-replacement policies does require the computation of the renewal function.

Renewal Function

Several characterizations have been developed for the renewal function (see e.g. Tijms [3]). We have a recursive integral formulation and an infinite series representation, i.e.

$$M(T) = F(T) + \int_0^T M(T-x) dF(x) \quad (6)$$

and

$$M(T) = \sum_{n=1}^{\infty} F^{(n)}(T) \quad (7)$$

where $F^{(n)}(T)$ denotes the n -fold convolution of $F(T)$, which can be interpreted as the probability of n or more failures in $[0, T]$. As lifetime distributions should have a positive support, the **normal distribution** has to be excluded. It appears that for many popular distributions in reliability, like the Weibull, taking the convolution can only be done numerically. For some specific distributions analytical expressions are available for the renewal function, viz

1. if X is exponentially distributed with mean $1/\lambda$, then $M(T) = \lambda T$;
2. if X follows a Coxian-2 distribution, that is with probability p it is an exponential distribution with rate λ_1 and with probability $(1-p)$ it is the sum of two exponential distributions with rates λ_1 and λ_2 , respectively, and so its pdf is given by

$$f(t) = \begin{cases} p\lambda \exp(-\lambda t) + (1-p)\lambda^2 t \times \exp(-\lambda t), & \lambda_1 = \lambda_2 = \lambda \\ \frac{p\lambda_1 - \lambda_2}{\lambda_1 - \lambda_2} \lambda_1 \times \exp(-\lambda_1 t) + \left(1 - \frac{p\lambda_1 - \lambda_2}{\lambda_1 - \lambda_2}\right) \lambda_2 \times \exp(-\lambda_2 t), & \lambda_1 \neq \lambda_2 \end{cases} \quad (8)$$

then its renewal function is given by

$$M(t) = \frac{\lambda_1 \lambda_2}{\lambda_1(1-p) + \lambda_2} t - \frac{\lambda_1(1-p)(\lambda_2 - p\lambda_1)}{(\lambda_1(1-p) + \lambda_2)^2} \times [1 - \exp\{-(\lambda_1(1-p) + \lambda_2)t\}] \quad (9)$$

For some distributions we can calculate the convolution easily. This is the case for Erlang distributions. For other distributions we have to apply numerical procedures. We would like to mention discretization of the distribution and using a recursive scheme for the convolutions. Other approaches work with cubic splines. There also exist very simple approximations, like $M(T) \approx F(T)$, for small T , but that does not work in the optimization as it will give an infinite optimal interval. It is however a good approximation for small T . This has led Smeitink and Dekker [4] to propose the following approximation

$$M(T) = F(T) + [\tilde{M}(T) - \tilde{F}(T)] \quad (10)$$

where for $\tilde{F}(T)$ a suitable distribution is chosen for which the renewal function can easily be evaluated and which has the same asymptotic behavior. The first can be done by choosing either mixed Erlang distributions or the hyper exponential and the latter by applying a two-moment fit.

Comparison with Other Policies

Age replacement is a natural competitor to block replacement. For a single component, age replacement is indeed easy to formulate: replace if the age of the part exceeds a critical age T . It avoids replacements of parts which have just been replaced in a failure replacement. Berg [5] showed that the optimal age-replacement policy has always lower average costs than the optimal block-replacement policy. The maximal difference is however, less than 10%. The comparison becomes more complex however, if multiple components have to be replaced. Several variants of multicomponent age replacement exist if there are set up costs involved. Variants can prescribe replacement of unfailed components also upon a failure or replacing a fixed or variable group preventively. However, in all cases the replacement policies depend on the individual component ages and such policies are much more complex, both in terms of description as well as in evaluation (see e.g. Dekker and Roelvink [6]). For multiple components block replacement however remains equally simple.

Although the term *block replacement* refers to component replacements, the model can also be used for cases where preventive maintenance is done which leads to a renewal of the component addressed, without really replacing it. Next the idea of fixed intervals can also be used for inspections.

Both block- and age-replacement policies involve a lifetime distribution which is obtained from data analysis or expert judgment. Quite often these distributions exhibit a large variation. Hence they are not that efficient and a lot of remaining life is thrown away. If one has specific information on the likelihood of the component on hand, one can better plan the replacement and avoid unnecessary preventive replacements. This is possible in the so-called condition-based maintenance, which focuses on monitoring techniques to predict failures, such as vibration monitoring, luboil analysis, and so on.

Practical Considerations

Although theoretically block replacement can be improved by age-types of replacement, it also has some other advantages. For example in age replacement one has to keep track of failure replacements of the individual components, which is more work than keeping track of the last execution of the preventive maintenance (PM) activity. A typical example is the replacement of fluorescent lamps in buildings. Another advantage of block replacement is the fact that it can be planned in advance, implying that all the necessary spare parts can be ordered in time and provisions can be made to ensure availability of appropriate manpower.

Extensions

Several extensions have been suggested to overcome the shortcoming of block replacement to replace relatively young components preventively. They have been termed *modified block replacement* (MBRP); see Archibald and Dekker [7]. Various options are available. Some authors suggest taking old unfailed components in failure replacements to lower costs and others to make the replacement flexible. In the so-called (t, T) MBRP one considers a preventive replacement every T time units, but at these times components are not replaced preventively if their age is below t time units. These policies are however, more difficult to analyze, but reasonable approximations are obtained by using the same interval as block replacement and taking t as about $T/2$.

References

- [1] Barlow, R.E. & Proschan, F. (1965). *Mathematical Theory of Reliability*, John Wiley & Sons, New York.
- [2] Jorgenson, D.W., McCall, J.J. & Radner, R. (1967). *Optimal Replacement Policy*, North Holland, Amsterdam.
- [3] Tijms, H.C. (2003). *A First Course in Stochastic Models*, John Wiley & Sons, New York.
- [4] Smeitink, E. & Dekker, R. (1990). A simple approximation to the renewal function, *IEEE Transactions on Reliability* **39**, 71–75.
- [5] Berg, M. (1976). A proof of optimality for age replacement policies, *Journal of Applied Probability* **13**, 751–759.
- [6] Dekker, R. & Roelvink, I.F.K. (1995). Marginal cost criteria for group replacement, *European Journal of Operational Research* **84**, 467–480.

4 Block Replacement

- [7] Archibald, T. & Dekker, R. (1996). Modified block replacement for multiple components, *IEEE Transactions on Reliability* **45**, 75–83.

Further Reading

Dekker, R. (1995). Integrating optimisation, priority setting, planning and combining of maintenance activities, *European Journal of Operational Research* **82**, 225–240.

Related Articles

Maintenance and Markov Decision Models; Multicomponent Maintenance; Multivariate Age and Multivariate Renewal Replacement.

ROMMERT DEKKER

System Downtime Distributions

Introduction

By system downtime we mean a period during which a system is not in a condition to perform its required function (see the glossary in [1]). In this article we examine the probability distribution of the system downtime, or system downtime distribution (SDD), considering two types of downtimes: the downtime in an interval and the downtime after a failure. The downtime in an interval is in general composed by the sum of the lengths of nonconnected subintervals during which the system is not functioning, whereas the downtime after a failure is the time it takes to complete the repair of the system and bring it back to the operating state.

Knowledge of the SDD can be useful for cost or risk analysis. The SDD in an interval would be used to calculate the expected cost of the out-of-service condition during the useful life of a system when the cost function is not a linear function of time. The 99th percentile of the SDD after a failure would give a repair time that is exceeded only with 1% probability, which is a useful information if long repair times involve some kind of risk (economic or of other nature).

This article is based mainly on the book by Aven and Jensen [2], where the authors give analytic formulas for both finite-time and asymptotic SDD. Such explicit solutions are obtained by assuming that the system can be partitioned into repairable and independent subsystems (or components) that evolve following an alternating renewal process [3, Chapter 3]. To our knowledge, reference [2] is the only systematic account giving a framework for addressing SDDs in complex systems. In [4] a dependence between failure time and repair time is allowed for, and the Laplace transform of the downtime in an interval is found, but for a system with one component only.

The expected downtime in an interval is related to system unavailability. In the next section we present a result on this, since it holds equally for one-component and multicomponent systems. Instead, in the successive sections we characterize the entire distribution of the downtime dealing with one-component systems and multicomponent systems

separately, beginning with the former, which are simpler.

It will be helpful to list some general notations at this point. More specialized notations will be introduced when needed. Let us denote by

- Y_u : a random variable representing the amount of time in $[0, u]$ during which the system is down;
- $\Phi(t)$: the structure function of the system, taking value 1 if the system is functioning at time t and value 0 otherwise;
- $A(t) = \Pr(\Phi(t) = 1)$: the availability of the system at time t , corresponding to the probability that the system is in operation;
- $F_i(\cdot)$ and $G_i(\cdot)$: the cumulative distribution functions of the failure time and of the repair time of component i , respectively, with finite expected values μ_{F_i} and μ_{G_i} , and variances $\sigma_{F_i}^2$ and $\sigma_{G_i}^2$.

In the systems we will be dealing with, the generic component i follows an alternating renewal process, with failure time distribution F_i and repair time distribution, G_i . The system starts with all components in operation. We assume that F_i is absolutely continuous with density f_i , whereas G_i can be any type of distribution. If the system has one component only, we suppress the index i .

Expected System Downtime in an Interval

In [2, Theorem 36] using the fact that the uptime in an interval is just the integral of the structure function of the system, one gets

$$Y_u = \int_0^u (1 - \Phi(t)) dt \quad (1)$$

and it can be shown that

$$E(Y_u) = \int_0^u (1 - A(t)) dt = \int_0^u \bar{A}(t) dt \quad (2)$$

This result is not very useful in practice, because the unavailability $\bar{A}(t)$ cannot be calculated explicitly even in the simplest cases. If there is only one component, an upper bound is given by

$$\bar{A}(t) \leq [\sup_{s \leq t} \lambda(s)] \mu_G \quad (3)$$

where $\lambda(s) = f(s)/\bar{F}(s) = f(s)/(1 - F(s))$ is the hazard function of F [2, Proposition 34]. As an

2 System Downtime Distributions

example, when F is Weibull, $\lambda(s) = \beta\lambda(\lambda s)^{\beta-1}$, for positive constants λ and β . If a cost function is proportional to the downtime through a constant C , an upper bound for the expected cost in the interval $[0, u]$ is given by $C \cdot (\lambda u)^\beta \cdot \mu_G$, for $\beta \geq 1$.

If a more general result is needed, we must resort to asymptotic approximations. Obviously, as $u \rightarrow \infty$, the expected downtime in $[0, u]$ will tend to infinity. However, the expected portion of time the system is down has the limiting unavailability as its limit, that is [2, Theorem 36],

$$\lim_{u \rightarrow \infty} \frac{E(Y_u)}{u} = \lim_{t \rightarrow \infty} \bar{A}(t) \equiv \bar{A} \quad (4)$$

Therefore, when the cost is proportional to the downtime, the expected cost over a long period of time can be approximated by $C \cdot u \cdot \bar{A}$, and the problem is now finding the limiting availability.

The availability topic is covered elsewhere in this encyclopedia (*see System Availability*). Very briefly, the limiting availability of a system of independent components is a function of the limiting availabilities of the components [2, Theorem 39]. The form of this function can be determined either from the **fault tree** or from the reliability block diagram representations of the systems. For example, the limiting availability of a series system with n independent components is given by $A = \prod_{i=1}^n A_i$, and $\bar{A} = 1 - A$. The limiting unavailability of another system with the same n components in parallel is $\bar{A} = \prod_{i=1}^n \bar{A}_i$. By the Key Renewal Theorem [3, Theorem 3.10], it can be shown that

$$A_i = \frac{\mu_{F_i}}{\mu_{F_i} + \mu_{G_i}} = \frac{\frac{1}{\mu_{G_i}}}{\frac{1}{\mu_{G_i}} + \frac{1}{\mu_{F_i}}} \quad (5)$$

The second expression of A_i is the ratio between the expected number of repairs per unit of repair time and the expected number of events (failures and repairs) per unit of calendar time. If the former is relatively large, it means that repairs are very short, so that it is likely to find the system in an operating state at inspection.

If we are not satisfied with the expected value of the system downtime, we should try to recover its entire distribution, and this is what we set out to do in the following sections.

SDD of a One-Component System

The SDD after a failure of a one-component system coincides trivially with G .

The SDD in an interval has an infinite series expansion, from [2, Theorem 37]:

$$\Pr(Y_u \leq y) = \sum_{n=0}^{\infty} G^{*n}(y) [F^{*n}(u - y) - F^{*(n+1)}(u - y)] \quad (6)$$

where G^{*n} and F^{*n} are the n -fold convolutions of G and F , respectively. To understand this formula, consider that $Y_u \leq y$ if and only if there are at least $u - y$ units of operating time in the interval $[0, u]$. During the first $u - y$ units of operating time one can observe any number n of failures, with probability $F^{*n}(u - y) - F^{*(n+1)}(u - y)$, which is derived from the properties of the renewal process determined by the failure times only. Then, for Y_u to be lower than or equal to y , the n repairs relative to these failures must take no more than y units of time, which happens with probability $G^{*n}(y)$, considering the renewal process of repair times alone.

Equation (6), while being elegant, is not of immediate use. As in the previous section, an asymptotic approximation is available as $u \rightarrow \infty$ [2, Theorem 38]:

$$Y_u \sim N(u\bar{A}, u\tau^2) \quad (7)$$

where

$$\tau^2 = \frac{\mu_F^2 \sigma_G^2 + \mu_G^2 \sigma_F^2}{(\mu_F + \mu_G)^3} \quad (8)$$

SDD of a Multicomponent System

SDD after a Failure of a Parallel System

In a multicomponent system there are many combinations of the states of the components that correspond to a failure state of the system. Hence the form of the SDD after a failure is not trivial. A building block for the derivation of approximate expressions for it is the SDD after a failure of a **parallel system**, which we will now briefly illustrate.

Consider a parallel system of n independent components, which fails at time t . Suppose we know that component i has caused the failure, meaning

it was the last component to fail: then the conditional SDD is the distribution of the minimum of the repair time of component i and the forward recurrence times in the failed state of the other components. The unconditional SDD is the weighted sum of the conditional SDDs. To express this formally we need some additional notation. We denote by

- Y : the downtime of the system after a failure;
- N_t : the counting process of failures of the system;
- $\Delta N_t = N_t - N_{t-}$: a quantity that is always zero except at jump times of N_t , when system failures occur;
- $\alpha_i(t)$: the forward recurrence time of component i , $i = 1, \dots, n$;
- $c_i(t)$: the probability that component i caused the system failure knowing that it has failed at time t .

The SDD after a failure for a parallel system is then

$$\begin{aligned} \Pr(Y \leq y \mid \Delta N_t = 1) \\ = \sum_{i=1}^n c_i(t) \left[1 - \bar{G}_i(y) \prod_{k \neq i} \bar{G}_{\alpha_i(k)}(y) \right] \end{aligned} \quad (9)$$

with Y indicating the downtime, and $G_{\alpha_i(k)}$ being the distribution of the forward recurrence time of component k in the failed state.

If the components are all identical, then $c_i(t) = 1/n$ for all values of i and

$$\Pr(Y \leq y \mid \Delta N_t = 1) = 1 - \bar{G}(y) [\bar{G}_{\alpha_i}(y)]^{n-1} \quad (10)$$

In the general case, it is still possible to write down a closed-form expression for the $c_i(t)$'s. Reading the proof of [2, Theorem 57], one finds that

$$c_i(t) = \frac{m_i(t) \prod_{k \neq i} \bar{A}_k(t)}{\sum_{j=1}^n m_j(t) \prod_{k \neq j} \bar{A}_k(t)} \quad (11)$$

where $m_i(t)$ is the renewal density of the modified renewal process of component i , that is, the process where the first failure time has distribution F_i and the interevent times between successive failures have distribution $G_i * F_i$. The renewal density $m(t)$ can be interpreted as the expected number of failures

per unit of time in an infinitesimal interval centered around t [2, Theorem 117]. We see that the numerator of equation (11) is in fact the expected number of system failures per unit of time caused by failures of component i in a tiny interval around time t .

The finite-time expression of equation (9) supplemented by equation (11) is not practically useful, because of the difficulty of calculating the quantities involved for specific lifetime and repair distributions. Taking the limit as $t \rightarrow \infty$, the unavailabilities tend to the limiting unavailabilities as in equation (5), $m_i(t)$ tends to the long-run frequency of failures, being $1/(\mu_{F_i} + \mu_{G_i})$ [2, Theorem 118], and $\bar{G}_{\alpha_i(k)}$ tends to

$$\bar{G}_{k,\infty}(y) = \frac{\int_y^\infty \bar{G}_k(x) dx}{\mu_{G_k}} \quad (12)$$

The asymptotic form of the SDD after a failure for a parallel system is then

$$\begin{aligned} \lim_{t \rightarrow \infty} \Pr(Y \leq y \mid \Delta N_t = 1) \\ = \sum_{i=1}^n c_i \left[1 - \bar{G}_i(y) \prod_{k \neq i} \bar{G}_{k,\infty}(y) \right] \end{aligned} \quad (13)$$

where

$$c_i = \frac{1/\mu_{G_i}}{\sum_{k=1}^n 1/\mu_{G_k}} \quad (14)$$

In the simplest possible case, that is when G_k is an exponential distribution for all k 's, the SDD takes the following form

$$\begin{aligned} \lim_{t \rightarrow \infty} \Pr(Y \leq y \mid \Delta N_t = 1) \\ = \sum_{i=1}^n c_i \left[1 - e^{-\left(\sum_{k=1}^n \frac{1}{\mu_{G_k}} \right) y} \right] \\ = 1 - e^{-\left(\sum_{k=1}^n \frac{1}{\mu_{G_k}} \right) y} \end{aligned} \quad (15)$$

that is, a mixture of exponential distributions, where each component of the mixture is the distribution of the minimum of n exponentially distributed variates. Since all the components of the mixture are the same, the SDD is simply an exponential.

SDD after a Failure of a Coherent System

We now turn to the more general case of a coherent system (see **Coherent Systems**), still with binary

4 System Downtime Distributions

components. Consider the collection of all the minimal cut sets of the system, K_j , $j = 1, \dots, J$ (see **Path Sets and Cut Sets in System Reliability Modeling**). Since the cut sets are minimal, one of them can cause system failure only if all its components fail, hence any minimal cut set can be considered as a parallel system and its asymptotic SDD after a failure is given by equation (13). In Figure 1 we show a very simple system with two minimal cut sets, $K_1 = \{C_1, C_2\}$ and $K_2 = \{C_1, C_3\}$, which will be helpful to illustrate the ideas that follow.

Now denote by $\{\Delta K_j(t) = 1\}$ the event that cut set K_j fails at time t , then we can state that

$$\{\Delta N_t = 1\} = \{\Delta K_1(t) = 1\} \cup \dots \cup \{\Delta K_J(t) = 1\} \quad (16)$$

Indeed, if a system failure occurs at time t , then at least one minimal cut set must have entered the failed state, and *vice versa*. However, the collection of events on the right-hand side is not a partition of $\{\Delta N_t = 1\}$ (in Figure 1, both $\{\Delta K_1(t) = 1\}$ and $\{\Delta K_2(t) = 1\}$ include the event that the cut set $\{C_1, C_2, C_3\}$ fails at time t) and the following inequality holds:

$$\begin{aligned} \Pr(Y \leq y \mid \Delta N_t = 1) \\ \leq \sum_{j=1}^J \Pr(Y \leq y \mid \Delta K_j(t) = 1, \Delta N_t = 1) \\ \times \Pr(\Delta K_j(t) = 1 \mid \Delta N_t = 1) \end{aligned} \quad (17)$$

To take advantage of our previous result on parallel systems, we substitute each term $\Pr(Y \leq y \mid \Delta K_j(t) = 1, \Delta N_t = 1)$ by $G_{K_j}(y)$, that is, the downtime distribution after a failure of the cut set K_j considered in isolation. When the failed cut set K_j is repaired, this does not imply that the system is repaired, whereas the converse is true, therefore, the downtime of K_j is lower than the downtime of the system, and we can state that $\Pr(Y \leq y \mid \Delta K_j(t) = 1, \Delta N_t = 1) \leq G_{K_j}(y)$. With reference to

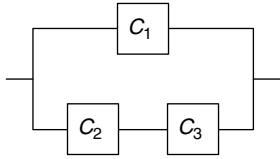


Figure 1 A three-component system

Figure 1, if C_2 and C_3 are failed and C_1 fails at time t , then both $\{\Delta K_1(t) = 1\}$ and $\{\Delta K_2(t) = 1\}$ occur simultaneously. If C_2 is repaired first, then K_1 is repaired, but the system is not.

Finally, let us examine $\Pr(\Delta K_j(t) = 1 \mid \Delta N_t = 1)$. This is the probability that K_j failed at the same time as the system, given that the latter failed at time t . As in equation (11), this quantity is given by the ratio of the expected frequency of system failures where cut set K_j also fails, divided by the expected frequency of all failure modes of the system at time t . Referring once more to Figure 1, letting $\lambda_i = (\mu_{F_i} + \mu_{G_i})^{-1}$, $i = 1, 2, 3$, we see that, as $t \rightarrow \infty$

$$\begin{aligned} \lim_{t \rightarrow \infty} \Pr(\Delta K_1(t) = 1 \mid \Delta N_t = 1) \\ = \frac{(\bar{A}_2 A_3 + \bar{A}_2 \bar{A}_3) \lambda_1 + \bar{A}_1 A_3 \lambda_2}{(A_2 \bar{A}_3 + \bar{A}_2 A_3 + \bar{A}_2 \bar{A}_3) \lambda_1 + \bar{A}_1 A_2 \lambda_3 + \bar{A}_1 A_3 \lambda_2} \\ = \frac{\bar{A}_2 \lambda_1 + \bar{A}_1 A_3 \lambda_2}{(A_2 \bar{A}_3 + \bar{A}_2) \lambda_1 + \bar{A}_1 A_2 \lambda_3 + \bar{A}_1 A_3 \lambda_2} \end{aligned} \quad (18)$$

If the components are highly available, we may approximate A_i by 1, $i = 1, 2, 3$, and equation (18) becomes

$$r_1 = \frac{\bar{A}_2 \lambda_1 + \bar{A}_1 \lambda_2}{(\bar{A}_3 + \bar{A}_2) \lambda_1 + \bar{A}_1 (\lambda_3 + \lambda_2)} \quad (19)$$

This is the asymptotic frequency of failures of the minimal cut set K_1 considered in isolation, normalized with the sum of such frequencies of minimal cut sets K_1 and K_2 , and it is much easier to calculate than equation (18). In the general case, supposing K_j has n_j components, we obtain

$$r_j = \frac{\sum_{l=1}^{n_j} \lambda_l^{(j)} \prod_{k \neq l} \bar{A}_k^{(j)}}{\sum_{j=1}^J \sum_{l=1}^{n_j} \lambda_l^{(j)} \prod_{k \neq l} \bar{A}_k^{(j)}}, \quad j = 1, \dots, J \quad (20)$$

where the necessary reindexing of components for the j th cut set is signaled by the superscript (j) . Then, the final asymptotic approximation to the SDD after a failure is

$$\Pr(Y \leq y) \simeq \sum_{j=1}^J r_j G_{K_j}(y) \quad (21)$$

where

$$G_{K_j}(y) = \sum_{i=1}^{n_j} c_i^{(j)} \left[1 - \bar{G}_i^{(j)}(y) \prod_{k \neq i} \frac{\int_y^\infty \bar{G}_k^{(j)}(x) dx}{\mu_{G_k^{(j)}}} \right] \quad (22)$$

Here $G_k^{(j)}$ denotes the repair time distribution of component k of cut set K_j , and $\mu_{G_k^{(j)}}$ is its mean.

The probabilities $c_i^{(j)}$, relative to K_j , are found from equation (14). We are not sure of the inequality sign in equation (17) anymore, because we do not know in general whether the coefficients r_j approximate by excess or by deficit.

SDD in an Interval

Let us reconsider equation (6). Suppose F is an exponential distribution with mean $\mu_F = 1/\lambda_F$. Then the SDD is

$$\Pr(Y_u \leq y) = \sum_{n=0}^{\infty} G^{*n}(y) \frac{[\lambda_F(u-y)]^n}{n!} e^{-\lambda_F(u-y)} \quad (23)$$

which says that the number of failure occurring in the operational time $u - y$ is Poisson distributed with mean λ_F . This would prompt us to consider an approximation to the SDD for a system with more than one component, where λ_F is substituted by the asymptotic frequency of system failures in operational time and G is replaced by the downtime distribution after a failure found in equations (21) and (22). This approach is justified in [2, Section 4.6], and we refer the reader to it for details, because, as in the case of a one-component system, this formula is not of immediate use.

There exists an asymptotic normal approximation that generalizes equation (7) under the assumption that the stochastic process of the joint states of the components is regenerative (see [3, Section 5.4]). In this case, let S denote the length of a regeneration cycle, and let Y_S be the downtime during the cycle, then

$$\bar{A} = \frac{E(Y_S)}{E(S)} \quad \text{and} \quad \tau^2 = \frac{\text{Var}(Y_S - \bar{A}S)}{E(S)} \quad (24)$$

The limiting unavailability $\bar{A}$ can be found starting from the structure function of the system or by considering the regeneration cycles directly. It can be

seen immediately that for a one-component system we reobtain the quantities in equation (7) with this latter approach. For **coherent systems**, when components are highly available, $\tau^2 \approx \lambda E(Y_1^2)$, where λ is the asymptotic frequency of system failures and Y_1 is the downtime after the first system failure. The asymptotic system failure rate is found as a weighted sum of the failure rates of each component. Each weight is the probability that, by failing, the component causes a system failure, for any state of the other components. The expression for λ is then

$$\lambda = \sum_{i=1}^n \frac{h(A_1, \dots, A_{i-1}, 1, A_i, \dots, A_n) - h(A_1, \dots, A_{i-1}, 0, A_i, \dots, A_n)}{\mu_{F_i} + \mu_{G_i}} \quad (25)$$

where

$$\begin{aligned} h(A_1, \dots, A_n) &= \lim_{t \rightarrow \infty} h(A_1(t), \dots, A_n(t)) \\ &= \lim_{t \rightarrow \infty} \Pr(\Phi(t) = 1) \end{aligned} \quad (26)$$

is the limiting **system availability** A . With regard to $E(Y_1^2)$, this can be further approximated via the asymptotic SDD after a failure in equations (21) and (22).

Coherent Systems with Multistate Components

So far we have considered systems of binary components. Every time a component changes its state, a new sojourn time in that state is chosen from the relevant distribution. With more than two states we must allow for a (possibly random) rule for the transitions among the states and a duration distribution for each state. This leads us to systems of independent components where each component evolves as a semi-Markov process (see **Markov Renewal Processes in Reliability Modeling**; [3, Chapter 5]). We may call these *semi-Markov systems*. Van Casteren *et al.* [5], Pievatolo and Valadè [6], and Pievatolo *et al.* [7] examined semi-Markov systems in the context of power system reliability. In [6] only the mean time between system failures is calculated, and the downtime distribution is found by simulation. In [5], the authors give analytic expressions for the duration distribution of each system state. Finally, in [7] an analytic approximation for the SDD after a failure is given, obtaining a result that is more general

than that of [5] because the system's out-of-service condition often corresponds to a set of system states, and not just to a single state. The analytic expression for the SDD resembles equations (21) and (22) closely, with two main differences: the concept of minimal cut set is extended to a set of component states where no state can be changed or no component can be removed without it losing its status of cut set; the probability that a cut set causes a failure given a system failure depends both on the transition matrix of the Markov chain and on the means of the duration distributions. The accuracy of this approximation is shown *via* a simulation study on the specific case. We refer the reader to these papers where analytic results are detailed without reference to the specific applications.

References

- [1] Rausand, M. & Høyland, A. (2004). *System Reliability Theory*, 2nd Edition, John Wiley & Sons, Hoboken.

- [2] Aven, T. & Jensen, U. (1999). *Stochastic Models in Reliability*, Springer-Verlag, New York.
- [3] Ross, S.M. (1970). *Applied Probability Models With Optimization Applications*, Holden Day, San Francisco.
- [4] Suyono & Van Der Weide, J.A.M. (2003). A method for computing total downtime distributions in repairable systems, *Journal of Applied Probability* **40**, 643–653.
- [5] van Casteren, J.F.L., Bollen, M.H.J. & Schmieg, M.E. (2000). Reliability assessment in electrical power systems: the Weibull-Markov stochastic model, *IEEE Transactions on Industry Applications* **36**, 911–915.
- [6] Pievatolo, A. & Valadè, I. (2003). UPS reliability analysis with non-exponential duration distributions, *Reliability Engineering and System Safety* **81**, 183–189.
- [7] Pievatolo, A., Tironi, E. & Valadè, I. (2004). Semi-Markov processes for power system reliability assessment with application to uninterruptible power supply, *IEEE Transactions on Power Systems* **19**, 1326–1333.

Related Articles

Availability Models; System Availability.

ANTONIO PIEVATOLO

Intensity Functions for Nonhomogeneous Poisson Processes

Introduction

The nonhomogeneous **Poisson process** (NHPP) has been used and developed for reliability development modeling since Duane [1] identified that plotting the cumulative number of a system's failures against cumulative time on log-log paper produced a straight-line relationship. Crow [2] developed this work for the US MIL HDBK-189 [3] in the late 1970s showing that the Duane model was an NHPP with a power law rate of occurrence of failures (ROCOF). Rigdon [4] subsequently showed that this was not the case. Parzen [5] presented the background theory to NHPPs and this has been developed further by Willis [6], Cox and Isham [7], and Thompson [8]. Cox and Lewis [9] stated a log-linear ROCOF and derived its maximum-likelihood estimates. Earlier, Rosner [10] had been looking into a software-applicable ROCOF and developed the IBM model, but only the power law ROCOF produces simple to calculate parameter estimates.

Definitions

An NHPP satisfies the following conditions

1. Since events start at $t = 0$, $N(0) = 0$.
2. The process $\{N(t), t \geq 0\}$ has independent increments. The time t is the cumulative time from zero.
3. For any $t > 0$, $0 < \text{Prob}(N(t) > 0) < 1$.
4. The number of failures in any interval of length t is distributed as a Poisson distribution with parameter $M(t) = E(N(t))$, also known as the *mean value function* and the expected cumulative number of failures in time t [11].
5. Only one event can occur at any one time.
6. $\lim_{h \rightarrow 0} \frac{P(N(t+h) - N(t) \geq 1)}{h} = \lambda(t)$.

This term $\lambda(t)$ is known as the *intensity*, *peril rate*, *ROCOF*, or *instantaneous failure rate* (not to be confused with hazard rate), which at the end

of a system development program is the failure rate which should remain constant through production and service use, if no more systematic design or manufacturing failures are encountered. Thompson [8] shows that if there are no simultaneous arrivals then $M(t) = \int_0^t \lambda(x) dx$, thus $\lambda(t) = \frac{dM(t)}{dt}$. So, as far as the models are concerned $E(N(0)) = 0$ and $\lim_{t \rightarrow \infty} E(N(t)) = \int_0^\infty \lambda(u) du \leq \infty$, where the less than sign usually applies to **software reliability** NHPPs. Also, another term used in calculations and for parameter **estimation** purposes is the *cumulative failure rate*, $\gamma(t_i) = \frac{M(t_i)}{t_i}$. The likelihood for a time-terminated process (at t_0 after the n th event is $L = (\prod_{i=1}^n \lambda(t_i)) \exp(-\int_0^{t_0} \lambda(x) dx)$. Discussion of sufficiency (of the log-linear model) appears in [9].

Parzen [5] and Thompson [8] discuss three distinct concepts of time between arrivals and results from these references are presented to highlight the difference between data, which is independently and identically distributed (i.i.d) and the NHPP. The distribution function of the first waiting time to failure is $F(w_1) = 1 - \exp(-M(w_1))$; however, the distribution function of w_r (the r th waiting time to failure) given the first $r - 1$ events is $F(w_r/T_i = t_1, \dots, T_{r-1} = t_{r-1}) = 1 - \exp(-M(t_r) + M(t_{r-1}))$. So, the first failure time is distributed with a **probability density function** (pdf) that can be derived from the intensity function but subsequent times will have different pdfs, however, all hazard rates are applicable intensities within the NHPP. The probability of n events in time t is $P(N(t) = n) = \frac{M(t)^n e^{-M(t)}}{n!}$ [12].

There is an inconsistency in naming intensities. For instance, the Duane model is also known as the *AMSAA model*, the *power law process*, and the *Weibull process*. If we decide to call the process by the hazard function of the first waiting time, then the Weibull process seems to be the most appropriate; however, the software Weibull model appears later in this contribution, and so we will refer to it as the *power law process*. Similarly, the log-linear model should be entitled the *Gompertz process* after Gompertz [13] and the IBM model or the Smiths Industries Ltd (SIL) model (see subsequently) should be known as the *Makeham process* after Makeham [14], although Gompertz and Makeham were actually working in actuarial science.

We shall address the following intensity functions in the areas where they have typically received the first usage.

The Simplest Models

The homogeneous Poisson process is given by $\lambda(t_i) = h(w_i) = b$ (ROCOF and hazard rate are a constant) and the expected number of failures is a linear function of cumulative time $M(t_i) = bt_i$ where $\lambda(t_i) = b$ (the exponential distribution). The cumulative failure rate also gives $\gamma(t_i) = b$. A check for data exhibiting **reliability growth** or decay is done by plotting number of failures against cumulative time and checking for curvature. Extensions to these are the linear ROCOF model: $\lambda(t_i) = at_i + b$ for which $M(t_i) = \frac{a}{2}t_i^2 + bt_i$ and $\gamma(t_i) = \frac{a}{2}t_i + b$ and the quadratic ROCOF model $\lambda(t_i) = 3at_i^2 + 2bt_i + c$ with $M(t_i) = at_i^3 + bt_i^2 + ct_i$ and $\gamma(t_i) = at_i^2 + bt_i + c$.

The Two Best-Known Models

Duane [1] and Crow [2] discuss the model $\lambda(t_i) = abt_i^{b-1} = b \frac{M(t_i)}{t_i} = b\gamma(t_i)$ so the ROCOF is the cumulative failure rate scaled by the shape parameter. This is straightforward to draw onto log-log paper commonly known as a *Duane plot*. The mean value $M(t_i) = at_i^b$ can also be plotted on log-log paper so that parameters a and b can be estimated. The approach to analysis of a Duane plot should follow procedures similar to those for the well-known Weibull **probability plot**, i.e., evidence of two or more slopes, curvature, outliers, and so on. Further work on this model appears in [4, 15–19].

The log-linear model [9] is given by $\lambda(t_i) = e^{a+bt_i}$ $M(t_i) = \frac{e^a}{b}(e^{bt_i} - 1)$. Writing this as $M(t_i) = \frac{a}{b}(e^{bt_i} - 1)$, we get the ROCOF as a linear function of the expected number of failures, $\lambda(t_i) = a + bM(t_i)$. Reasonably straightforward programming provides maximum-likelihood estimates of the parameters. Nonlinear least-squares estimates can be calculated easily for this model.

We can consider the various hardware ROCOFs as having the same functional form as hazard functions and determine their pdfs. Addressing the hazard rate for the log-linear model, we get the pdf $f(w_i) = e^{a+bw_i} \exp(-\frac{e^a}{b}(e^{bw_i} - 1))$, an extreme-value (or Gumbel [20]) distribution [21, p. 139] with a Gompertz hazard rate (see below). Thus, the two best-known models give rise to the type I and type III extreme-value distributions. For completeness, a type II intensity can be

based on the arctangent survival distribution [22], where $\lambda(t_i) = \frac{a}{(\tan^{-1}(a(b-t_i)) + \frac{\pi}{2})(1+(a(t_i-b))^2)}$ and $M(t_i) = \ln\left(\frac{\tan^{-1}ab}{\tan^{-1}a(b-t_i)}\right)$, $0 < t_i < b$. Another trigonometric intensity $\lambda(t_i) = a(1 + b \sin ct_i)$ was developed by Willis [6], who showed that the expected number of events between t and $t + \tau$ is $a\left(\tau + \frac{2b}{c} \sin c\left(t + \frac{\tau}{2}\right) \sin \frac{c\tau}{2}\right)$.

Extensions to the log-linear model are in [23] $\lambda(t_i) = e^{a+bt_i+ct_i^2}$ and [24] $\lambda(t_i) = e^{\sum_{j=0}^k at_i^j}$. Simpler extensions of this latter one are given by $\lambda(t_i) = ae^{bt_i}(1 + abt_i)$, $M(t_i) = at_i e^{bt_i}$, $\gamma(t_i) = ae^{bt_i}$ and the simple exponential model $M(t_i) = at_i e^{-bt_i}$, $\gamma(t_i) = ae^{-bt_i}$ and least-squares estimates can be obtained from $\ln(\gamma(t_i)) = \ln a + bt_i$, $\ln(\gamma(t_i)) = \ln a - bt_i$, respectively.

Extensions to the power law process include the square root model [25], which is the power law process with $b = 0.5$: $\lambda(t_i) = \frac{a}{\sqrt{t_i}}$, $M(t_i) = 2a\sqrt{t_i}$ and the modified Duane model [26]: $\lambda(t_i) = abc^b(c + t_i)^{-b-1}M(t_i) = a\left(1 - \left(\frac{c}{c+t_i}\right)^b\right)$. An extension of this software reliability model is the hardware equivalent: $\lambda(t_i) = \frac{b(M(t_i)+a)}{c+t_i}$, $M(t_i) = a\left(\left(\frac{c+t_i}{c}\right)^b - 1\right)$. This intensity is of a similar functional form to the hazard rate developed in [27]: $h(w_i) = \frac{ab}{(c-w_i)^{b-1}}$, $0 < w_i < c$.

Reference [16] contains a model $\lambda(t_i) = act_i^{a-1}e^{bt_i}$, which incorporates both the best-known ones, but requires some programming to calculate the parameter estimates.

NHPPs Derived from Actuarial Studies

Historically, much work has been carried out on hazard rates in the actuarial field and some NHPPs are given below based on the best-known models in this field.

Gompertz [13] initially defined a hazard rate $h(w_i) = ab^{w_i}$, which we can write as the ROCOF $\lambda(t_i) = ab^{t_i}$ giving $M(t_i) = \frac{a}{\ln b}(b^{t_i} - 1)$. Makeham [14] extended this to $h(w_i) = c + ab^{w_i}$, which provides the ROCOF $\lambda(t_i) = c + ab^{t_i}$ with $M(t_i) = ct_i + \frac{a}{\ln b}(b^{t_i} - 1)$ described by Ascher [28].

Now, denote $T_1 = \sum t_i$, $T_2 = \sum t_i^2$, $B_1 = \sum b^{t_i}$, $B_2 = \sum b^{2t_i}$, $B_3 = \sum t_i b^{t_i}$, $B_4 = \sum t_i b^{2t_i}$, $B_5 = \sum t_i^2 b^{t_i}$, $M_1 = \sum m_i$, $M_2 = \sum m_i t_i$, $M_3 = \sum m_i t_i^2$, $M_4 = \sum m_i b^{t_i}$, and $M_5 = \sum m_i t_i b^{t_i}$.

We can obtain the Ascher model least-squares estimates for b by numerically solving

$$\frac{M_5(B_3 - T_1) - M_2(B_4 - B_3)}{B_5(B_3 - T_1) - T_2(B_4 - B_3)} = \frac{(M_4 - M_1)(B_4 - B_3) - M_5(B_2 - 2B_1 + n)}{(B_3 - T_1)(B_4 - B_3) - B_5(B_2 - 2B_1 + n)},$$

$$\text{then solve } c = \frac{M_5(B_3 - T_1) - M_2(B_4 - B_3)}{B_5(B_3 - T_1) - T_2(B_4 - B_3)}$$

$$\text{and finally } a = \frac{M_5 - cB_5}{B_4 - B_3}.$$

Rosner [10] describes the IBM model $\lambda(t_i) = c + abe^{bt_i}$ with $M(t_i) = ct_i + a(e^{bt_i} - 1)$ latter renamed the *SIL model* [29]. If we write $b^{t_i} = e^{t_i \ln b}$ in the Makeham model, we get the IBM model.

Other actuarial hazard rates discussed in Benjamin and Pollard [30] and converted to ROCOFs include $\lambda(t_i) = a + bt_i + cd^{t_i}$, (an addition of a linear model with time with a Gompertz process); $\lambda(t_i) = e^{a+bt_i} + e^{c+dt_i}$, (an addition of two log-linear ROCOFs);

Perk's process

$$\lambda(t_i) = \frac{a+bc^{t_i}}{1+dc^{t_i}} \text{ with } M(t_i) = \alpha + \beta t_i + \delta \ln(1 + \varepsilon c^{t_i});$$

Perk's process (extended)

$$\lambda(t_i) = \frac{a+bc^{t_i}}{Kc^{-t_i} + 1 + dc^{t_i}} \text{ with } M(t_i) = \alpha \ln \frac{c^{-t_i} + \beta}{1 + \beta} + \delta \ln \frac{c^{-t_i} + \phi}{1 + \phi}; \text{ and}$$

Barnett's process

$$\lambda(t_i) = \frac{-H + Bc^{t_i} \ln c}{1 + A - Ht_i + Bc^{t_i}} \text{ (original notation used)}$$

with $M(t_i) = \ln \left(\frac{1 + A - Ht_i + Bc^{t_i}}{1 + A + B} \right).$

The **repairable systems** bathtub curve [31] can be modeled by the following intensities which are based on the nonrepairable bathtub functions mentioned in [32] and [33], respectively: $\lambda(t_i) = a + bt + \frac{c}{t_i + d}$, $M(t_i) = at_i + bt_i^2 + c \ln \left(\frac{t_i + d}{d} \right)$, and $\lambda(t_i) = at + \frac{c}{1 + bt_i}$ (bathtub shaped for $0 < a < bc$) and $\lambda(t_i) = ab(at_i)^{b-1} \exp(at_i^b)$ (see the software Weibull model below).

The NHPPs Associated with Well-Known Probability Density Functions

The list below gives the ROCOFs for some well-known distributions:

The lognormal ROCOF $\lambda(t_i) \propto \frac{\phi(\ln(t_i))}{1 - \Phi(\ln(t_i))}$, where $\Phi(\cdot)$ is the standardized normal integral.

The gamma process [34, 35], $\lambda(t) = \frac{t^{\psi-1} e^{-\gamma t}}{\int_t^\infty u^{\psi-1} e^{-\gamma u} du}$

The Erlang process (based on the hazard rate in [36]) $\lambda(t_i) = \frac{(t_i/a)^{b-1}}{a((b-1)!)} \left(\sum_{i=0}^{c-1} \frac{(t_i/a)^i}{i!} \right)$

The generalized Wald process [37]

$$\lambda(t_i) = a + \frac{b}{(1 + ct_i)(1 + \ln(1 + ct_i))},$$

$$M(t_i) = at_i + b \ln(1 + \ln(1 + ct_i))$$

The Burr process based on the Burr distribution [38] is

$$\lambda(t_i) = \frac{abt_i^{b-1}}{c + t_i^b}, M(t_i) = a \ln \left(\frac{c + t_i^b}{c} \right)$$

The log-logistic process [39] $\lambda(t_i) = \frac{abt_i^{b-1}}{1 + at_i^b}$, $M(t_i) = \ln(1 + at_i^b)$ is the Burr process with $c = 1$. We may plot $\ln(e^M - 1)$ versus $\ln t$ to obtain least-squares estimates of the parameters.

The Pareto process $\lambda(t_i) = ab(1 + bt_i)^{-1}$, $M(t_i) = a \ln(1 + bt_i)$ is the same as $\lambda(t_i) = \frac{ab}{1 + at_i}$, $M(t_i) = b \ln(1 + at_i)$ [39] and is similar to the logarithmic model by Musa and Okumoto [40] $\lambda(t_i) = a(1 + abt_i)^{-1}$, $M(t_i) = \frac{1}{b} \ln(1 + abt_i)$.

The inverse Exponential process intensity [37] is $\lambda(t_i) = \frac{abe^{-b/t_i}}{t_i^2}$, $M(t_i) = ae^{-b/t_i}$. We plot $\ln M(t_i) = \ln a - \frac{b}{t_i}$ for parameter estimates.

The inverse Gaussian process intensity is

$$\lambda(t_i) = \frac{\left(\frac{a}{b} \right)^{1/2} \exp \left(-\frac{ab(t_i - \frac{1}{b})^2}{2t_i} \right)}{\left[1 - \Phi \left[\left(\frac{a}{bt_i} \right)^{1/2} (-1 + bt_i) \right] \right] \left[-e^{2a} \Phi \left[-\left(\frac{a}{bt_i} \right)^{1/2} (1 + bt_i) \right] \right]}$$

where $\Phi(\cdot)$ is the standardized normal integral.

Some Logarithmic Models

The simple logarithmic process [39] has intensity $\lambda(t_i) = \frac{a}{t_i}$ with $M(t_i) = a \ln(bt_i) = a \ln b + a \ln(t_i)$, $t \geq 1$. We can plot $M(t_i)$ against t_i on log-linear paper to obtain estimates of a, b . Weitz and Fraser

[41] also considered the hazard rate $h(w_i) = \frac{a}{1-aw_i}$, which if also considered as the intensity $\lambda(t_i) = \frac{a}{1-at_i}$ produces a logarithmic mean value function of $M(t_i) = -a \ln(1 - at_i)$.

Another model based on the hazard rate [42] is given by $\lambda(t_i) = at_i(b^2 + t_i^2)^{-1}$ with $M(t_i) = \frac{a}{2} \ln\left(1 + \frac{t_i^2}{b^2}\right)$.

A two-parameter logarithmic intensity function was developed as a hazard rate by Dhillon [43] as $\lambda(t_i) = \frac{(a+1)(\ln(1+bt_i))^a}{1+bt_i}$, which has mean value function $M(t_i) = \frac{1}{b}(\ln(1+bt_i))^{a+1}$.

Five variations of this model include

1. $\lambda(t_i) = a \ln(t_i) + a + b = \frac{M(t_i)}{t_i} + a$,
 $M(t_i) = at_i \ln(t_i) + bt_i$, $\gamma(t_i) = \frac{M(t_i)}{t_i} =$
 $a \ln(t_i) + b$ since as $t \rightarrow 0$,
 $\lim(t \ln t) = \lim\left(\frac{\ln t}{t^{-1}}\right) = \lim\left(\frac{1/t}{1/t^2}\right) =$
 $\lim(t) = 0$ so $M(0) = 0$; however, for calculation of λ we require $t \geq 1$.

2. A reduced form is

$$\lambda(t_i) = a \ln(t_i) + a = \frac{M(t_i)}{t_i} + a = \gamma(t_i) + a,$$

$$M(t_i) = at_i \ln(t_i), \gamma(t_i) = a \ln(t_i).$$

3. The logarithmic NHPP [44] intensity is

$$\lambda(t_i) = a \ln(1 + bt_i) + c \text{ with}$$

$$M(t_i) = \frac{a}{b}(1 + bt_i) \ln(1 + bt_i) + (c - a)t_i$$

4. Pievatolo *et al.* [45] used

$$\lambda(t_i) = \frac{a \ln(1 + bt_i)}{1 + bt_i}, M(t_i) = \frac{a}{2b}(\ln(1 + bt_i))^2$$

in the modeling of underground train failures.

5. Xie and Zhao [39] came up with

$$\lambda(t_i) = \frac{ab \ln(t_i)^{b-1}}{t_i}, M(t_i) = a \ln(t_i)^b, t \geq 1 \text{ and}$$

$$\lambda(t_i) = \frac{ab \ln(t_i + 1)^{b-1}}{t_i + 1}, M(t_i) = a(\ln(t_i + 1))^b,$$

$$t \geq 0.$$

Estimates of the parameters can be obtained from plotting $\ln M(t_i)$ against $\ln(\ln(t_i + 1))$.

Software Models and Creating a Nonrepairable Hazard Rate from them

The well-known software NHPPs have the distinction of a ROCOF, which does not tend to infinity as operational time does. So the ROCOF of the first failure time is not a hazard rate. Some of the well-known ones are listed below. The parameter a is the expected number of initial defects as $t \rightarrow \infty$.

The Littlewood model [46] $\lambda(t_i) = abc(1 + bt_i)^{-c-1}$, $M(t_i) = a(1 - (1 + bt_i)^{-c})$ is similar to the “initial defects” model (generalized version) [47] $\lambda(t_i) = \frac{b}{a} \left(\frac{a}{a+t_i}\right)^{b+1}$, $M(t_i) = a \left(1 - \left(\frac{a}{a+t_i}\right)^b\right)$.

This is Littlewood’s modified Duane model with $c = a$ [26].

In the Musa [48] $\lambda(t_i) = ae^{-bt_i}$, $M(t_i) = \frac{a}{b}(1 - e^{-bt_i})$ and Goel–Okumoto [49] $\lambda(t_i) = abe^{-bt_i}$, $M(t_i) = a(1 - e^{-bt_i})$ models, the number of faults would have a Poisson distribution with mean a , the n th interval having a pdf of $f_n(t) = (a - n + 1)be^{-(a-n+1)bt}$, ($1 \leq n \leq a$). Thompson [8] showed that the probability that the number of errors remaining equal to some value x given that y errors have been found would be a Poisson distribution with mean $a - M(t)$, and that the expectation of the errors remaining, given y errors have been found is $a - M(t)$, which is independent of y .

The logarithmic Poisson model described in detail in Musa *et al.* [50] is $\lambda(M(t_i)) = ae^{-bM(t_i)}$, where $M(t_i) = \frac{1}{b} \ln(abt_i + 1)$ with the t metric being execution time.

The S-shaped inflection [51] $\lambda(t_i) = \frac{ab(c+1)e^{-bt_i}}{(1+ce^{-bt_i})^2}$, $M(t_i) = \frac{a(1-e^{-bt_i})}{(1+ce^{-bt_i})}$. The S-shaped model as proposed by Yamada *et al.* [52] $\lambda(t_i) = ab^2t_i e^{-bt_i}$, $M(t_i) = a(1 - (1 + bt_i)e^{-bt_i})$, the bounded intensity model [53] $\lambda(t_i) = a(1 - (1 + at_i)^{-1}) = \frac{a^2t_i}{1+at_i}$, $M(t_i) = at_i - \ln(1 + at_i)$ and the log-logistic reliability growth model [54] $\lambda(t_i) = ab \frac{(ct_i)^{b-1}}{(1+(ct_i)^b)^2}$, $M(t_i) = a \frac{(ct_i)^b}{1+(ct_i)^b}$ offer different shapes to the intensity function. The latter model is a special case of the model $\lambda(t_i) = ah(t_i)$.

Other models discussed in [26] are characterized by the intensity $\lambda(t_i) = af(x)$ where a is the mean of the Poisson random variable N and X_i are i.i.d. with pdf $f(x)$, e.g., $\lambda(t_i) = abc(c + t_i)^{-b-1}$, the modified inverse power law (see [46]). The software Weibull model [40], [55] $\lambda(t_i) = abct_i^{c-1}e^{-bt_i^c}$, $M(t_i) = a(1 - e^{-bt_i^c})$ also comes into

this category. A classification of software reliability models is given in [50].

It is possible to make pdfs from these models by multiplying by functions of time, which tend to infinity as operational time tends to infinity. As an example, if we take the software Weibull model (not to be confused with the Weibull process/distribution) and multiply the mean value function by $e^{bt_i^c}$, we obtain $M(t_i) = a(e^{bt_i^c} - 1)$, which then provides us with a pdf of $f(w_i) = abcw_i^{c-1}e^{-bw_i^c}e^{-a(e^{bw_i^c}-1)}$ which, to avoid confusion, should be known as the *Musa–Okumoto distribution*. Dhillon [43] and Smith and Bain [56] discussed exponential power distributions, which were characterized by the hazard rate (intensity) $\lambda(t_i) = ab(at_i)^{b-1} \exp at_i^b$ and it can be seen that the software Weibull model in [40] is of a similar form.

Some Other Models

Some other models are listed below which require some development and application:

Ascher and Feingold [31] discuss three differential equation models (all NHPPs can be written as differential equations in $\frac{dM}{dt}$) and there is scope for more development here.

Exponential single-term power series model [57]: $\lambda(t_i) = a \left(\frac{(1 + bt_i)}{(1 - ce^{-bt_i})} - \frac{bt_i}{(1 - ce^{-bt_i})^2} \right)$, $M(t_i) = at_i(1 - ce^{-bt_i})^{-1}$, $\gamma(t_i) = a(1 - ce^{-bt_i})^{-1}$ [21]: $\lambda(t_i) = \frac{t_i(bt_i - 2a)}{(bt_i - a)^2}$, $t_i \geq \frac{2a}{b}$, $M(t_i) = \frac{t_i^2}{bt_i - a}$, $t_i \geq \frac{a}{b}$, $\gamma(t_i) = \frac{t_i}{bt_i - a}$ [58]: $\frac{1}{\gamma(t_i)} = ae^{-b/t_i}$. This model requires b to be negative for $M(0) = 0$. If we rewrite the model as $M(t_i) = at_i e^{-b/t_i}$, then $\lambda(t_i) = ae^{-b/t_i} \left(1 + \frac{b}{t_i} \right)$ and $\gamma(t_i) = ae^{-b/t_i}$. Parameter estimates can be obtained by plotting $\ln \gamma(t_i)$ against $\frac{1}{t_i}$.

Other differential equation models include $\lambda(t_i) = \frac{ae^{ct_i}}{1 + \frac{ab}{c}(e^{ct_i} - 1)}$, $M(t_i) = \frac{1}{b} \ln(1 + \frac{ab}{c}(e^{ct_i} - 1))$

$\lambda(t_i) = at_i M(t_i) + b$ as $M(t_i) = be^{\frac{at_i^2}{2}} \sqrt{\frac{\pi}{2a}} \int_0^{\sqrt{\frac{a}{2}t_i}} e^{-x^2} dx$ shown by [59].

Reference [60] lists a table of other differential equation forms. Other truncated models are discussed in [37].

Extensions

In many applications, where work has been carried out on the hazard rate, there is usually a reference to work on the intensity. A nonparametric estimate of intensity is given by [61] and other intensity estimates appear in [62–64]. Proportional intensity modeling (PIM) references are [65–67]. Aalen [68] discusses multiplicative intensities in depth. Bayesian inference references include [60, 69–72]. Tests and **confidence intervals** for parameters are covered in [31], and [73]. Martingales and intensities are discussed in [74, 75]. A review which includes topics for further study is provided by Lindqvist [76].

References

- [1] Duane, J.T. (1964). Learning curve approach to reliability monitoring, *IEEE Transactions on Aerospace* **2**, 563–566.
- [2] Crow, L. (1975). On tracking reliability growth. Proceedings of the twentieth conference on the design of experiments, ARO Report 75-2, ARO, pp. 741–754.
- [3] Department of Defence Washington, OC 20301. (1981). US MIL-HDBK-189. Military standard for reliability growth plans, RADC. See http://www.barringer1.com/mil_files/MIL-HDBK-189.pdf.
- [4] Rigdon, S.E. (2002). Properties of the Duane plot for repairable systems, *Quality and Reliability Engineering International* **18**(part 1), 1–4.
- [5] Parzen, E. (1962). *Stochastic Processes*, Holden Day, San Francisco.
- [6] Willis, D.M. (1964). The statistics of a particular non-homogeneous poisson process. *Biometrika* **51**(3 and 4), 399–404.
- [7] Cox, D.R. & Isham, V. (1980). *Point Processes*, Chapman & Hall, Cambridge.
- [8] Thompson, W.A. (1988). *Point Process Models with Applications to Safety and Reliability*, Chapman & Hall, New York.
- [9] Cox, D.R. & Lewis, P.A.W. (1966). *The Statistical Analysis of Series of Events*, Methuen, London.
- [10] Rosner, N. (1961). System analysis-non-linear estimation techniques, *Proceedings of the Seventh National Symposium. on Reliability and Quality Control Institute of Radio Engineers*, IEEE.
- [11] Fry, T.C. (1965). *Probability and its Engineering Uses*, 2nd Edition, Van Nostrand, Princeton.

- [12] Cinlar, E. (1975). *Introduction to Stochastic Processes*, Prentice-Hall, Englewood.
- [13] Gompertz, B. (1825). On the nature of the function of the law of human mortality and on a new mode of determining the value of life contingencies, *Philosophical Transactions of the Royal Society* **115**, 513–585.
- [14] Makeham, W.M. (1860). On the law of mortality, *Journal of the Institute of Actuaries* **13**, 325–358.
- [15] Bain, L.J. (1978). *Statistical Analysis of Reliability and Life-testing Models*, Marcel Dekker, New York.
- [16] Lee, L. (1980). Comparing rates of several independent Weibull processes, *Technometrics* **22**(3), 427–430.
- [17] Bain, L.J. & Engelhardt, M. (1982). Sequential probability ratio tests for the shape parameter of a nonhomogeneous poisson process, *IEEE Transactions on Reliability* **R-31**(1), 79–83.
- [18] Jang, J.S. & Bai, D.S. (1987). Efficient sequential estimation in a nonhomogeneous poisson process, *IEEE Transactions on Reliability* **R-36**(2), 235–258.
- [19] Engelhardt, M., Guffey, J.M. & Wright, F.T. (1990). Tests for positive jumps in the intensity of a poisson process: a power study. *IEEE Transactions on Reliability* **39**(3), 356–360.
- [20] Gumbel, E.J. (1958). *Statistics of Extremes*, Columbia University Press, New York.
- [21] Lloyd, D.K. & Lipow, M. (1977). *Reliability: Management, Methods, and Mathematics*, 2nd Edition Miramar, Redondo Beach, CA.
- [22] Glen, A.G. & Leemis, L.M. (1997). The arctangent survival distribution, *Journal of Quality Technology* **29**(2), 205–210.
- [23] Lewis, P.A. (1972). Recent results in the statistical analysis of univariate point processes, in *Stochastic Point Processes*, P.A. Lewis, ed, Wiley-Interscience, New York, pp. 1–54.
- [24] Maclean, C.J. (1974). Estimation and testing of an exponential polynomial rate function with the nonstationary poisson process, *Biometrika* **61**, 81–85.
- [25] Kremer, W. (1983). Birth-death and bug counting, *IEEE Transactions on Reliability* **R-32**(1), 37–47.
- [26] Littlewood, B. (1984). Rationale for a modified Duane model, *IEEE Transactions on Reliability* **R-33**(2), 157–159.
- [27] Zelterman, D. (1992). A statistical distribution with an unbounded hazard function and its application to a theory from demography, *Biometrics* **48**, 807–818.
- [28] Ascher, H.E. (1968). Evaluation of repairable system reliability using the “bad-as-old” concept, *IEEE Transactions on Reliability* **R-17**, 103–110.
- [29] Lloyd, D.A., Staley, J.E. & Sutcliffe, P.S. (1977). A theoretical study of methods for analyzing reliability growth, MOD (PE) Report No. GR/77/06 MOD(PE), UK.
- [30] Benjamin, B. & Pollard, J.H. (1986). *The Analysis of Mortality and Other Actuarial Statistics*, Heinemann, London.
- [31] Ascher, H.E. & Feingold, H. (1984). *Repairable Systems Reliability Modeling, Inference, Misconceptions and their Causes*, Marcel Dekker.
- [32] Gaver, D.P. & Acar, M. (1979). Analytical hazard representations for use in reliability, mortality and simulation studies, *Communications in Statistics-Simulation and Computation* **8**, 91–111.
- [33] Hjorth, U. (1980). A reliability distribution with increasing, decreasing, constant and bath-tub shaped failure rate, Report LITH-MAT-R-77-14, Matematiska Institutionen, Tekniska Högskolan I Linköping, Linköping.
- [34] Barlow, R.E. & Davis, B. (1977). *Analysis of Time between Failures for Repairable Components*, Operations Research Center, University of California, Berkeley. ORC 77-20, Also appeared in Proc. Of SIAM Conference on Nuclear Systems Reliability and Risk Assessment, 1978.
- [35] Barlow, R.E. & Davis, B. (1979). Total time on test plots. Proceedings of the twenty-fourth conference. On the design of experiments in army research development and testing, ARO Report 79-2, ARO, pp. 361–380.
- [36] Hastings, N.A.J. & Peacock, J.B. (1975). *Statistical Distributions*, Butterworth, Bath.
- [37] McCollin, C. (2003). The Third International Symposium in Business and Industrial Statistics (ISBIS), *A Review of the Rate of Occurrence of Failures of the Non-Homogeneous Poisson Process Joint ENBIS and ISI conference*, Barcelona, August.
- [38] Burr, I.W. (1942). Cumulative frequency functions, *Annals of Mathematical Statistics* **13**, 215–232.
- [39] Xie, M. & Zhao, M. (1993). On some reliability growth models with graphical interpretations, *Microelectronics and Reliability* **33**(2), 149–167.
- [40] Musa, J.D. & Okumoto, K. (1984). A logarithmic Poisson execution time model for software reliability measurement, *Proceedings of the 7th International Conference on Software Engineering*, IEEE Computer Society Press, Washington, DC, pp. 230–238.
- [41] Weitz, J.S. & Fraser, H.B. (2001). Explaining mortality rate plateaus, *Proceedings of the National Academy of Sciences of the United States of America* **98**(26), 15383–15386.
- [42] Greenwich, M. (1992). A uni-modal hazard rate function and its failure distribution, *Statistische Hefte* **33**, 187–202.
- [43] Dhillon, B.S. (1981). Life distributions, *IEEE Transactions on Reliability* **30**, 457–459.
- [44] Cavallo, D. & Ruggeri, F. (2001). Bayesian models for failures in a gas network, Proceedings of ESREL2001, Torino, pp. 1963–1970.
- [45] Pievatolo, A., Argiento, R. & Ruggeri, F. (2002). Bayesian analysis and prediction of failures in underground trains, in *Proceedings of the ENBIS Second Annual Conference*, Rimini.
- [46] Littlewood, B. (1981). Stochastic reliability growth: a model for fault-removal in computer programs and hardware designs, *IEEE Transactions on Reliability* **R-30**, 313–320.

- [47] Cozzolino, J.M. (1968). Probabilistic models of decreasing failure rate processes, *Naval Research Logistics Quarterly* **15**, 361–374.
- [48] Musa, J.D. (1975). A theory of software reliability and its application, *IEEE Transactions on Software Engineering* **SE-1**(3), 312–327.
- [49] Goel, A.L. & Okumoto, K. (1979). Time-dependent error detection rate model for software reliability and other performance measures, *IEEE Transactions on Reliability* **R-28**(3), 206–211.
- [50] Musa, J.D., Iannino, A. & Okumoto, K. (1987). *Software Reliability Measurement, Prediction, Application*, McGraw-Hill, New York.
- [51] Ohba, M. (1984). Software reliability analysis models, *IBM Journal of Research and Development* **28**(4), 428–443.
- [52] Yamada, S., Ohba, M. & Osaki, S. (1983). Reliability growth models for hardware and software systems based on nonhomogeneous Poisson processes, *Microelectron and Reliability* **23**, 91–112.
- [53] Hartler, G. (1989). The nonhomogeneous Poisson process – a model for the reliability of complex repairable systems, *Microelectron and Reliability* **29**(3), 381–386.
- [54] Gokhale, S.S. & Trivedi, K.S. (1998). Log-logistic software reliability growth model, *Third IEEE International Symposium on High-Assurance Systems Engineering*, IEEE, Washington DC, ISBN: 0-8186-9221-9 pp. 34–41.
- [55] Lyu M.R., ed. (1996). *Handbook of Software Reliability Engineering*, McGraw-Hill, New York.
- [56] Smith, R.M. & Bain, L.J. (1975). An exponential power life-testing distribution, *Communications in Statistics* **A4**, 469–481.
- [57] Perkowski, N. & Hartvigsen, D.E. (1962). Derivations and discussions of the mathematical properties of various candidate growth functions, Report No. CRA 62-8, Aerojet-General Corporation, Azusa.
- [58] Aroef, M. (1957). Study of learning curves of industrial manual Operations, Unpublished Master's Thesis. Cornell University, Ithaca.
- [59] Sackfield, A. (2003). *Personal Communication*, Nottingham Trent University.
- [60] Ruggeri, F. (2005). *On the Reliability of Repairable Systems: Methods and Applications. International Workshop/Conference on Bayesian Statistics and its Applications*, Banaras Hindu University, Varanasi, January pp. 6–8.
- [61] Crowder, M.J., Kimber, A.C., Smith, R.L. & Sweeting, T.J. (1991). *Statistical Analysis of Reliability Data*, Chapman & Hall, London.
- [62] Leemis, L.M. (1991). Nonparametric estimation of the cumulative intensity function for a nonhomogeneous Poisson process, *Management Science* **37**(7), 886–900.
- [63] Arkin, B.L. & Leemis, L.M. (2000). Nonparametric estimation of the cumulative intensity function for a nonhomogeneous Poisson process from overlapping realizations, *Management Science* **46**(7), 989–998.
- [64] Leemis, L.M. (2004). Nonparametric estimation and variate generation for a nonhomogeneous Poisson process from event count data, *IIE Transactions* **36**(12), 1155–1160.
- [65] Lawless, J.F. (1987). Regression methods for Poisson process data, *Journal of the American Statistical Association* **82**(399), 808–814.
- [66] Wightman, D.W., Bendell, A. & McCollin, C. (1994). *Comparison of Proportional Hazards Modelling, Proportional Intensity Modelling and Additive Hazards Modelling in a Software Reliability Context*, Second International Applied Statistics in Industry Conference (Continuing Continuous Improvement) ACG Press Wichita Kansas, pp. 176–185.
- [67] Landers, T.L. & Soroudi, H.E. (1991). Robustness of a semi-parametric proportional intensity model, *IEEE Transactions on Reliability* **40**(2), 161–164.
- [68] Aalen, O.O. (1989). A linear regression model for the analysis of life times, *Statistics in Medicine* **8**, 907–925.
- [69] Guida, M., Calabria, R. & Pulcini, G. (1989). Bayes inference for a non-homogeneous Poisson process with power intensity law, *IEEE Transactions on Reliability* **38**(5), 603–609.
- [70] Bunday, B.D. & Al-Ayoubi, I.D. (1990). Likelihood and Bayesian estimation methods for Poisson process models in software reliability international, *Journal of Quality and Reliability Management* **7**(5), 9–18.
- [71] Ke, H. & Shen, F. (1999). Integrated Bayesian reliability assessment during equipment development, *International Journal of Quality and Reliability Management* **16**(9), 892–902.
- [72] Ryan, K.J. (2003). Some flexible families of intensities for non-homogeneous Poisson process models and their Bayes inference, *Quality and Reliability Engineering International* **19**(2), 171–181.
- [73] Meeker, W.Q. & Escobar, L.A. (1998). *Statistical Models for Reliability Data*, Wiley-Interscience, New York.
- [74] Fleming, T.R. & Harrington, D.P. (1991). *Counting Processes and Survival Analysis*, Wiley-Interscience, New York.
- [75] Andersen, P., Borgan, O., Gill, R. & Keiding, N. (1992). *Statistical Models Based on Counting Processes*, Springer.
- [76] Lindqvist, B.H. (2006). On the statistical modeling and analysis of repairable systems, *Statistical Science* **21**(4).

Related Article

Poisson Processes.

CHRISTOPHER MCCOLLIN

Repairable Systems: Bayesian Analysis

Introduction

The typical Bayesian paradigm, whether it is applied to repairable or nonrepairable systems, is to:

1. Formulate a prior distribution $p(\theta)$ that expresses our degree of belief in a parameter θ before any data are collected. Note that θ could be a vector.
2. Collect data $\mathbf{x}$ and form the likelihood function $L(\theta|\mathbf{x}) = f(\mathbf{x}|\theta)$.
3. Apply Bayes theorem to obtain the posterior distribution

$$p(\theta|\mathbf{x}) = cp(\theta)L(\theta|\mathbf{x}) \quad (1)$$

4. Use the posterior distribution to make inference for θ . For example, a point estimate for θ_k (the k th component of θ , if θ is a vector) is often taken to be the posterior mean

$$\hat{\theta}_k = E(\theta|\mathbf{x}) = \int \dots \int \theta_k p(\theta|\mathbf{x}) d\theta \quad (2)$$

We begin by presenting a Bayesian analysis for the homogeneous Poisson process, the simplest model for repairable systems. We then discuss the power law process (PLP). The concept of availability follows next and we present a Bayesian analysis of availability for the homogeneous **Poisson process** with gamma repair times. For multiple systems, we present **empirical Bayes** and hierarchical Bayes methods.

The Homogeneous Poisson Process

If we define $N(t)$ to be the (random) number of failures in the interval $[0, t]$, and more generally, $N(a, b)$ to be the number of failures in the interval (a, b) , then the homogeneous Poisson process is characterized by the following:

1. $N(0) = 0$ (the system begins with no failures).
2. If $a < b \leq c < d$, then $N(a, b)$ and $N(c, d)$ are independent random variables. This is called the *independent increments* property.

3. $\lim_{\Delta t \rightarrow 0} \frac{P(N(t, t + \Delta t) = 1)}{\Delta t} = \lambda$ and λ is independent of t .
4. $\lim_{\Delta t \rightarrow 0} \frac{P(N(t, t + \Delta t) \geq 2)}{\Delta t} = 0$. This precludes simultaneous failures.

These assumptions imply that the number of failures in an interval $[a, b]$ has a Poisson distribution with mean $\lambda(b - a)$. Note that the expected number of failures in an interval depends only on λ and the length of the interval, not on when the interval occurs. A consequence of this is that the number of failures in an interval of length t does not depend on whether the time interval covers the period when the system is young or old. Thus, the homogeneous Poisson process cannot be used to model systems that are improving or deteriorating.

Suppose now that the number n is predetermined and the system is observed until n failures occur. This is called *failure truncation*. We observe the failure times $t_1 < t_2 < \dots < t_n$ or the times between failures $x_1 = t_1, x_2 = t_2 - t_1, \dots, x_n = t_n - t_{n-1}$. The times between failure are independent exponential random variables with mean $1/\lambda$, so the likelihood function is

$$L(\mathbf{x}|\lambda) = \lambda^n \exp\left(-\lambda \sum_{i=1}^n x_i\right) = \lambda^n \exp(-\lambda t_n) \quad (3)$$

where t_n is the sum of the x_i 's and is a sufficient statistic. The conjugate prior for λ is the $\Gamma(a, b)$ distribution, which has probability density function (pdf)

$$p(\lambda) = \frac{b^a \lambda^{a-1}}{\Gamma(a)} e^{-b\lambda}, \quad \lambda > 0 \quad (4)$$

For this distribution, the mean and variance are

$$\mu = \frac{a}{b} \quad \text{and} \quad \sigma^2 = \frac{a}{b^2} \quad (5)$$

Thus, if we have prior knowledge about the parameter λ , especially concerning the mean and variance, then we can solve these equations for a and b to get

$$a = \frac{\mu^2}{\sigma^2} \quad \text{and} \quad b = \frac{\mu}{\sigma^2} \quad (6)$$

Combining equations (3) and (4) through Bayes theorem, we obtain the posterior

$$p(\lambda|\mathbf{x}) = c \lambda^{a+n-1} \exp[-(b + t_n)\lambda], \quad \lambda > 0 \quad (7)$$

2 Repairable Systems: Bayesian Analysis

which is seen to be the $\Gamma(a + n, b + t_n)$ distribution. Here c is the constant, independent of λ , that makes $p(\lambda|\mathbf{x})$ integrate to one. Inference is then based on this posterior distribution which reflects our knowledge about λ after seeing the data $\mathbf{x}$. For example, the posterior mean can be a point estimate of λ :

$$\hat{\lambda} = E[\lambda|\mathbf{x}] = \frac{a + n}{b + t_n} \quad (8)$$

The posterior variance can be used to assess the uncertainty in λ :

$$V[\lambda|\mathbf{x}] = \frac{a + n}{(b + t_n)^2} \quad (9)$$

A posterior probability interval for λ can be obtained using a normal approximation for the posterior distribution

$$\hat{\lambda} \pm z_{\alpha/2} \sqrt{V[\lambda|\mathbf{x}]} \quad (10)$$

If little or no prior information about λ is known, we can use the noninformative prior

$$p(\lambda) = \frac{1}{\lambda}, \quad \lambda > 0 \quad (11)$$

Combining equations (11) and (4), again using Bayes theorem, we obtain the posterior distribution

$$p(\lambda|\mathbf{x}) = c\lambda^{n-1} \exp[-t_n\lambda], \quad \lambda > 0 \quad (12)$$

showing that $\lambda|\mathbf{x}$ has a $\Gamma(n, t_n)$ distribution. Point and interval estimates for λ can be obtained from this result.

If testing of the repairable system stops at a predetermined time τ , rather than after a predetermined number of failures, then the sufficient statistic is $N(\tau)$, the number of observed failures. This is called *time truncation*. The likelihood is therefore

$$L(n|\lambda) = \frac{(\lambda\tau)^n e^{-\lambda\tau}}{n!}, \quad n = 0, 1, 2, \dots \quad (13)$$

Once again, the gamma distribution is the conjugate prior for λ . If this conjugate prior is chosen, then the posterior distribution has the form

$$p(\lambda|n) = c\lambda^{a+n-1} \exp[-(b + t)\lambda] \quad (14)$$

which is the $\Gamma(a + n, b + t)$ distribution.

Predicting the Number of Failures in a Future Interval

Suppose that a repairable system, whose failures are modeled by the homogeneous Poisson process, is observed until the time of the n th failure, t_n . The number of failures in the next τ time units, $N(t_n, t_n + \tau) = M$, is a random variable that has a Poisson distribution with mean $\lambda\tau$. Predicting the value of M then involves the randomness of this random variable *and* the uncertainty about the parameter λ . How can we combine these two sources of uncertainty? This is a difficult problem in the classical or frequentist framework, but it is straightforward in the Bayesian framework.

We can write the posterior of (λ, m) given the observed times of failure $\mathbf{t}$ as

$$p(\lambda, m|\mathbf{t}) = p(\lambda|\mathbf{t}) P(M = m|\lambda) \quad (15)$$

The marginal posterior distribution for m can then be obtained by integrating out λ :

$$p(m|\mathbf{t}) = \int_0^\infty p(\lambda|\mathbf{t}) P(M = m|\lambda) d\lambda \quad (16)$$

For the case of the noninformative prior $p(\lambda) = 1/\lambda$, it can be shown, (Rigdon and Basu [1]), that the posterior for M has the negative binomial distribution

$$P(M = m|\mathbf{t}) = \binom{m + n - 1}{m} \left(\frac{t_n}{t_n + \tau}\right)^n \left(\frac{\tau}{t_n + \tau}\right)^m, \quad m = 0, 1, 2, \dots \quad (17)$$

Point and interval predictions can be based on this marginal posterior distribution. Beiser and Rigdon [2] have shown that the coverage rates for these Bayesian intervals are slightly above the nominal levels when the noninformative prior is used.

Power Law Process

The assumptions for a nonhomogeneous Poisson process (NHPP) are the same as described in the section titled “The Homogeneous Poisson Process” above, except $\lim_{\Delta t \rightarrow 0} \frac{P(N(t, t + \Delta t) = 1)}{\Delta t} = \lambda(t)$ that is, the intensity is time dependent. The PLP is an NHPP whose **intensity function** has the parametric form

$$\lambda(t) = \frac{\beta}{\theta} \left(\frac{t}{\theta}\right)^{\beta-1}, \quad t > 0 \quad (18)$$

The Bayesian approach for the PLP follows the same general paradigm as described in the previous sections. The β parameter has a nice interpretation as a measure of improvement or deterioration:

1. $\beta = 1$ implies a constant intensity and there is no deterioration (or improvement),
2. $\beta < 1$ means that the intensity is decreasing and therefore there is reliability improvement, and
3. $\beta > 1$ means that the intensity is increasing and therefore there is reliability deterioration.

It may be reasonable to give prior information about the parameter β directly. The interpretation of the θ parameter is more difficult, especially in the absence of knowledge about β . ($\theta^{-\beta}$ can be interpreted as the expected number of failures through time $t = 1$, but this expression depends on both θ and β .) Guida *et al.* [3] have suggested that prior information about β and θ can be obtained, not directly, but rather in the following way. First, select a (marginal) prior distribution $g(\beta)$ for β . Then select some time T and consider the expected number of failures before time T . The expected number of failures η is

$$\eta = E(N(T)) = \int_0^T \lambda(t) dt = \left(\frac{T}{\theta}\right)^\beta \quad (19)$$

In a practical situation, it may be more reasonable to have prior information about η than about θ directly. We can then select a gamma prior distribution for η :

$$h(\eta) = \frac{b^a \eta^{a-1}}{\Gamma(a)} e^{-b\eta}, \quad \eta > 0 \quad (20)$$

We can select a and b to match the mean and variability of our prior knowledge about η .

If we assume independence regarding β and η , then

$$p(\beta, \eta) = g(\beta) h(\eta) \quad (21)$$

Since $\theta = T/\eta^{1/\beta}$, we can use the Jacobian to find that the joint prior for (β, θ) is

$$p(\beta, \theta) = g(\beta) \frac{\beta b^a}{\theta \Gamma(a)} \left(\frac{T}{\theta}\right)^{\beta a} \exp\left[-b \left(\frac{T}{\theta}\right)^\beta\right], \quad \beta > 0, \theta > 0 \quad (22)$$

The joint posterior can then be written as

$$p(\beta, \theta | \mathbf{t}) = c p(\beta, \theta) L(\beta, \theta | \mathbf{t})$$

$$= c g(\beta) \left(\frac{\beta}{\theta}\right)^n \left(\frac{P_n}{\theta^n}\right)^{\beta-1} \left(\frac{T}{\theta}\right)^{\beta a} \frac{\beta b^a}{\theta \Gamma(a)} \times \exp\left[-\left(\frac{t_n}{\theta}\right)^\beta - b \left(\frac{T}{\theta}\right)^\beta\right] \quad (23)$$

where

$$P_n = \prod_{i=1}^n t_i \quad (24)$$

Point and interval estimates for θ and β can be obtained from the joint posterior distribution.

The value of the intensity function at time t is often important, especially if t is the time of the conclusion of testing. If no further improvements are made to the system (which would be the case, for example, with software) then this would be the constant value of the intensity after release. Since $\lambda(t) = (\beta/\theta) (t/\theta)^{\beta-1}$, the posterior mean

$$E(\lambda(t) | \mathbf{t}) = \int_0^\infty \int_0^\infty (\beta/\theta) (t/\theta)^{\beta-1} p(\beta, \theta | \mathbf{t}) d\beta d\theta \quad (25)$$

would serve as a point estimate for $\lambda(t)$.

Other Models

The Bayesian methods described in the previous section can be applied to other models for repairable systems. The Bayesian approach will be the same in principle, but details will differ. For example, if the Weibull renewal process is assumed, then a prior would be placed on the parameters of the Weibull and the posterior distribution would then be computed *via* Bayes theorem. There are numerous other models that can be applied to repairable systems. See Ascher and Feingold [4] and Rigdon and Basu [1]. Also, *see Repairable Systems Reliability*.

Assessment of System Performance

There are often measures of system performance that involve more than one type of input. For example, the availability of a repairable system depends on both the failure process and the repair process, both of which must be estimated from the available data. It is often more straightforward to assess the uncertainty of such an estimate using the Bayesian paradigm.

4 Repairable Systems: Bayesian Analysis

Consider for example, the situation where the failure process is a homogeneous Poisson process with intensity λ , which is unknown, and the repair times are gamma random variables with pdf

$$g(x|\alpha, \beta) = \frac{\beta^\alpha x^{\alpha-1}}{\Gamma(\alpha)} e^{-\beta x}, \quad x > 0 \quad (26)$$

where α and β are also unknown. Suppose that the prior distribution for λ , the parameter of the failure process, is $p_1(\lambda)$ and the joint prior for (α, β) , the parameters of the repair process, is $p_2(\alpha, \beta)$. Suppose further, that we have observed the failure times $t_1 < t_2 < \dots < t_n$ (measured in operating time, which excludes repair times) and repair times $x_1, x_2, \dots, x_m$. Assuming independence between the failure process and the repair process, and also between the prior information about λ and (α, β) , we find that the posterior distribution for (λ, α, β) is

$$\begin{aligned} p(\lambda, \alpha, \beta | \mathbf{t}, \mathbf{x}) &= c p_1(\lambda) L(\lambda | \mathbf{t}) p_2(\alpha, \beta) L(\alpha, \beta | \mathbf{x}) \\ &= c p_1(\lambda) \lambda^n \exp[-\lambda t_n] \\ &\quad \times p_2(\alpha, \beta) \frac{\beta^{m\alpha} \left(\prod_{j=1}^m x_j \right)^{\alpha-1}}{[\Gamma(\alpha)]^m} \\ &\quad \times \exp \left[-\beta \sum_{j=1}^m x_j \right] \end{aligned} \quad (27)$$

The steady-state availability is

$$\begin{aligned} A_\infty &= \lim_{t \rightarrow \infty} P(\text{system works at time } t) \\ &= \frac{E(\text{time between failure})}{E(\text{time between failure}) + E(\text{repair time})} \\ &= \frac{1/\lambda}{1/\lambda + \alpha/\beta} = \frac{\beta}{\beta + \alpha\lambda} \end{aligned} \quad (28)$$

If we had exact values for λ , α , and β , we would have an exact value for A_∞ , but we have just estimates of these parameters. We can obtain a point estimate for A_∞ as the posterior mean,

$$\begin{aligned} \hat{A}_\infty &= E \left(\frac{\beta}{\beta + \alpha\lambda} \middle| \mathbf{t}, \mathbf{x} \right) \\ &= \int_0^\infty \int_0^\infty \int_0^\infty \frac{\beta}{\beta + \alpha\lambda} p(\lambda, \alpha, \beta | \mathbf{t}, \mathbf{x}) d\lambda d\alpha d\beta \end{aligned} \quad (29)$$

We could also find the posterior variance

$$\hat{V} = E \left(\left(\frac{\beta}{\beta + \alpha\lambda} \right)^2 \middle| \mathbf{t}, \mathbf{x} \right) - \hat{A}_\infty^2 \quad (30)$$

and use its square root (the posterior standard deviation) to assess the uncertainty in our estimate of the availability. Assessing the uncertainty in our estimate $\hat{A}_\infty$ would be difficult in the frequentist paradigm.

See **Queues in Reliability** for a description of how **queues** are used in reliability, especially regarding system performance measures such as availability. Bayesian methods offer a practical method for assessing system performance in such cases.

Empirical and Hierarchical Bayes Models

So far we have assumed that there is just a single repairable system. Suppose now that there are k copies of a system and all are simultaneously under test. The appropriate analysis depends on the assumptions about the similarity or dissimilarity of the systems.

To illustrate, suppose that the homogeneous Poisson process applies to each system. The following assumptions can be made:

1. All systems are the same. Times between failures are independent and identically distributed (i.i.d.) exponential random variables with mean $1/\lambda$ for all systems (λ is the same for all systems). In this case, we can combine the likelihood for each of the systems, along with the prior distribution for λ to obtain the posterior. This essentially leads to the superposition of point processes (see Pievatolo and Ruggeri [5]).
2. Systems are all different. Each system has its own failure rate λ and we have a prior distribution for each. In effect, we would analyze these systems separately, using Bayes theorem for each system.
3. Systems are different but similar enough that we may assume that the failure rates $\lambda_1, \lambda_2, \dots, \lambda_n$ are i.i.d. from some prior distribution $p(\lambda)$. If this prior distribution contains the parameter ϕ (which could be a vector), then we would write $p(\lambda|\phi)$. Figure 1 illustrates this situation.

If we assume that ϕ is fixed, but unknown, then we would have an empirical Bayes model. We could estimate ϕ by the method of maximum likelihood

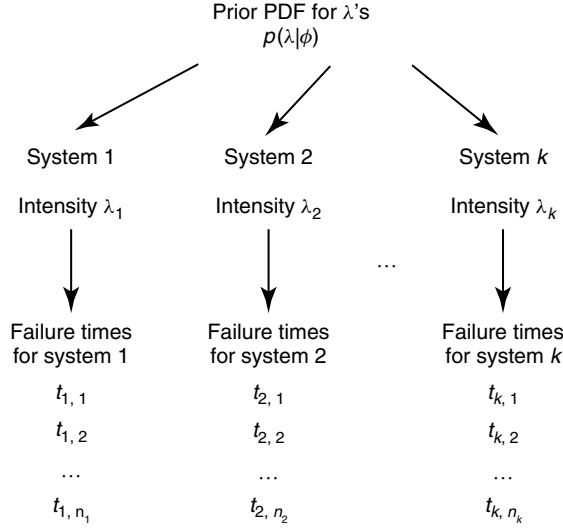


Figure 1 An empirical Bayes model for system failure intensities

and then we would obtain an estimated posterior distribution for each system:

$$p_i(\lambda_i | \mathbf{t}) = p(\lambda_i | \hat{\phi}) L(\lambda_i | \mathbf{t}) \quad (31)$$

where $\mathbf{t}$ is the array of all failure times. Note that the λ 's are unobservable random variables in this case.

Alternatively, we might assume that the hyperparameter ϕ is also unknown and place a prior $g(\phi)$ on it. This is a fully Bayesian model and is called a *hierarchical Bayes model*. See Figure 2. From this setup, we can obtain the marginal posterior distributions for λ_i , $i = 1, 2, \dots, k$.

If the PLP, rather than the homogeneous Poisson process were assumed, the approach would be similar, but a bit more involved. In Figure 1, the prior distribution would be for the pair (β, θ) , so we would write $p(\beta, \theta | \phi)$. An alternative assumption, reasonable for some situations, would be that the β parameter is the same for all systems, but the θ 's vary from system to system; we might further assume that $\theta_1, \theta_2, \dots, \theta_k$ are i.i.d. from some prior distribution $p(\theta | \phi)$. This would lead to an empirical Bayes model where the (fixed but unknown) parameters would be (β, ϕ) . Various other empirical Bayes

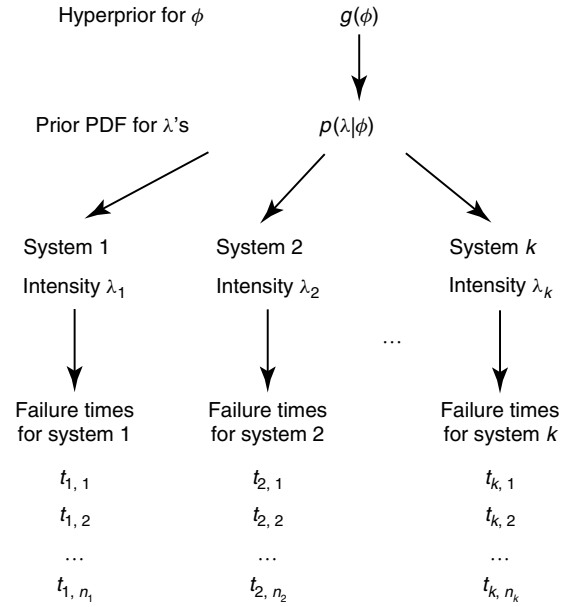


Figure 2 A hierarchical Bayes model for system failure intensities

or hierarchical Bayes models are possible. See Rigdon and Basu [1], Chapter 5 for a further discussion.

References

- [1] Rigdon, S.E. & Basu, A.P. (2000). *Statistical Methods for the Reliability of Repairable Systems*, John Wiley & Sons, New York.
- [2] Beiser, J.A. & Rigdon, S.E. (1997). Bayes prediction for the number of failures of a repairable system modeled by a nonhomogeneous Poisson process, *IEEE Transactions on Reliability* **R-46**, 291–295.
- [3] Guida, M., Calabria, R. & Pulcini, G. (1989). Bayes inference for a nonhomogeneous Poisson process with power intensity law, *IEEE Transactions on Reliability* **R-38**, 603–609.
- [4] Ascher, H. & Feingold, H. (1984). *Repairable Systems Reliability: Modeling, Inference, Misconceptions and Their Causes*, Marcel Dekker, New York.
- [5] Pievatolo, A. & Ruggeri, F. (2004). Bayesian reliability analysis of complex repairable systems, *Applied Stochastic Models in Business and Industry* **20**, 253–264.

STEVEN E. RIGDON

Flowgraph Models

Flowgraph Models

Flowgraph models are multistate models that are used to describe longitudinal, time-to-event data. They model stochastic processes that progress through various stages. Today's complex systems make the analysis of multistate models very important in reliability. Flowgraphs model potential outcomes, probabilities of outcomes, and waiting times for the outcomes to occur. Flowgraphs provide a natural model for complex engineering systems and reliability. A recent book on the subject is [1].

Examples of processes that have been modeled using flowgraphs include the internal mechanisms of cellular telephone networks [2], stages in the maintenance and repair of aircraft, or improving stages of assembly and rework in a manufacturing process. For example, modeling the stages of the cellular telephone network begins with a fully functioning network that proceeds through a series of degradations to a partially functioning network and eventually to a fully failed network. The network itself may be a local network composed of only a few cells or it may be composed of subsystems that comprise a network over an entire state or country. In aircraft maintenance, the macro level process can be represented by five main systems. However, each system consists of subsystems with thousands of elements that represent tasks involved in repair and maintenance. Data may be available in detail for the entire process or, more realistically, data may be available only partially at the component, subsystem, or system level. Combining such information across various levels of a system has become increasingly important in complex systems (*cf.* [3]).

Flowgraphs allow each component or set of components to have its own distribution, and flowgraph algebra provides a way to combine these varied distributions. The examples of the previous paragraph involve modeling the time until the occurrence of some event or events. The main interest is in modeling large, complex systems or subsystems of such systems.

Figure 1 presents a flowgraph model for a complex system. The states of the system, 0, 1, and 2,

represent the *functioning*, *degraded*, and *failed* system respectively. Alternatively, each of these nodes could also represent a subsystem. Our interest is in modeling the time to total failure, partial failure, or repair of the system. In addition we may also be interested in the number of times the system must be repaired after failure. We may also have various types of data on the system. We may have data on all of the transitions but the data may be censored, or we may have data on portions of the system but missing transition information. All of these situations can be handled by a flowgraph model. The end result from a flowgraph is a Bayes predictive density, survivor/reliability function, and hazard function. Additional quantities such as availability can also be calculated. If one prefers, the maximum-likelihood estimates can be computed.

This article begins by providing an introduction to the analysis of basic flowgraph series-parallel-loop systems. Then we analyze the complex system of Figure 1 and describe the notation used in the figure. The labels on the branches are transition probabilities, p_{ij} 's, and moment generating functions (MGFs), $M_{ij}(s)$'s, of the waiting time distributions for transitions to occur. These are ultimately used with flowgraph algebra to obtain the MGF of the distribution of the waiting time of interest. We also present a general rule to analyze larger systems. Most of the basic flowgraph analysis software is available on the web site associated with [1].

Series, Parallel, and Loop Systems

In a flowgraph model, the states or nodes represent outcomes. This is distinct from a graphical model where the nodes represent random variables. Flowgraphs are outcomes graphs, whereas graphical models are variables graphs. The nodes are connected by directed line segments called *branches*. Each branch has a transition probability and waiting time distribution associated with a transition from its beginning and ending nodes. The branches are labeled with *transmittances*, which involve probabilities and MGFs. In engineering, block diagrams are analyzed using Laplace transforms of functional blocks and the primary interest is system stability. In statistics, we

2 Flowgraph Models

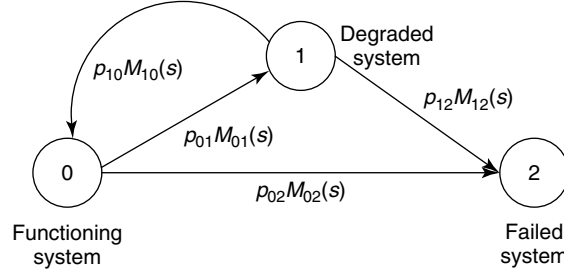


Figure 1 Flowgraph model for an overview of a complex system

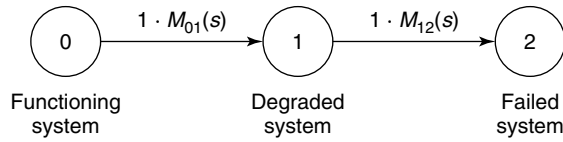


Figure 2 Flowgraph model for a series system

are interesting in quantities based on probability distributions so we use the MGF in place of the Laplace transform.

Series Systems

Figure 2 shows a series system with states 0, 1, and 2 representing functioning, degraded, and failed system states respectively.

Let Y_0 be the random waiting time in state 0 until the system reaches the degraded state, 1. Let Y_1 , independent of Y_0 , be the random waiting time in state 1 until the system fails, state 2. Starting in state 0, the system degrades with probability 1, so the transition probability from $0 \rightarrow 1$ is $p_{01} = 1$. The associated waiting time is represented by its MGF, $M_{01}(s) = E(e^{sY_0})$, provided that the expectation exists for s in an open neighborhood of 0. Once in state 1, the system eventually fails so that $p_{12} = 1$ and the waiting time MGF of the time until state 2 is reached is $M_{12}(s) = E(e^{sY_1})$. The overall waiting time, time to system failure, has MGF $M_{01}(s)M_{12}(s)$. This is exactly what one would do to compute the sum of two independent random variables, $Y_0 + Y_1$.

A few other definitions that are useful when working flowgraphs are those of a transmittance and a path. A *transmittance* consists of the transition probability multiplied by the MGF of the waiting time distribution. An *overall transmittance* is the transmittance of the entire flowgraph from the initial

to the end state. A *path* from state i to state j is any possible sequence of nodes from i to j that does not pass through any intermediate state more than once. A *path transmittance* is the product of all the branch transmittances for that path. In the simple flowgraph of Figure 2, there is only one path, $0 \rightarrow 1 \rightarrow 2$, the transmittance of the $0 \rightarrow 1$ branch is $1 \cdot M_{01}(s)$, and the overall transmittance is the same as the MGF of the overall waiting time distribution.

Flowgraph Models for Parallel Structures

The next basic component of a flowgraph is one where branches are in parallel (see Figure 3). In a parallel flowgraph, transition from a state is allowed to one of a set of outcomes.

State 0 represents the functioning system and from this state, we wait for one of three possible outcomes: state 1, 2, or 3. This flowgraph represents a *competing-risks* situation. The flowgraph models the observable parts of the competing-risks distribution, that is, the actual competitive waiting times and the probabilities of each outcome. The time to each individual type of failure is given by the branch transmittances. The time to failure regardless of type has MGF given by $p_{01}M_{01}(s) + p_{02}M_{02}(s) + p_{03}M_{03}(s)$. Series and parallel flowgraphs can be combined to model more complex processes.

Closed Parallel System

Figure 4 shows a closed parallel system. State 0 represents the functioning system, and states 1 and 2 represent the degraded and the failed system. The path $0 \rightarrow 1 \rightarrow 2$ is a series system and it is in parallel with the direct path from $0 \rightarrow 2$.

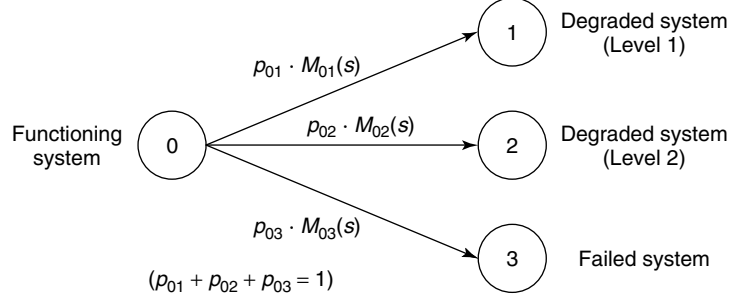


Figure 3 Flowgraph model for a parallel system

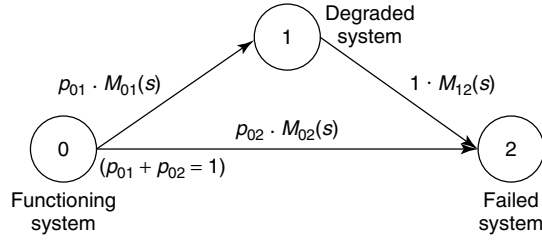


Figure 4 Flowgraph model for a closed parallel system

This three-state model is a classic competing-risks model and is also called the *illness–death model* in the multistate models literature [4].

We solve this flowgraph for the MGF of the overall time to system failure by first considering the upper path. Note that the path $0 \rightarrow 1 \rightarrow 2$ is a series flowgraph. Thus the transmittance of this path is $p_{01}M_{01}(s)M_{12}(s)$, where $p_{12} = 1$. We can replace the original flowgraph with the one in Figure 5, a flowgraph that has two paths in parallel.

We reduce this further by treating the flowgraph as having just two branches in parallel, as in Figure 3. If transition to state 1 occurs prior to transition to state 2, the $0 \rightarrow 1$ branch is taken with probability p_{01} , leading to the series path through the flowgraph. If transition to state 2 occurs prior to state 1, the transition is direct passage from $0 \rightarrow 2$. Solving this gives the overall MGF,

$$M(s) = p_{01}M_{01}(s)M_{12}(s) + p_{02}M_{02}(s) \quad (1)$$

Series and parallel elements are two basic components of flowgraphs. Series elements lead to distributions of sums of independent random variables, whereas parallel elements lead to finite mixtures of distributions of independent random variables. The

feedback loop is the third basic element of a flowgraph model. A *loop* is any closed path that returns to the initiating state without passing through any state more than once. Feedback loops arise naturally in the engineering setting and are generally used to represent recurrent events such as repairs to a system or a part.

Series Model with Feedback

Consider a portion of the system of Figure 1 given in Figure 6. This is a series system with feedback that allows for repair of the system.

To solve this flowgraph for the MGF of T , $M_T(s)$, we set up a balance equation at state 0. A balance equation equates the inputs into a state with the outputs of the state. In our case, the input into state 0 is $p_{10}M_{10}(s)$. The output from state 0 is $M_{01}(s)$. This is magnified by the factor $M_T(s)$, which is the MGF of the loop that is fed back into itself. In systems analysis, this factor is called the *gain* of the system. The balance equation for state 0 in Figure 6 is

$$M_T(s) = M_{01}(s)p_{10}M_{10}(s)M_T(s) + M_{01}(s)p_{12}M_{12}(s) \quad (2)$$

4 Flowgraph Models

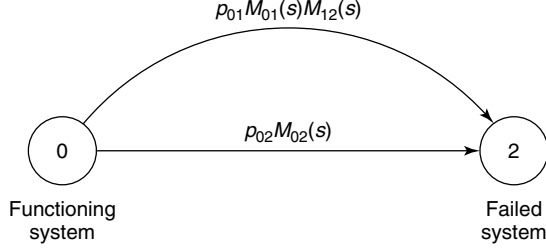


Figure 5 Reduced flowgraph model for a closed parallel system

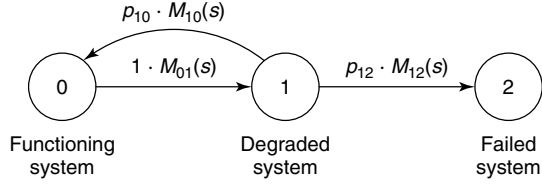


Figure 6 Flowgraph model with feedback

Note that the term $M_{01}(s)$ appears in both terms on the right-hand side since it is a transmittance passed through the system. Solving for $M_T(s)$ gives

$$M_T(s) = \frac{p_{12}M_{01}(s)M_{12}(s)}{1 - p_{10}M_{01}(s)M_{10}(s)} \quad (3)$$

where $p_{12} + p_{10} = 1$. This is the overall MGF of the distribution of T .

Complex Systems: Automated Analysis

In this section we present a detailed solution of the flowgraph of Figure 1. Note that this flowgraph has a combination of series, parallel, and loop structures, requiring several steps for reducing the flowgraph to solve for the overall waiting time MGF of interest. Mason [5] developed a rule in the context of graph theory for solving systems of linear equations. In its original implementation, Mason's rule did not involve probabilities or MGFs. However, flowgraph models can be solved by applying Mason's rule to the branch transmittances. This provides a systematic procedure for computing the transmittance of a flowgraph from any state i to any state j .

Mason's rule entails enumerating the paths and loops involved in traversing the system from any state i to any other state j . Often i and j are taken to be the input and output respectively but

that is not a requirement. In Figure 1, there are two paths from 0 to 2: $0 \rightarrow 1 \rightarrow 2$ and direct passage from $0 \rightarrow 2$. The transmittances of these paths are $p_{01}p_{12}M_{01}(s)M_{12}(s)$ and $p_{02}M_{02}(s)$. A first-order loop is any closed path that returns to the initiating state without passing through any state more than once. In Figure 1, $0 \rightarrow 1 \rightarrow 0$ constitutes a first-order loop. In general, a j th-order loop consists of j nontouching first-order loops. There are no loops of order 2 or more in Figure 1. The transmittance of a first-order loop is the product of the individual branch transmittances involved in the path. The transmittance of a higher-order loop is the product of the transmittances of the first-order loops it contains.

The general form of Mason's rule gives the overall transmittance, i.e. the MGF from input to output, as

$$M(s) = \frac{\sum_i P_i(s) \left[1 + \sum_j (-1)^j L_j^i(s) \right]}{1 + \sum_j (-1)^j L_j(s)} \quad (4)$$

where $P_i(s)$ is the transmittance for the i th path, $L_j(s)$ in the denominator is the sum of the transmittances over the j th-order loops, and $L_j^i(s)$ is the sum of the transmittances over j th-order loops sharing no common nodes with the i th path. For this system, $P_1(s) = p_{01}p_{12}M_{01}(s)M_{12}(s)$ and $P_2(s) =$

$p_{02}M_{02}(s)$. There is only one first-order loop and it touches both the paths at node 0 so that $L_1(s) = p_{01}p_{10}M_{01}M_{10}$ and $L_j^i(s) = 0$. Using equation (4), the overall waiting time from 0 to 2 is

$$M(s) = \frac{p_{01}p_{12}M_{01}(s)M_{12}(s) + p_{02}M_{02}(s)}{1 - p_{01}p_{10}M_{01}(s)M_{10}(s)} \quad (5)$$

In general, for systems with many states, paths, and loops, equation (4) is programmed using symbolic algebra. When the interest is in partial passage from state i to state j , equation (4) must be modified for the event of interest conditional on the occurrence of the event of interest. The conditional MGF in this case is

$$\frac{M_{ij}(s)}{M_{ij}(0)}$$

for partial passage from $i \rightarrow j$. While such flowgraph algebra provides the MGF of interest, it must be converted to a waiting time density, reliability, or hazard function in order to be useful. This can be performed *via* either a Laplace transform inversion routine [6] or a saddle point approximation [7].

Bayesian Prediction with Flowgraphs

The Bayes predictive density of a future observable Z given data $\mathcal{D}$ is

$$\begin{aligned} f_Z(z|\mathcal{D}) &= \int f_Z(z, \theta|\mathcal{D}) d\theta \\ &= \int f_Z(z|\theta)\pi(\theta|\mathcal{D}) d\theta \\ &\equiv E_{\theta|\mathcal{D}}[f_Z(z|\theta)] \end{aligned} \quad (6)$$

where Z has density $f_Z(z|\theta)$ and $\pi(\theta|\mathcal{D})$ is the posterior distribution of θ .

The flowgraph model gives us the MGF of the density $f_Z(z|\theta)$, so that for computational purposes, equation (6) can be written as

$$f_Z(z|\mathcal{D}) = \frac{\int f_Z(z|\theta)\mathcal{L}(\theta|\mathcal{D})\pi(\theta) d\theta}{\int \mathcal{L}(\theta|\mathcal{D})\pi(\theta) d\theta} \quad (7)$$

Parametric assumptions about the flowgraph allow us to specify the likelihood function for complete data. *Complete data* on the flowgraph model

consists of having complete transition information for each observation; however, the associated waiting time may be censored. *Incomplete data* refers to data with incomplete transition information in addition to censoring. For example, if an observation on the system is made in state 0 and 2 and we have additional information that the system did degrade to state 1, but we do not know when that occurred, we have incomplete data. In that case, likelihood construction is required [8].

Numerical Example: Bayes Predictive Distribution

For illustration we use the flowgraph of Figure 1 with simulated data. This example appears as a brief illustration in [9] with full details in [1]. The $0 \rightarrow 1$ waiting time is $IG(\lambda_1, \lambda_2)$ with mean λ_1/λ_2 , the $0 \rightarrow 2$ and $1 \rightarrow 0$ waiting times are $\Gamma(\alpha_1, \beta_1)$ (with mean α_1/β_1) and $\Gamma(\alpha_2, \beta_2)$ (with mean α_2/β_2), and the $1 \rightarrow 2$ waiting time is $Exp(\mu)$ (with mean $1/\mu$). The corresponding MGFs for the parametric distributions are

$$M_{01}(s) = \exp\{\lambda_1\lambda_2 - [\lambda_1^2(\lambda_2^2 - 2s)]^{1/2}\} \quad \text{for } s < \frac{\lambda_2^2}{2} \quad (8)$$

$$M_{02}(s) = \left(\frac{\beta_1}{\beta_1 - s}\right)^{\alpha_1} \quad \text{for } s < \beta_1 \quad (9)$$

$$M_{10}(s) = \left(\frac{\beta_2}{\beta_2 - s}\right)^{\alpha_2} \quad \text{for } s < \beta_2 \quad (10)$$

$$M_{12}(s) = \frac{\mu}{\mu - s} \quad \text{for } s < \mu \quad (11)$$

For real data, we use censored data histograms [10] to assist in choosing a parametric model for each branch waiting time in the flowgraph. Fifty runs of the system were simulated with parameter values $p_{01} = 0.7$, $p_{10} = 0.4$, $\lambda_1 = 3$, $\lambda_2 = 5$, $\alpha_1 = 4$, $\beta_1 = 6$, $\alpha_2 = 2$, $\beta_2 = 3$, and $\mu = 0.2$. Approximately 20% of the realizations were randomly right censored. For these data, all of the censored observations were in state 1. The system made 54 transitions from $0 \rightarrow 1$, 24 transitions from $0 \rightarrow 2$, 15 transitions from $1 \rightarrow 2$, and 28 transitions from $1 \rightarrow 0$. There were 11 censored observations.

Let $f_{gh}(\cdot)$ and $F_{gh}(\cdot)$ be the density and cumulative distribution function (**cdf**), respectively, of the waiting time in state g until transition to state h .

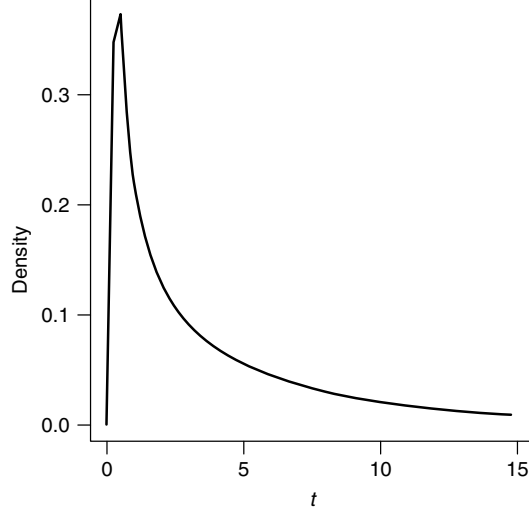


Figure 7 Bayes predictive density using slice sampling for the reversible illness–death model

Let x_{igh} be the i th uncensored observation from state g to state h , and let x_{jg}^* be the j th censored observation where censoring occurred in state g . Let n_{gh} be the number of individual transitions from state g to state h , and let n_g^* be the number of observations censored in state g . For this example, $n_0^* = 0$ and $n_1^* = 11$. Then the likelihood function is

$$\begin{aligned}
 \mathcal{L}(p_{01}, p_{10}, \lambda_1, \lambda_2, \alpha_1, \beta_1, \alpha_2, \beta_2, \mu | \mathcal{D}) \\
 &= \prod_{i=1}^{n_{01}} p_{01} f_{01}(x_{i01} | \lambda_1, \lambda_2) \prod_{i=1}^{n_{12}} (1 - p_{10}) f_{12}(x_{i12} | \mu) \\
 &\quad \times \prod_{i=1}^{n_{10}} p_{10} f_{10}(x_{i10} | \alpha_2, \beta_2) \\
 &\quad \times \prod_{i=1}^{n_{02}} (1 - p_{01}) f_{02}(x_{i02} | \alpha_1, \beta_1) \\
 &\quad \times \prod_{j=1}^{n_0^*} \left\{ 1 - [p_{01} F_{01}(x_{j0}^* | \lambda_1, \lambda_2) \right. \\
 &\quad \left. + (1 - p_{01}) F_{02}(x_{j0}^* | \alpha_1, \beta_1)] \right\} \\
 &\quad \times \prod_{j=1}^{n_1^*} \left\{ 1 - [p_{10} F_{10}(x_{j1}^* | \alpha_2, \beta_2) \right. \\
 &\quad \left. + (1 - p_{10}) F_{12}(x_{j1}^* | \mu)] \right\}
 \end{aligned} \tag{12}$$

The MGF of the overall survival time from $0 \rightarrow 2$ is found using the symbolic algebra software maple. We use noninformative priors on $\theta = (p_{01}, p_{10}, \lambda_1, \lambda_2, \alpha_1, \beta_1, \alpha_2, \beta_2, \mu)$. A *Unif*(0, 1) prior is used for both p_{01} and p_{10} . Each of the remaining parameters has a $\Gamma(1/1000, 1/1000)$ prior. Sampling from the posterior was done using slice sampling [11], a **Markov chain Monte Carlo** method that is globally applicable to flowgraph models. The resulting Bayes predictive density function is presented in Figure 7. Notice that the mixture of the inverse Gaussian and gamma distributions creates a very skewed overall waiting time distribution. Other functions of interest such as the reliability function, hazard function, availability function, or specific **quantiles** of the cdf can also be calculated as required. Maximum-likelihood estimates of these quantities can also be calculated. Software for flowgraph models is freely available at the web site associated with [1].

Summary

Flowgraph models provide a data analytic technique to model complex systems. Flowgraph models are semi-**Markov processes** but the transform methodology used with flowgraphs circumvents the difficulty associated with semi-Markov processes. The method is dynamic in that the systems and sub-systems are allowed to change and be updated

as the requirements change. The method allows for the incorporation of time-to-event and transition information data at all levels of the system, which makes it more than just a probabilistic method. The flowgraph models observable waiting times and does not make a proportional hazards assumption.

References

- [1] Huzurbazar, A.V. (2005). *Flowgraph Models for Multistate Time-to-Event Data*, John Wiley & Sons, New York.
- [2] Huzurbazar, A.V. (2000). Modeling and analysis of engineering systems data using flowgraph models, *Technometrics* **42**, 300–306.
- [3] Hamada, M., Martz, H., Reese, C.S., Graves, T., Johnson, V. & Wilson, A. (2004). A fully Bayesian approach for combining multilevel failure information in fault tree quantification and optimal follow-up resource quantification, *Reliability Engineering and System Safety* **86**, 297–305.
- [4] Andersen, P.K. & Keiding, N. (2002). Multi-state models for event history analysis, *Statistical Methods in Medical Research* **11**, 91–115.
- [5] Mason, S.J. (1953). Feedback theory: some properties of signal flow graphs, *Proceedings of the Institute of Radio Engineers* **41**, 1144–1156.
- [6] Abate, J. & Whitt, W. (1992). The Fourier-series method for inverting transforms of probability distributions, *Queueing Systems* **10**, 5–88.
- [7] Jensen, J. (1995). *Saddlepoint Approximation*, Oxford Press, New York.
- [8] Williams, B.J. & Huzurbazar, A.V. (2006). Posterior sampling with constructed likelihood functions: an application to flowgraph models, *Applied Stochastic Models in Business and Industry* **22**(2), 127–137.
- [9] Huzurbazar, A.V. (2004). Multistate models, flowgraph models and semi-Markov processes, *Communications in Statistics: Theory and Methods* **33**, 457–474.
- [10] Huzurbazar, A.V. (2005). A censored data histogram, *Communications in Statistics: Simulation and Computation* **34**(1), 113–120.
- [11] Neal, R. (2003). Slice sampling, *Annals of Statistics* **31**, 705–767.

Related Articles

Coherent Systems; Parallel, Series, and Series-Parallel Systems; Repairable Systems Reliability; System Reliability: Monte Carlo Estimation; Repairable Systems: Bayesian Analysis; Queues in Reliability.

APARNA V. HUZURBAZAR

Queues in Reliability

Introduction

A queuing system is a mathematical model to characterize situations where certain units, called *customers*, arrive in continuous time in order to receive a service or facility provided by other units, called *servers*. When customers arrive, they are immediately served if there is an empty server. Otherwise, the customers must spend some time waiting in line for service until a server becomes available. Queuing theory is the methodology devoted to the stochastic description of queuing systems and the study of their performance measures, such as the queue length, waiting times, and so on. Some essential references on queuing theory are found in [1–5].

Queuing theory can be very useful in solving many complex reliability problems. Component failures or machine breakdowns can be treated as arrivals of customers and the repair or replacement of failed components may be seen as the service facility in a queuing system. The number of repair facilities or maintenance crews is equivalent to the number of servers. It is useful to find the analogies between reliability and queuing theory because it is frequently found in practice that a given reliability problem has been previously studied with an equivalent problem in queuing theory [6]. For example, the number of failed components corresponds to the *queue length* and the time elapsed from a failure to repair is equivalent to the waiting time in the system or *sojourn time*. Note that the notion of repair may also refer to preventive maintenance activities performed before system failures. For example, a deteriorated machine can be repaired to an upgraded status before it fails.

A queuing model is completely described by six characteristics that, following Kendall's notation [7], are given by $A/B/c/K/m/R$, where A and B describe the arrival and service pattern, respectively, c the number of servers, K the system capacity or maximum number of customers allowed in the system, m the size of the customer population, and R the discipline or the order in which customers are selected for service. Different symbols are traditionally used for the type of distribution of A and B . For example, M denotes exponential (Markovian) random variables, D degenerate (constant) variables,

and G general distributions. For instance, the $M/G/1$ queue denotes a single-server system with exponential interarrival times and general service time distribution. Note that when nothing is said about K , m , and R , it is assumed that there is an infinite capacity, an infinite customer population, and a first-come-first-served discipline.

In the case of reliability, the main characteristic of queues is that requests for service are usually generated by a finite customer population, that is, $m < \infty$. This is because, in general, there is a limited number of units (machines, devices, etc.) which can fail, and when they are all in the system (being repaired or waiting for repair), no more can arrive. Then, the arrivals of requests do not form a renewal process as the arrivals may depend on the number of units staying at the system. This is an essential difference from typical queuing systems, where the population of potential arrivals can be considered to be effectively limitless. For example, in telephone operations the number of customers is usually large enough that the probability of having all of them wanting service at once is extremely small. However, in repair of machinery, the whole set of machines may be simultaneously out of order. Queuing systems with finite population are known as *finite source queues*, which have been extensively studied in the queuing literature, see e.g., [8] and the detailed bibliography provided in [9]. Finite-source queuing models are used to describe the classical *machine repairmen problem* which is described in the following section.

The Machine Repairmen Problem

Consider a repairable system (see **Repairable Systems Reliability**) consisting of r machines which operate independently of each other and may eventually fail. When a machine breaks down, it has to be repaired by one of the c repairmen, if there are any available. Otherwise, it has to wait for service until a server becomes available. Thus, the operating machines are outside the queuing system and enter the system only when they break down and require repair, as shown in Figure 1. Machine lifetimes and repair times are independent random variables with rates λ and μ , respectively. Then, the so-called mean time between failures (MTBF) and mean time to repair (MTTR) are given by $1/\lambda$ and $1/\mu$, respectively.

2 Queues in Reliability

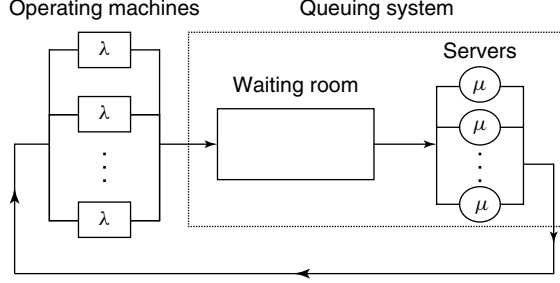


Figure 1 Machine repairmen/interference problem

This model is known as *the machine repairmen problem* or the *machine interference problem* since the machines interfere with each other when all the servers are busy and new requests for service arrive. It can be used to describe a large variety of real situations, such as manufacturing systems (see e.g., [10]), telecommunication traffic (see e.g., [11]), computer systems (see e.g., [12]), inventory models (see e.g., [13]), and so on. This broad range of applications has motivated a large amount of research on this problem in the literature. For two reviews, see [14, 15].

The simplest machine repairmen model supposes that both lifetimes and service times are exponentially distributed, there is no limit in the waiting room, the repair is carried out by a first-in-first-out discipline and after having been served, each machine starts working again as good as new. The queuing model describing this system is denoted by $M/M/c/r/r$ or, equivalently, $M/M/c//r$. Note that there is an important difference from the notation used for the arrival process in infinite source queues, such as the $M/M/c$ queue, where the symbol M indicates a Poisson arrival process. However, in the $M/M/c/r/r$ queue, the time between failures of each machine is exponential but the arrival of failures from the whole set of machines does not follow a **Poisson process** because, as mentioned before, these arrivals depend on the number of machines waiting for repair. In particular, when the number of machines waiting for service increases, the arrival rate of further service demands decreases. More specifically, when n machines are “down” for repair, there are $(r - n)$ operating machines and the time until the next breakdown is the minimum of $(r - n)$ exponential distributions with rate λ , which is also an exponential distribution of rate $(r - n)\lambda$. Then, given that there

are n machines being repaired or waiting for repair, the arrival rate or average number of fails per unit time is given by,

$$\lambda_n = \begin{cases} (r - n)\lambda, & \text{if } 0 \leq n \leq r \\ 0, & \text{if } n \geq r \end{cases} \quad (1)$$

and the average number of repaired machines per unit time is given by,

$$\mu_n = \begin{cases} n\mu, & \text{if } 0 \leq n \leq c \\ c\mu, & \text{if } n \geq c \end{cases} \quad (2)$$

Thus, it can be shown that the number of machines in the $M/M/c/r/r$ system follows a birth–death process, in which an arrival of a failure is interpreted as a birth and the completion of a repair is interpreted as a death, see e.g., [3]. Then, the long-run probability of having n machines in the system can be shown to be,

$$p_n = \begin{cases} \binom{r}{n} \left(\frac{\lambda}{\mu}\right)^n p_0, & \text{if } 1 \leq n \leq c \\ \binom{r}{n} \frac{n!}{c^{n-c}c!} \left(\frac{\lambda}{\mu}\right)^n p_0, & \text{if } c \leq n \leq r \end{cases} \quad (3)$$

where p_0 is obtained from $p_0 + \dots + p_r = 1$. Using these probabilities, we can calculate for example the mean number of machines in the system (being repaired or waiting for repair), L , and the mean number of machines waiting in the queue for repair, L_q , using,

$$L = \sum_{n=1}^r np_n \quad \text{and} \quad L_q = \sum_{n=c+1}^r (n - c) p_n \quad (4)$$

As noted in equation (1), the arrival rate, given that the system is in state n , is $(r - n)\lambda$. Then, the unconditional average rate at which machines arrive to the system (also called the *effective mean arrival Rate*) is given by,

$$\lambda_{\text{eff}} = \sum_{n=0}^{r-1} (r - n)\lambda p_n = \lambda(r - L) \quad (5)$$

Thus, we can also obtain the mean waiting time, W , a machine spends in the system (from failure to recovery) and the mean queuing time, W_q , a machine

waits in the queue for repair using the well-known Little's formulas,

$$L = \lambda_{\text{eff}} W \quad \text{and} \quad L_q = \lambda_{\text{eff}} W_q \quad (6)$$

Moreover, it is also possible to derive the distribution functions of the waiting time and queuing time in addition to their mean values, W and W_q , see e.g., [2, 3]. Some other performance measures such as the probability that a broken machine must wait for repair can also be obtained. Finally, it is worth observing that these results also hold even though the lifetimes are not exponentially distributed. That is, equation (3) is also valid for the $G/M/c/r/r$ queuing system. Furthermore, equation (3) also holds for systems with exponential lifetimes, general service times, and the same number of repairmen as machines, i.e. for the $M/G/c/c/c$ queue, see [3].

These results have also been extended to more complicated situations where the lifetimes and repair durations follow more general distributions than the exponential one, such as the Erlang distribution, mixtures of exponentials, Coxian and phase-type distributions (see e.g., [16, 17]). Many other features have been introduced to extend the applicability of the machine repairmen problem, including for example *vacation* models where the servers can take a vacation of random duration when the system is empty (see e.g., [18]), *retrial* models where machines that find all the servers busy do not enter the queue but instead reapply for service at random intervals (see e.g., [19]), *unreliable* servers that may themselves break down (see e.g., [20]), *priority* models where a group of machines are served with priority over the others (see e.g., [21]), *k-out-of-r* systems where at least k machines must be functioning for the whole system to work properly (see e.g., [22]), heterogeneous failure modes (see e.g., [23]), and so on. Another important extension is the use of *sparers*, also called *stand-by* or *sparing redundancy* (see **Reliability of Redundant Systems**), which is introduced in the following section.

The Machine Repairmen Problem with Sparers

Assume now that in addition to the r working machines considered previously, there is another set of s machines standing by as spares such that, when a working machine fails, it is immediately substituted by a spare machine if any is available. When a machine is repaired, it becomes a spare unless the

number of working machines is less than r , in which case the repaired machine is restarted immediately.

Consider an $M/M/c/r/r$ with s spares where the lifetimes and repair times are exponentially distributed with rates λ and μ , respectively. Then, given that there are n machines being repaired or waiting for repair, the arrival rate is now given by,

$$\lambda_n = \begin{cases} r\lambda, & \text{if } 0 \leq n \leq s \\ (r + s - n)\lambda, & \text{if } s \leq n \leq s + r \\ 0, & \text{if } n \geq s + r \end{cases} \quad (7)$$

and μ_n is the same as given in equation (2). Thus, this model can also be described by a birth-death process [3]. And then, it can be shown that the probability of having n machines in the system when $c \leq s$ is given by,

$$p_n = \begin{cases} \frac{r^n}{n!} \left(\frac{\lambda}{\mu}\right)^n p_0, & \text{if } 1 \leq n \leq c \\ \frac{r^n}{c^{n-c} c!} \left(\frac{\lambda}{\mu}\right)^n p_0, & \text{if } c \leq n \leq s \\ \frac{r^s r!}{(r + s - n)! c^{n-c} c!} \left(\frac{\lambda}{\mu}\right)^n p_0, & \text{if } s \leq n \leq s + r \end{cases} \quad (8)$$

and when $c > s$,

$$p_n = \begin{cases} \frac{r^n}{n!} \left(\frac{\lambda}{\mu}\right)^n p_0, & \text{if } 1 \leq n \leq s \\ \frac{r^s r!}{(r + s - n)! n!} \left(\frac{\lambda}{\mu}\right)^n p_0, & \text{if } s \leq n \leq c \\ \frac{r^s r!}{(r + s - n)! c^{n-c} c!} \left(\frac{\lambda}{\mu}\right)^n p_0, & \text{if } c \leq n \leq s + r \end{cases} \quad (9)$$

Using these probabilities, we can obtain the mean number of machines in the system and in the queue, L and L_q , respectively, as in equation (4). Also, we can obtain the effective mean arrival rate to the system by,

$$\begin{aligned} \lambda_{\text{eff}} &= \sum_{n=0}^{s-1} r\lambda p_n + \sum_{n=s}^{s+r} (r + s - n)\lambda p_n \\ &= \lambda \left(r - \sum_{n=s}^{s+r} (n - s) p_n \right) \end{aligned} \quad (10)$$

which allows us to obtain the mean waiting time in the system and the mean queuing time, W and

W_q , respectively, using Little's formula given in equation (6).

An interesting particular case appears when the number of spares, s , is very large. Then, the arrival rate λ_n given in equation (7) will be essentially constant and equal to λr , so that the arrival process will be well approximated by a Poisson process of rate λr . Thus, the system will converge to an infinite source $M/M/c$ model, whose stationary distribution can be obtained as the limit of equation (8) when s goes to infinity [2, 3], that is,

$$p_n = \begin{cases} \frac{(\lambda r)^n}{n! \mu^n} p_0, & \text{if } n \leq c \\ \frac{(\lambda r)^n}{c^{n-c} c! \mu^n} p_0, & \text{if } n \geq c \end{cases} \quad (11)$$

As shown in this case, infinite-source queues are in general much simpler than finite-source queues and consequently, queuing theory on infinite systems is much more developed. Thus, in practice, it may be convenient to assume an infinite source if the customer population is finite but large. Note that this will be a good approximation when the arrival process depends only in a negligible way on the number of customers already in the system.

Inference for Queues

Failure and repair times are random variables and the distributions and parameters describing their behavior are unknown in practice. Thus, the use of statistical methods (*see Repairable Systems: Statistical Inference*) becomes necessary to estimate the quantities of interest in the system. First, an experiment to collect data from the system must be designed. If we assume independence between the lifetimes and repair durations, a simple experiment providing complete information about the system consists in observing independently n lifetimes $\{x_1, \dots, x_n\}$ and m repair times $\{y_1, \dots, y_m\}$. Using these data, the traditional approach to queuing inference is based on **maximum-likelihood estimation** of the failure and service parameters. For example, the likelihood function of the parameters in the $M/M/c/r/r$ queue is given by,

$$l(\lambda, \mu) \propto \lambda^n \exp\left(-\sum_{i=1}^n x_i\right) \mu^m \exp\left(-\sum_{i=1}^m y_i\right) \quad (12)$$

and then, the maximum-likelihood estimators of λ and μ are given by $\hat{\lambda} = \bar{x}^{-1}$ and $\hat{\mu} = \bar{y}^{-1}$, respectively, where $\bar{x}$ and $\bar{y}$ denote the means of the lifetimes and repair times, respectively. These estimators can be plugged in to obtain estimations of the performance measures. For example, the stationary probabilities, p_n , can be estimated by replacing λ and μ by $\hat{\lambda}$ and $\hat{\mu}$, respectively, in equation (3). However, using this classical approach, it is often difficult to obtain measures of the uncertainty in the estimated performance measures.

An alternative approach is the use of the Bayesian methodology (*see Repairable Systems: Bayesian Analysis*), which is specially well suited for queues [24]. Consider again the $M/M/c/r/r$ model and assume independent gamma **prior distributions** for λ and μ ; say $\lambda \sim G(\alpha_0, \beta_0)$ and $\mu \sim G(a_0, b_0)$. From equation (12), it is easy to see that the posterior distributions are also independent gamma distributions, $\lambda | \text{data} \sim G(\alpha, \beta)$ and $\mu | \text{data} \sim G(a, b)$, where $\alpha = \alpha_0 + n$; $\beta = \beta_0 + \sum_{i=1}^n x_i$; $a = a_0 + m$; $b = b_0 + \sum_{j=1}^m y_j$. Using these posterior distributions, we can approximate the posterior predictive distributions of any performance measure. For example, we can simulate K values $\{\lambda^{(k)}, \mu^{(k)}\}_{k=1}^K$ from the joint posterior distribution of (λ, μ) , and approximate the posterior mean of p_n with,

$$E[p_n | \text{data}] \approx \frac{1}{K} \sum_{k=1}^K p_n^{(k)} \quad (13)$$

where $p_n^{(k)}$ is the value of p_n for $\lambda = \lambda^{(k)}$ and $\mu = \mu^{(k)}$. Analogously, we can easily approximate the posterior variance, **percentiles**, and credible intervals of p_n . For further details of Bayesian methods for the $M/M/c/r/r$ model, see [25]. Also, Bayesian analysis of more general queues, where the exponential assumption for arrivals and services is relaxed, can be found in e.g., [26] and the references therein.

References

- [1] Kleinrock, L. (1976). *Queuing Systems*, John Wiley & Sons, New York.
- [2] Allen, A.O. (1976). *Probability, Statistics, and Queuing Theory with Computer Science Applications*, 2nd Edition, Academic Press, Boston.
- [3] Gross, D. & Harris, C.M. (1998). *Fundamentals of Queuing Theory*, 3rd Edition, John Wiley & Sons, New York.

- [4] Trivedi, K.S. (2002). *Probability and Statistics with Reliability, Queuing, and Computer Science Applications*, 2nd Edition, John Wiley & Sons, New York.
- [5] Medhi, J. (2002). *Stochastic Models in Queuing Theory*, 2nd Edition, Academic Press, Boston.
- [6] Kovalenko, I.N., Kuznetsov, N.Y. & Pegg, P.A. (1997). *Mathematical Theory of Reliability of Time Dependent Systems with Practical Applications*, John Wiley & Sons, New York.
- [7] Kendall, D.G. (1953). Stochastic processes occurring in the theory of queues and their analysis by the method of imbedded Markov chains, *Annals of Mathematical Statistics* **24**, 338–354.
- [8] Takagi, H. (1993). *Queuing Analysis: A Foundation of Performance Evaluation Finite Systems*, North Holland, Amsterdam, Vol. 2.
- [9] Sztrik, J. (2002). *Finite Source Queuing Systems and Their Applications: A Bibliography*. Working paper 2002, Institute of Mathematics & Informatics, University of Debrecen. Available from: <http://irh.inf.unideb.hu/user/jsztrik/research/fsqreview.pdf>.
- [10] Buzacott, J. & Shanthikumar, J. (1993). *Stochastic Models of Manufacturing Systems*, Prentice Hall, Englewood Cliffs.
- [11] Daigle, J.N. (2004). *Queuing Theory with Applications to Packet Telecommunication*, Springer, New York.
- [12] Fortier, P.J. (2003). *Computer Systems Performance Evaluation and Prediction*, Digital Press, Burlington.
- [13] Perlman, Y., Mehrez, A. & Kaspi, M. (2001). Setting repair policy in a multi-echelon repairable-item inventory system with limited repair capacity, *The Journal of the Operational Research Society* **52**, 198–209.
- [14] Steckle, K.E. & Aronson, J.E. (1985). Review of operator/machine interference models, *International Journal of Production Research* **23**, 129–151.
- [15] Haque, L. & Armstrong, M.J. (2007). A survey of the machine interference problem, *European Journal of Operational Research* **179**, 469–482.
- [16] Neuts, M.F. & Meier, K.S. (1981). On the use of phase type distributions in reliability modelling of systems with two components, *Operation Research Spectrum* **2**, 227–234.
- [17] Carmichael, D.G. (1987). Machine interference with general repair and running times, *Mathematical Methods of Operations Research* **31**, B115–B133.
- [18] Gupta, S.M. (1997). Machine interference problem with warm spares, server vacations and exhaustive service, *Performance Evaluation* **29**, 195–211.
- [19] Falin, G. & Templeton, J. (1997). *Retrial Queues*, Chapman & Hall, London.
- [20] Wang, K.H. & Hsu, L.Y. (1995). Cost analysis of the machine repair problem with r non-reliable service stations, *Microelectronics Reliability* **35**, 923–934.
- [21] Posafalvi, A. & Sztrik, J. (1989). On the heterogeneous machine interference problem with priority and ordinary machines, *European Journal of Operational Research* **41**, 54–63.
- [22] Krishnamoorthy, A. & Ushakumari, P.V. (2001). K-out-of-n:G system with repair: the D-policy, *Computers and Operations Research* **28**, 973–981.
- [23] Palesano, J. & Chandra, J. (1986). A machine interference problem with multiple types of failures, *International Journal of Production Research* **24**, 567–582.
- [24] Armero, C. & Bayarri, M.J. (1999). Dealing with uncertainties in queues and networks of queues: a Bayesian approach, in *Multivariate Analysis, Design of Experiments, and Survey Sampling*, S. Ghosh, ed, Marcel Dekker, New York, pp. 579–608.
- [25] Castellanos, M.E., Morales, J., Mayoral, A.M., Fried, R. & Armero, C. (2006). On Bayesian design in finite source queues, in *COMPSTAT 2006 Proceedings in Computational Statistics*, A. Rizzi & M. Vichi, eds, Physica-Verlag, Heidelberg, pp. 1381–1388.
- [26] Ausin, M.C., Lillo, R.E. & Wiper, M.P. (2007). Bayesian control of the number of servers in a GI/M/c queueing system, *Journal of Statistical Planning and Inference* **137**, 3043–3057.

Related Articles

***k*-out-of-*n* Systems; Parallel, Series, and Series-Parallel Systems; Repairable Systems Reliability; Reliability of Redundant Systems; Repairable Systems: Statistical Inference; Repairable Systems: Bayesian Analysis**

M. CONCEPCIÓN AUSÍN

Imperfect Repair, Counting Processes

Introduction

The effect of a failure and a subsequent repair on the performance of a repairable system is what distinguishes many models for **repairable systems**. For example, if the repair brings the system to the condition just before the failure, then the repair is said to be minimal. This leads to the nonhomogeneous **Poisson process** (NHPP). By contrast, if a unit is replaced with a new one, then the times between failure are independent and identically distributed (i.i.d.). Such a repair is called a *renewal* and the resulting process is called a *renewal process*. **Imperfect repair** models are ones in which a renewal is performed with a specified probability; otherwise **minimal repair** is done. Many models are compromises between renewal and minimal repair. For example, the modulated power law process (MPLP) assumes that a repaired unit is in better condition after a repair than just before the failure, but is not as good as new. Another class of models, called *reduction of failure intensity models*, is another compromise between renewal and minimal repair. These models are discussed further in the following sections.

Minimal Repair and the NHPP

Let $N(t)$ denote the number of failures of the system through time t . The assumptions behind the NHPP are

1. $N(0) = 0$ (the system begins with no failures).
2. If $a < b \leq c < d$, then $N(a, b)$ and $N(c, d)$ are independent random variables. This is called the *independent increments* property.
3. $\lim_{\Delta t \rightarrow 0} \frac{P(N(t, t+\Delta t)=1)}{\Delta t} = \lambda(t)$.
4. $\lim_{\Delta t \rightarrow 0} \frac{P(N(t, t+\Delta t) \geq 2)}{\Delta t} = 0$, precluding simultaneous failures.

Repair time is negligible, or is not included in operating time. The **intensity function** in 3 above is independent of the previous history of the process. Thus, it does not matter whether a failure occurred just **moments** before time t , or whether there has

been no failure since the initial start-up of the system. The failure intensity will be the same in either case. The repair does nothing to the performance of the system except to bring it back to the state just before the failure. Such a model is often called *bad-as-old*, because a repairable system is as bad as it was just before the failure. Models for which the system is renewed after each failure, in other words, **renewal processes**, are sometimes called *good-as-new*.

If $\mathcal{H}_t$ denotes the entire failure history of the process, then we call

$$\lambda(t|\mathcal{H}_t) = \lim_{\Delta t \rightarrow 0} \frac{P(N(t, t + \Delta t) = 1|\mathcal{H}_t)}{\Delta t} \quad (1)$$

the *generalized intensity function*. Thus, for the NHPP,

$$\lambda(t|\mathcal{H}_t) = \lambda(t) \quad (2)$$

irrespective of the failure history, and for the renewal process,

$$\lambda(t|\mathcal{H}_t) = \lambda(x) \quad (3)$$

where x is the time since the most recent failure ($x = t - t_n$, where t_n is the time of the most recent failure). There are models, called *self-exciting processes*, for which the failure history affects the status of the system in a more direct way (see [1]).

A common form for the intensity function is

$$\lambda(t) = \frac{\beta}{\theta} \left(\frac{t}{\theta} \right)^{\beta-1}, \quad t > 0 \quad (4)$$

An NHPP with this intensity is called the *power law process* (see **Repairable Systems: Statistical Inference**). Typically, β and θ are unknown parameters that must be estimated from the observed failure times.

Imperfect Repair

Brown and Proschan [2] developed a model under the following assumptions. Failures occur according to an NHPP with intensity function $\lambda(t)$. When a failure occurs, the repair is minimal (in the sense of the previous section) with probability p , and the repair is perfect with probability $q = 1 - p$. Thus, the system is brought to a bad-as-old condition with probability p , and to a good-as-new condition with probability q . Of course, special cases include

$$p = 0 \Rightarrow \text{NHPP} \quad (5)$$

$$p = 1 \Rightarrow \text{Renewal process} \quad (6)$$

The time epochs of perfect repairs are regeneration points. If $0 < p \leq 1$, then the system will eventually renew itself; in fact, it will renew itself infinitely many times.

Brown and Proschan developed a number of properties for this imperfect repair model, including the following. If $\lambda(t)$ is the intensity function for the NHPP, then the intensity function for the next renewal (measured in terms of times since the previous renewal) is an NHPP (not a renewal process) with intensity function

$$\lambda_p(t) = p\lambda(t) \quad (7)$$

The survival function for the next renewal is

$$\begin{aligned} S_p(t) &= P(\text{no renewals in } [0, t]) \\ &= \exp\left(-\int_0^t \lambda_p(x) dx\right) \\ &= \exp\left(-p \int_0^t \lambda(x) dx\right) = [S(t)]^p \end{aligned} \quad (8)$$

Many of the monotonicity properties of the intensity (e.g., increasing intensity, increasing intensity on average, etc.) follow from this result and from the fact that $\lambda_p(t) = p\lambda(t)$.

Block *et al.* [3] generalized the Brown–Proschan minimal repair model by allowing for the probability p of perfect repairs to depend on the age t of the system. This model is called the *age-dependent minimal repair model*. Block *et al.* [3] showed that the successive renewal times form a renewal process with cumulative distribution function (cdf)

$$F_p(t) = 1 - \exp\left(-\int_0^t p(x) S^{-1}(x) dF(x)\right) \quad (9)$$

where we have assumed that the failure process is an NHPP with intensity function $\lambda(t)$ and

$$S(x) = \exp\left(-\int_0^x \lambda(u) du\right) \quad (10)$$

and

$$F(x) = 1 - S(x) \quad (11)$$

Block *et al.* showed that if $p(t)$ is increasing, then many of the monotonicity properties of F_p , such as increasing failure rate, increasing failure rate on average, new better than used, and so on, hold whenever they hold for F .

Other Imperfect Repair Models

Berman [4] proposed a class of models where shocks affect a system, and a failure occurs not at every shock, but at every κ th shock, where κ is a positive integer. If the shocks occur according to an NHPP with intensity function $\lambda(t)$ and mean function $\Lambda(t)$, then the joint **probability density function** (pdf) of the first n failure times (or the likelihood function, depending on the perspective) is

$$\begin{aligned} f(t_1, t_2, \dots, t_n) &= \left\{ \prod_{i=1}^n \lambda(t_i) [\Lambda(t_i) - \Lambda(t_{i-1})]^\kappa \right\} \\ &\quad \times [\Gamma(\kappa)]^{-n} \exp(-\Lambda(t_n)) \end{aligned} \quad (12)$$

This is a density as long as $\kappa > 0$, so we may drop the assumption that κ is an integer.

Lakey and Rigdon [5] proposed a special case of Berman's inhomogeneous Poisson process as a model for repairable systems. The model, called the *modulated power law process*, assumes that the intensity function is

$$\lambda(t) = \frac{\beta}{\theta} \left(\frac{t}{\theta}\right)^{\beta-1}, \quad t > 0 \quad (13)$$

that is, the intensity function for the power law process. If $\kappa = 1$, then the MPLP reduces to the power law process since a failure occurs at every shock. If $\beta = 1$, then the process of shocks is a homogeneous Poisson process with intensity $1/\theta$. In this case, the times between failures are sums of i.i.d. exponential random variables, which is a gamma random variable. Thus, $\beta = 1$, reduces the MPLP to the gamma renewal process. Of course if $\kappa = \beta = 1$, then the model reduces to the homogeneous Poisson process.

The MPLP is a compromise between the bad-as-old and good-as-new models (assuming $\kappa > 1$ and $\beta > 1$). A failed/repared unit is in better condition than just before the failure because the counter for the number of shocks is reset to zero, and the number of shocks must again total κ . Also, the system is still not good as new because the intensity of shocks is higher than for a new system, which means that the shocks will tend to arrive at a faster rate.

Black and Rigdon [6] studied inference for the MPLP, including **maximum-likelihood estimation** and **confidence intervals** for the parameters. Maximum likelihood is a bit unstable and several iterations of the Nelder and Meade simplex algorithm (see Buchanon and Turner [7]) are often needed before

starting with the Newton–Raphson method. Asymptotic confidence intervals based on the log of the maximum likelihood estimates (MLEs), obtained using the delta method, have coverage rates closer to the nominal than the asymptotic intervals based on the MLEs themselves.

Another class of imperfect repair models is the reduction-in-intensity model of Doyen and Gaudoin [8]. Their arithmetic reduction-in-intensity (ARI) model assumes

$$\lambda_{T_i^+} = \lambda_{T_i^-} - S(i, T_1, \dots, T_i) \quad (14)$$

where $\lambda_{T_i^+}$ and $\lambda_{T_i^-}$ are, respectively, the limits of the values of the intensity function from the right (just after the failure) and from the left (just before the failure), and S is some function, which in general can depend on the failure number i and all previous failures. This general model has a number of special cases, including the following which may be the simplest. They call this the ARI_∞ model and it assumes

$$\lambda_{T_i^+} = \lambda_{T_i^-} - \rho \lambda_{T_i^-} = (1 - \rho) \lambda_{T_i^-} \quad (15)$$

In other words, after a failure/repair, the intensity is reduced by a factor of $1 - \rho$. Special cases of the ARI_∞ model include the NHPP ($\rho = 0$) and the renewal process ($\rho = 1$). Once again, this model is a compromise between NHPP (bad as old) and renewal process (good as new).

Figure 1 illustrates the idea for a power law intensity with $\rho = \frac{1}{3}$. The dashed line indicates the intensity function if there were no reduction in the intensity function. There are discrete jumps, reductions, in the intensity function at the time of each failure.

Doyen and Gaudoin [8] presented methods for maximum likelihood for the ARI model. Pan and Rigdon [9] described Bayesian methods for the ARI model for repairable systems.

Summary

Any model for which the performance of a failed/rep-
aired unit is between good-as-new (optimistic) and
bad-as-old (pessimistic) can be called an *imperfect repair model*. We have discussed a number of
models for repairable systems that fall into this cat-
egory, including the NHPP (minimal repair), the
Brown–Proschan imperfect repair model, the Block,

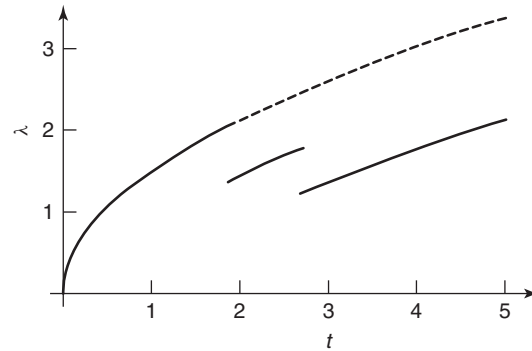


Figure 1 Intensity function for ARI_∞ model

Borges, and Savits age-dependent minimal repair model, the inhomogeneous gamma process (and the special case, the MPLP), and finally the reduction-in-intensity model.

References

- [1] Cox, D.R. & Isham, V. (1980). *Point Processes*, Chapman and Hall, London.
- [2] Brown, M. & Proschan, F. (1983). Imperfect repair, *Journal of Applied Probability* **20**, 851–859.
- [3] Block, H.W., Borges, W.S. & Savits, T.H. (1985). Age-dependent minimal repair, *Journal of Applied Probability* **22**, 370–385.
- [4] Berman, M. (1981). Inhomogeneous and modulated gamma processes, *Biometrika* **68**, 143–152.
- [5] Lakey, M.J. & Rigdon, S.E. (1992). The modulated power law process, in *Proceedings of the 45th Annual Congress, American Society for Quality*, Milwaukee, pp. 559–563.
- [6] Black, S.E. & Rigdon, S.E. (1996). Statistical inference for a modulated power law process, *Journal of Quality Technology* **28**, 81–90.
- [7] Buchanon, J.T. & Turner, P.R. (1992). *Numerical Methods and Analysis*, McGraw-Hill, New York.
- [8] Doyen, L. & Gaudoin, O. (2004). Classes of imperfect repair models based on reduction of failure intensity or virtual age, *Reliability Engineering & System Safety* **45**–56.
- [9] Pan, R. & Rigdon, S.E. (2007). Bayes Inference for General Repairable Systems, Unpublished manuscript, available from S. Rigdon.

Related Articles

General Minimal Repair Models; Imperfect Repair.

STEVEN E. RIGDON

Reliability Allocation

Introduction

Many systems are implemented by using a set of interconnected subsystems. While the architecture of the overall system may often be fixed, individual subsystems may be implemented differently. A designer needs to either achieve the target reliability while minimizing the total cost, or maximize the reliability while using only the available budget. Intuitively, some of the lowest-reliability components may need special attention to raise the overall reliability level. Such an optimization problem may arise while designing a complex software or a computer system. Such problems also arise in mechanical or electrical systems. A number of studies since 1960 have examined such problems [1].

In a nonredundant system, all the subsystems are essential. Often, however, an individual subsystem can be made more reliable by using a more costly implementation. This additional cost may represent wider columns in a building or more thorough testing of software. In redundant implementations, higher reliability can sometimes be achieved by using several copies of a subsystem, such that the system forms a parallel or *k-out-of-n* configuration.

The next section considers the problem formulation, followed by approaches used for setting up an optimization problem. As an example, software reliability allocation is examined in detail with two numerical illustrations. The last section considers reliability allocation in complex systems.

Problem Formulation

We assume that a system has been designed at a higher level as an assembly of appropriately connected subsystems. In general, the functionality of each subsystem can be unique; however, there can be several choices for many of the subsystems providing the same functionality, but differently reliability levels.

Here we consider the problem formulation for a common and widely applicable case. Let there be n subsystems SS_i , $i = 1, \dots, n$, each with reliability R_i and cost C_i . Let the cost C_i be a function of the reliability given by $f_i(R_i)$, it is referred to as the *cost*

function below. Let C_s and R_s represent the total system cost and the overall reliability and R_{ST} be the specified target reliability. If all the subsystems are essential to the system and if their failures are statistically independent, the system can be modeled as a *series system*. The cost minimization problem can be stated as follows:

$$\text{Minimize } C_s = \sum_{i=1}^n C_i = \sum_{i=1}^n f_i(R_i) \quad (1)$$

Subject to

$$R_{ST} \leq R_s$$

$$\text{i.e. } R_{ST} \leq \prod_{i=1}^n R_i \text{ since } R_s = \prod_{i=1}^n R_i \quad (2)$$

Note that equation (1) assumes that the cost of interconnecting the subsystems is negligible. An alternative problem would be to maximize R_s while keeping C_s less than or equal to the allocated cost budget.

The i th subsystem SS_i may have several implementation choices with different reliability values:

1. By extending a continuous attribute (for example, diameter of a column in building or time spent for software testing) the subsystem can be made more reliable.
2. Different vendors may offer their own implementations of SS_i at different costs.
3. It may be possible to use multiple copies of SS_i (for example, double wheels of a truck) to achieve higher reliability. Often the number of copies is constrained between a minimum (often one) and a maximum number because of implementation issues.

Note that in the first case, both the cost and the reliability can be varied continuously, whereas in the other two cases, the choices are discrete. In the first case, we can define a continuous cost function. In the second case too, the market forces may impose a cost function. In the third case, the subsystem may be modeled as a *parallel* or *k-out-of-n system* for reliability evaluation, provided the failures are statistically independent. A number of publications on reliability allocation consider only the third case, where the optimization problem becomes an integer optimization problem. It becomes a 0–1 optimization problem

2 Reliability Allocation

when choices are discrete and a component from a given list of candidates is either used or not used [2].

Approaches for Problem Setup

It is reasonable to assume that the cost function f_i would satisfy the following three conditions [3]:

1. f_i is a positive function
2. f_i is nondecreasing
3. f_i increases at a higher rate for higher values of R_i .

The third condition suggests that it can be very expensive to achieve the reliability value of one. In fact, for software, it has to been shown that under some assumptions, it is infeasible to achieve ultrahigh reliability [4].

In some cases, the cost function can be derived from basic considerations, as we will do below for **software reliability**. In other cases, it may be derived empirically by compiling data for different choices. The cost function is often stated in terms of the reliability, for example, the cost function proposed by Mettas [3] is given by

$$C_i(R_i, p_i, R_{i,\min}, R_{i,\max}) = \exp\left((1 - p_i) \frac{R_i - R_{i,\min}}{R_{i,\max} - R_i}\right) \quad (3)$$

where $R_{i,\min}$ and $R_{i,\max}$ are the minimum and maximum values of R_i and p_i is parameter ranging between zero and one that represents the relative difficulty of increasing a component's reliability. The cost function can also be given in terms of the failure rate as illustrated below for software reliability allocation.

A useful transformation of equation (2) can be obtained by taking the logarithms of both sides of the equation [5–7].

$$\ln(R_{ST}) \leq \sum_{i=1}^n \ln(R_i) \quad (4)$$

The transformation in equation (4) can sometimes reduce the problem to a linear optimization problem. The term $\ln(R_i)$ can also have a well-defined physical significance, in some cases, as the failure rate. When the failure rate of a subsystem SS_i is constant, its

reliability is given by an exponential relationship $R_i(t) = \exp(-\lambda_i t)$, the system failure rate is given by the summation of the subsystem failure rates and hence equation (4) can be restated as

$$\lambda_{ST} \geq \sum_{i=1}^n \lambda_i \quad (5)$$

The failure rate itself is a major reliability attribute. In some cases, such as in software reliability engineering, it is the failure rate that is often specified [8, 9]. The cost function of a subsystem can also be given in terms of its failure rate. If the cost C_i is given by the function $f_i(\lambda_i)$, equation (1) can be restated as

$$\text{Minimize } C_S = \sum_{i=1}^n C_i = \sum_{i=1}^n f_i(\lambda_i) \quad (6)$$

For example, let us consider software reliability. When a software is tested, defects in it are found and removed by debugging. A program tested more thoroughly will have fewer bugs and hence higher reliability. Several models relating software **reliability growth**, with the time spent in testing, have been proposed and validated. These are termed *software reliability growth models* (SRGMs). For the popular *exponential software reliability growth model* [8–10], the failure rate as a function of testing time d is given by

$$\lambda_i(d) = \lambda_{0i} \exp(-\beta_i d) \quad (7)$$

where λ_0 and β_1 are the SRGM parameters. If we assume that the cost is dominated by the testing time, the cost is given by the following function, which satisfies the three conditions mentioned above.

$$d(\lambda_i) = \frac{1}{\beta_i} \ln\left(\frac{\lambda_{0i}}{\lambda_i}\right) \quad (8)$$

Reliability Allocation Approaches for Basic Series and Parallel Systems

The reliability allocation problem for two basic reliability structures *series* and *parallel* can be solved by linearizing the constraints [1, 2, 7]. In a *series* system, the constraint is given in equation (2) above, which can be linearized by rewriting it as

$$\ln(R_{ST}) \leq \sum_{i=1}^n \ln(R_i) \quad (9)$$

which may then be solved relatively easily. An example is given below for software reliability allocation.

In a *parallel system*, functionally identical subsystems are configured such that correct operation of at least one of them assures a correctly functioning system. It is assumed that any overhead in implementing such a system is negligible. In real systems, the overhead involved will result in a lower level of reliability. In a number of studies, the problem assumes that the reliability of a subsystem can be increased by using a functionally identical component in parallel [2, 7]. For **parallel systems**, the constraint is given by

$$R_{ST} \leq 1 - \prod_{i=1}^n (1 - R_i) \quad (10)$$

The constraint can be linearized by using logarithms of the complements of reliability. Therefore, equation (10) can be rewritten as

$$\ln(1 - R_{ST}) \geq \sum_{i=1}^n \ln(1 - R_i) \quad (11)$$

Elegbede *et al.* have recently shown [7] that if the cost function satisfies the three properties given above, the cost is optimal if all the parallel components have the same cost. For software, computer hardware, and mechanical systems, the number of discrete parallel component is likely to be very small.

Reliability Allocation for Software Systems

As a detailed example, we examine the problem of software reliability allocation [6]. Typically, a software consists of sequentially executed blocks, such that only one is under execution at a time. Each block can be independently tested and debugged to reduce the failure rate below a target value. In some cases, the reliability of a block can be further increased by replication. For replication to be effective, each replicated version must be developed independently such that the failures are relatively independent. The impact of replication can be evaluated by assuming statistical independence. However, it has been shown that sometimes the statistical **correlation** can be significant, requiring more complex analysis.

Here, in this example we consider the common case, a nonredundant implementation of software,

divided into n sequential blocks [6]. Let us assume that a block i is under execution for a fraction x_i of the time where $\sum x_i = 1$ [5]. Then the reliability allocation problem can be written as

$$\text{Minimize } C = \sum_{i=1}^n \frac{1}{\beta_i} \ln \left(\frac{\lambda_{0i}}{\lambda_i} \right) \quad (12)$$

$$\text{Subject to } \lambda_{ST} \leq \sum_{i=1}^n x_i \lambda_i \quad (13)$$

Solution: Let us solve the problem posed by equations (12) and (13) by using the Lagrange multiplier approach by finding the minimum of

$$F(\lambda_1, \dots, \lambda_n) = C + \theta(\lambda_1 + \dots + \lambda_n) - \lambda_{ST} \quad (14)$$

where θ is the Lagrange multiplier. The necessary conditions for the minimum to exist are (a) the partial derivatives of the function F are equal to zero, (b) $\theta > 0$, and (c) $x_1 \lambda_1 + x_2 \lambda_2 + \dots + x_n \lambda_n = \lambda_{ST}$ [6]. Equating the partial derivatives to zero and using the third condition, the solutions for the optimal failure rates are found as follows:

$$\lambda_1 = \frac{\frac{\lambda_{ST}}{x_1}}{\sum_{i=1}^n \frac{\beta_i}{x_i}} \quad \lambda_2 = \frac{\beta_1 x_1}{\beta_2 x_2} \lambda_1 \quad \dots \quad \lambda_n = \frac{\beta_1 x_1}{\beta_n x_n} \lambda_1 \quad (15)$$

The testing times of the individual modules are given by equation (12). The optimal values of d_1 and $d_i, i \neq 1$ are given by

$$d_1 = \frac{1}{\beta_1} \ln \left(\frac{\lambda_{10} x_1 \sum_{i=1}^n \frac{\beta_i}{x_i}}{\lambda_{ST}} \right) \quad \text{and} \quad d_i = \frac{1}{\beta_i} \ln \left(\frac{\lambda_{i0} \beta_i x_i}{\lambda_1 \beta_1 x_1} \right) \quad (16)$$

Note that d_i is positive if $\lambda_i \leq \lambda_{i0}$. The testing time for a block must be nonnegative. If any of the testing times in equation (16) is negative, the optimization problem must be solved iteratively [6].

In software reliability engineering, the assumptions involved in formulation of the exponential model imply that the parameter β_i is inversely proportional to the software size [8, 11], when measured in terms of the lines of code. The value of x_i can be

4 Reliability Allocation

Table 1 Input data and optimal values for Example 1

Block	B ₁	B ₂	B ₃	B ₄	B ₅
β_i	7×10^{-3}	3.5×10^{-3}	2.333×10^{-3}	7×10^{-4}	3.5×10^{-4}
λ_{i0}	0.14	0.14	0.14	0.175	0.21
x_i	0.028	0.056	0.083	0.278	0.556
Optimal λ_i	0.06	0.06	0.06	0.06	0.06
Optimal d_i	121.043	242.085	363.128	1.529×10^3	3.579×10^3

reasonably assumed to be proportional to the code size. The values of λ_i and λ_{i0} do not depend on size but depend on the initial defect densities [11]. Therefore, if the exponential model indeed holds, the equation (15) states that the optimal values of the post-test failure rates $\lambda_1, \dots, \lambda_n$ are equal. In addition, if all the initial defect densities are also equal for all the blocks, then the optimal test times for each module is proportional to its size.

Example 1 A software system uses five functional blocks B1–B5. We construct this example assuming sizes 1, 2, 3, 10, and 20 KLOC (thousand lines of code) respectively, and the initial defect densities of 20, 20, 20, 25, and 30 defects per KLOC respectively. Let us assume that measured parameter values are given in the top three rows, which are the inputs to the optimization problem. The solution obtained using equations (15) and (16) are given in the two bottom rows. Let us now minimize the test cost such that the overall failure rate is less than or equal to 0.06 per unit time. Here the time units can be hours of testing time or hours of CPU time used for testing. (See Table 1).

Note that the optimal values of λ_i for the five modules are equal, even though they start with different initial values. This requires a substantial part of the test effort allocated to largest blocks. The total cost in terms of testing is 5.835×10^3 h.

Example 2 This is identical to the previous example with one difference, the numbers in the second

row for λ_{i0} are obtained assuming all five modules have the same defect density value of 20 per KLOC. (See Table 2).

We note that for this example, the optimal testing time is exactly proportional to the software size, block B₅ is tested for 20 times the time used for B₁. The final failure rates are identical for the five blocks.

The above discussion suggests that some preliminary rules may be used for obtaining initial apportionments. Some apportionment rules have been suggested in the literature [5].

Equal reliability apportionment For example, one can test a set of software blocks, such that at the end they all individually have the failure rate equal to the target failure rate for the system.

Complexity-based apportionment For example, the software size itself is a complexity metric. Therefore, the available test time can be apportioned in proportion to the software size.

Impact-based apportionment A block that is executed more frequently, or is more critical in terms of failures, should be assigned more resources.

Reliability Allocation for Complex Systems

In practice, many cases can be complex and may require an iterative approach [1, 2, 7, 12]. Such an approach is also needed if the objective function has multiple objectives and includes both the total cost

Table 2 Input data and optimal values for Example 2

Block	B ₁	B ₂	B ₃	B ₄	B ₅
β_i	7×10^{-3}	3.5×10^{-3}	2.333×10^{-3}	7×10^{-4}	3.5×10^{-4}
λ_{i0}	0.14	0.14	0.14	0.14	0.14
x_i	0.028	0.056	0.083	0.278	0.556
Optimal λ_i	0.06	0.06	0.06	0.06	0.06
Optimal d_i	121.043	242.085	363.128	1.21×10^3	2.421×10^3

and the system reliability [1]. These steps are based on [12]:

1. Design the system using functional subsystems.
2. Perform an initial apportionment of cost or reliability attributes based on suitable apportionment rules or preliminary computation.
3. Predict system reliability.
4. Determine if reallocation is feasible and will enhance the objective function. If so, perform reallocation.
5. Repeat until optimality is achieved.
6. See if this meets the objectives. If not, consider returning to step 1 and revising the design.
7. Finalize the design with recommended reliability allocation and the cost projections.

The optimization methods used in steps 2–5 above can be classified into three approaches [1].

1. **Exact methods** When the problem is not large, exact methods can be desirable. In general, the problem can be a nonlinear optimization problem. In a few cases, the problem can be transformed into a linear problem, as shown in the example above.
2. **Heuristics-based methods** Several heuristics for reliability allocation have been developed. Many of them are based on identifying the variable to which the solution is most sensitive and incrementing its value.
3. **Metaheuristic algorithms** These algorithms are based on artificial reasoning. The best known of them are *genetic algorithms*, *simulated annealing*, and *tabu-search*. These algorithms can be useful when the search space is large and approximate results are sought.

Some general purpose [13] and special purpose [6] software tools have been developed that can simplify setting up and solving the optimal allocation problem. The reliability allocation problem can also be formulated to address other reliability attributes like availability [14] or maintainability.

References

- [1] Kuo, W. & Prasad, V.R. (2000). An annotated overview of system-reliability optimization, *IEEE Transactions on Reliability* **49**, 176–187.
- [2] Majety, S.R.V., Dawande, M. & Rajgopal, J. (1999). Optimal reliability allocation with discrete cost-reliability data for components, *Operations Research* **47**, 899–906.
- [3] Mettas, A. (2000). Reliability allocation and optimization for complex systems, in *Proceedings Annual Reliability and Maintainability Symposium*, Los Angeles, January 2000, pp. 216–221.
- [4] Butler, R.W. & Finelli, G.B. (1993). The infeasibility of quantifying the reliability of life-critical real-time software, *IEEE Transactions on Software Engineering* **19**, 3–12.
- [5] Lakey, P.B. & Neufelder, A.M. (1996). *System and Software Reliability Assurance Notebook*, Rome Laboratory, Rome, pp. 6.1–6.24.
- [6] Lyu, M.R., Rangarajan, S. & van Moorsel, A.P.A. (2002). Optimal allocation of test resources for software reliability growth modeling in software development, *IEEE Transactions on Reliability* **51**, 183–192.
- [7] Elegbede, A.O.C., Chengbin, C., Adjallah, K.H. & Yalaoui, F. (2003). Reliability allocation through cost minimization, *IEEE Transactions on Reliability* **52**, 106–111.
- [8] Musa, J.D., Iannini, A. & Okumoto, K. (1997). *Software Reliability, Measurement, Prediction, Application*, McGraw-Hill.
- [9] Lyu, M.R. (ed) (1995). *Handbook of Software Reliability Engineering*, McGraw-Hill.
- [10] Goel, A.L. & Okumoto, K. (1979). Time-dependent error detection rate model for software and other performance measures, *IEEE Transactions on Reliability* **28**, 206–211.
- [11] Malaiya, Y.K. & Denton, J. (1997). What do the software reliability growth model parameters represent? in *International Symposium on Software Reliability Engineering*, Albuquerque, pp. 124–135.
- [12] NASA Document. (2004). Reliability allocation, Oct 1, 2004, <http://www.foia.msfc.nasa.gov/docs/NAS8-00179/S&MAOI/R-Reliability/QD-R-008.pdf>.
- [13] ReliaSoft. (2003). Reliability importance and optimized reliability allocation (analytical), <http://www.wei-bull.com/SystemRelWeb/blocksimtheory.htm>.
- [14] Gurov, S.V., Utkin, L.V. & Shubinsky, I.B. (1995). Optimal reliability allocation of redundant units and repair facilities by arbitrary failure and repair distributions, *Microelectronics Reliability* **35**(12), 1451–1460.

Related Articles

***k*-out-of-*n* Systems; Parallel, Series, and Series-Parallel Systems.**

YASHWANT K. MALAIYA

Accelerated Life Tests: Classical Methods for Design and Analysis

Introduction

In today's marketplace, new high-technology components are often required to be very reliable. In addition, manufacturers are usually compelled to develop new products in record time, so that they do not lose market opportunities. For instance, items deployed in environments where replacement costs are prohibitive, such as undersea and space, have to be designed in order to operate without failure for years whereas time spent to develop new products cannot exceed half a dozen months.

Reliability of this kind of products is very important; it is vital for manufacturers to satisfy in full the requirements of the customers.

To assure achievement of the desired reliability levels, it is necessary to design for reliability since the earliest stage of the product development and complete at least the first reliability loop. In this loop, failure data collection and analysis play the essential role of feeding information back to designers.

Unfortunately, in the case of high-reliability products, the exiguous amount of failure time data that it is possible to collect carrying out tests at normal use conditions does not allow statistical analyses of practical utility. In the above cases, the only way to obtain the needed amount of data in time is to test the products in conditions more severe than normal, forcing the products to fail more quickly. Over the years, such kind of tests have been indicated as **accelerated life tests** (ALTs). The main purpose of such testing is to estimate product's reliability or long-term performances at normal use condition. This risky but unavoidable extrapolation has to be performed using appropriate statistical models and methods.

Methods for Accelerating Life Tests

There are fundamentally two different ways to compress the life of products:

- *Increasing the usage rate*: This approach can be adopted in the case of products as toasters, refrigerator pumps, rolling bearings, or automotives, whose duration is defined in terms of cycles, mileages, and so on. Time to failure can be compressed by running the product faster, as in the case of the rolling bearings, or reducing the off time, as in the case of refrigerator pumps and toasters;
- *Increasing the stress*: This approach can produce two main effects: increasing the aging rate and/or shortening product life. The first effect can be observed in the case of products and materials whose failure is due to a chemical or physical **degradation** process. The aging process can be accelerated by increasing the level of experimental factors such as temperature and humidity. The second effect can be observed in the case of materials and components that fail when their strength falls below a critical bound that increase with some experimental factors such as voltage or mechanical load. Thus, higher than normal levels of these experimental factors force these products to fail more quickly than they would under normal use conditions. For example [1, p. 479–480], certain types of insulation present a dielectric strength (i.e., breakdown voltage) that decreases with time due to a chemical degradation process. Failure occurs when breakdown voltage falls below a threshold that increases with the applied voltage level.

For many different products these strategies enable the collection of failure data in a very short time.

In the case of products whose life is compressed by increasing the usage rate, usually, the desired estimates are obtained by simply assuming that the number of cycles to failure at higher usage rate remains the same as at normal use conditions.

In the case of overstress testing, extrapolation is usually performed by adopting statistical models that can account for the effect of the stress level on the failure mechanism.

In this article, only fully parametric models for ALT data with a single constant stress are considered. Moreover, only classical **estimation** procedures are discussed. Step-stress testing and/or other estimation methods are described in **Accelerated Life Tests: Bayesian Design, Accelerated Life Tests: Stress**

Step (Classical Methods), Accelerated Life Tests: Designs Comparison within a Bayesian Framework, and Accelerated Life Tests: Nonparametric Approach.

Fully Parametric Models for ALT with a Single Constant Stress

Fully parametric ALT models are special regression models in which the failure time has a distribution that explicitly depends on the value of a vector of, possibly transformed, accelerating variables, $\mathbf{z}$, also called *covariates* or *regressors*.

Most of these models are obtained by combining a standard parametric lifetime model and a life–stress relationship. The first model describes the failure time distribution at each fixed stress level, $\mathbf{z}$, the second accounts for the effect of the **covariate** vector, $\mathbf{z}$, on the parameters of this distribution.

In statistics, regression models are well studied and have a long history; the ones used to model ALT data have the two following peculiarities:

- owing to the nonnegative nature of lifetime, ALT models are usually based on distributions such as exponential, Weibull, or lognormal, rather than on **normal distribution**;
- owing to the purpose of the analyses, these models are used to estimate the product response for values of the covariates that fall outside the test region.

The first characteristic has created the need for appropriate statistical procedures. The second one exposes to unavoidable risks of error, which can be mitigated only if sound physical or chemical justifications can be made about the adequacy of the assumed model at normal use conditions. Nevertheless, many empirical ALT models are successfully used in applications.

In this section, the most commonly used simple failure time regression models are presented, where “simple” is used to indicate that only one accelerating variable is considered. All the models presented in this paper usually apply to a single failure mode; they may be not appropriate for products that fail from more causes. Models for these products are discussed in **Accelerated Life Tests: Analysis with Competing Failure Modes**.

Simple Accelerated Failure Time Model

The accelerated failure time (AFT) model is a widely used regression model. It expresses the effect of the accelerating variable, z , on lifetime, X , as follows:

$$X_z = \frac{X_0}{\varphi(z)} \quad (1)$$

where $\varphi(z)$ is a positive function called *acceleration factor*, X_0 is the lifetime at condition z_0 for which $\varphi(z_0) = 1$, and X_z is the corresponding lifetime at z .

According to this model, the shape of the reliability function is the same at any given level of the accelerating variable, but products fail faster or slower according to the value of z . In particular, the dependence of the p th quantile of $\ln(X)$ on the stress level, z , can be expressed as:

$$\ln[x_p(z)] = \ln[x_p(0)] - \ln[\varphi(z)] \quad (2)$$

and the acceleration factor, $\varphi(z_u, z)$, between the use stress level, z_u , and a reference stress level, z , can be expressed as:

$$\varphi(z_u, z) = \frac{x_p(z_u)}{x_p(z)} = \frac{\varphi(z)}{\varphi(z_u)} \quad (3)$$

Finally, the reliability function of X_z , $R(x; z)$, and that of X_0 , $R_0(x)$, are related by the following equation:

$$R(x; z) = R_0[x \cdot \varphi(z)] \quad (4)$$

One of the most-used forms for the acceleration factor is the log-linear model:

$$\varphi(z) = \exp(-a \cdot z) \quad (5)$$

which includes, as special cases, many widely applied life–stress relationships, such as the Arrhenius, the inverse power law, and the exponential models that are presented and justified in the Appendix.

In the case of failure-causing chemical or physical reactions the AFT model can be considered adequate if the following assumptions are satisfied [1, pp. 476–477, 566; 2, p. 152]:

- there is a single rate-controlling step;
- specimens under study are degrading products that fail when degradation overcomes a critical threshold;

- both the initial degradation level and the degradation threshold for which failure occurs do not depend on stress;
- rate-acceleration parameters are *fixed-effect* parameters. For example, the use of the Arrhenius model suggests that the activation energy does not vary from unit to unit (see the section titled “Simple Log-Location Scale Regression Model with Nonconstant Spread”).

Simple Log-Location Scale Regression Model

The simple log-location scale (LLS) regression model is probably the most commonly used parametric lifetime regression model. It can be formulated as follows:

$$R(x; z) = 1 - \Phi \left[\frac{\ln x - \mu(z)}{\sigma} \right] \quad (6)$$

where $\Phi(\cdot)$ is a fully specified cumulative distribution function (**cdf**).

This model assumes that the scale parameter, σ , is the same at each stress level and that the location parameter $\mu(z)$ is a linear function of a, possibly transformed, accelerating variable:

$$\mu(z) = \gamma_0 + \gamma_1 \cdot z \quad (7)$$

where $\gamma_0 = \mu(0)$ is the location parameter at $z = 0$ and γ_1 is a constant that depends on product geometry, test method, and other factors.

The resulting model can be expressed as follows:

$$R(x; z) = 1 - \Phi \left[\frac{\ln x - (\gamma_0 + \gamma_1 \cdot z)}{\sigma} \right] \quad (8)$$

An alternative formulation of equation (8):

$$R(x; z) = 1 - \Phi_* \left[\left(\frac{x}{\exp(\gamma_0) \cdot \exp(\gamma_1 \cdot z)} \right)^\beta \right] \quad (9)$$

where $\beta = \sigma^{-1}$ and $\Phi_*(\cdot) = \Phi[\ln(\cdot)]$, shows that the LLS model is in the AFT class.

Under model (8) the p th quantile, $\ln[x_p(z)]$, of $\ln(X)$ at any given stress level z , can be obtained as:

$$\ln[x_p(z)] = \sigma \cdot \Phi^{-1}(p) + \gamma_0 + \gamma_1 \cdot z \quad (10)$$

and the acceleration factor (see equation (2)), results in the following log-linear model:

$$\varphi(z) = \exp(-\gamma_1 \cdot z) \quad (11)$$

To obtain a fully parametric model it is necessary to select the shape of $\Phi(\cdot)$, this determines the shape of the lifetime distribution at each given stress level.

The well-known Weibull and lognormal regression models are obtained by choosing a standardized smallest extreme value cdf, $\Phi_{\text{SEV}}(\cdot)$ and a standard normal cdf, $\Phi_{\text{N}}(\cdot)$, respectively:

$$\Phi_{\text{SEV}}(y) = 1 - \exp[-\exp(y)] \quad (12)$$

$$\Phi_{\text{N}}(y) = \frac{1}{\sqrt{2\pi}} \cdot \int_{-\infty}^y \exp[-(u^2/2)] \cdot du \quad (13)$$

The exponential regression model is obtained as a special case of the Weibull regression model for $\sigma = 1$. These regression models have been successfully used to fit data from single constant-stress ALTs in the case of a very broad variety of products [3, Section 2; 4; 5, p. 28; 6].

Simple Log-Location Scale Regression Model with Nonconstant Spread

The LLS regression model is not always adequate for applications. Theoretical and experimental studies have shown that the scale parameter of the log failure time often depends on stress [7, 8]. A simple model that can account for this kind of dependence is an LLS regression model in which both location and scale parameters depend on stress:

$$R(x; z) = 1 - \Phi \left[\frac{\ln x - \mu(z)}{\sigma(z)} \right] \quad (14)$$

where $\Phi[\cdot]$ is a fully specified cdf. Many authors assume for $\mu(z)$ and $\sigma(z)$ [3, p. 106; 7; 9] the simple forms:

$$\mu(z) = \gamma_0 + \gamma_1 \cdot z \quad (15)$$

$$\ln[\sigma(z)] = \nu_0 + \nu_1 \cdot z \quad (16)$$

A justification for both constant and nonconstant spread models is given in [10], where, assuming that log lifetime, $\ln X(z)$, scales according to the Arrhenius model (see Appendix):

$$\ln X(z) = \ln A + \frac{E_a}{k} \cdot z \quad (17)$$

$$z = \frac{1}{T} \quad (18)$$

it is remarked that:

- the constant spread model is obtained in the case in which A is a random variable (i.e. it varies among test units) and E_a is a constant;
- the nonconstant spread model is obtained in the case in which both A and E_a are random variables.

A potential pitfall of this model is discussed in [1, p. 442; 3, p. 107].

Simple Proportional Hazard Model

The simple proportional hazard (PH) model [11] is a very general regression model. It assumes that the effect of a possibly transformed accelerating variable, z , on the lifetime distribution can be expressed through the following product:

$$h(x; \mathbf{z}) = h_0(x) \cdot \psi(z) \quad (19)$$

where $h(x; z)$ is the hazard function for units tested at z and $h_0(x)$ is the function obtained at condition z_0 for which $\psi(z_0) = 1$.

Owing to this assumption, also called the *proportionality assumption*, the ratio of the hazard functions $h(x; z_1)$ and $h(x; z_2)$, obtained in correspondence of two different stress levels, z_1 and z_2 , respectively, are proportional to one another:

$$\frac{h(x; z_1)}{h(x; z_2)} = \frac{h_0(x) \cdot \psi(z_1)}{h_0(x) \cdot \psi(z_2)} = \frac{\psi(z_1)}{\psi(z_2)} \quad (20)$$

The reliability function at z can be obtained from the reliability function $R_0(t)$ as follows:

$$R(x; z) = R_0(x)^{\psi(z)} \quad (21)$$

where $R_0(t) = \exp\left[-\int_0^t h_0(u) \cdot du\right]$.

In the case of a fully parametric approach both $h_0(x)$ and $\psi(z)$ have to be considered specified function of unknown parameters. Most of the PH models used in ALT with a single constant stress are obtained by combining baseline hazard function of nonnegative random variables such as exponential, Weibull, or lognormal with the log-linear relation:

$$\psi(z) = \exp(b \cdot z) \quad (22)$$

Assumptions of the PH model are very different from those of the AFT model. In fact, the Weibull regression model is the only model in both classes. The PH model is largely applied in the

biomedical field [12, p. 276]. In ALTs physical and chemical considerations generally suggest the use of models such as (6) or (14) are usually adopted. Therefore the PH model, in reliability, has been prevalently adopted in the analysis of field data to model the effect of environmental factors and/or of other sources of nonhomogeneity [1, p. 458; 13, 14].

Estimation Procedures for the Simple Log-Location Scale Regression Model

ALT data usually consist of a certain number of, complete or right-censored observations, $n = \sum_{i=1}^m n_i$, taken at m different stress levels, $z_1, z_2, \dots, z_m$, with n_i observations at z_i .

These data constitute the basis for performing point and interval estimates and for checking validity of the selected regression model. Statistical procedures employed to make inference from ALT data are slightly different from those used in “classical” regression analysis for two main reasons:

- the residual distribution is nonnormal and
- observations are usually censored.

So, analysis of these data is usually performed by adopting a combination of graphical and analytical methods. Moreover the analytical approach is prevalently based on the maximum-likelihood (ML) method rather than on other classical approaches, such as least-square estimators, which in the case of ALT data are often unsatisfactory [1, p. 154; 3, pp. 240–241; 12, pp. 169, 299, 306].

Application of graphical, ML and least-square methods to models presented in the section titled “Fully Parametric Models for ALT with a Single Constant Stress” is largely discussed in literature [1, Chapters 17–19; 3, Sections 3–5; 7; 9; 12, Chapters 6, 7; 15, Chapter 9; 16, Chapter 5; 17, Chapter 7]. Herein, attention will be focused on the application of ML methods to the model (8).

Point Estimation

The likelihood function for n s -independent exact and right-censored ALT data, obtained at m different

stress levels, $z_1, z_2, \dots, z_m$, is

$$L(\gamma_0, \gamma_1, \sigma) = \prod_{i=1}^m \prod_{j=1}^{n_i} \left[\frac{1}{\sigma \cdot x_{i,j}} \phi \left(\frac{\ln x_{i,j} - \gamma_0 - \gamma_1 \cdot z_i}{\sigma} \right) \right]^{\delta_{i,j}} \times \left[1 - \Phi \left(\frac{\ln x_{i,j} - \gamma_0 - \gamma_1 \cdot z_i}{\sigma} \right) \right]^{1-\delta_{i,j}} \quad (23)$$

where $\phi(y) = \frac{d}{dy} \Phi(y)$, n_i is the number of units allocated at stress z_i , $x_{i,j}$ is the duration of the j th test run at stress z_i , and $\delta_{i,j}$ is an indicator that assumes the value 1 if the $x_{i,j}$ is uncensored and 0 if it is censored. Given the data, the maximum-likelihood estimates (MLEs) of γ_0 , γ_1 and σ are the values, $\hat{\gamma}_0$, $\hat{\gamma}_1$ and $\hat{\sigma}$, that maximize equation (23) or the corresponding log-likelihood, $l(\gamma_0, \gamma_1, \sigma)$:

$$l(\gamma_0, \gamma_1, \sigma) = \ln[L(\gamma_0, \gamma_1, \sigma)] \quad (24)$$

An alternative formulation for the likelihood function that uses the equation (9) is the following:

$$L(\gamma_0, \gamma_1, \beta) = \prod_{i=1}^m \prod_{j=1}^{n_i} \left\{ \frac{\beta \cdot x_{i,j}^{\beta-1}}{[\exp(\gamma_0 + \gamma_1 \cdot z_i)]^\beta} \times \phi_* \left[\left(\frac{x_{i,j}}{\exp(\gamma_0 + \gamma_1 \cdot z_i)} \right)^\beta \right] \right\}^{\delta_{i,j}} \times \left\{ 1 - \Phi_* \left[\left(\frac{x_{i,j}}{\exp(\gamma_0 + \gamma_1 \cdot z_i)} \right)^\beta \right] \right\}^{1-\delta_{i,j}} \quad (25)$$

where $\phi_*(y) = \frac{d}{dy} \Phi_*(y)$.

As a rule, to maximize the likelihood it is necessary to use numerical procedures. The Newton–Raphson method usually gives good results. Nevertheless, when censoring is heavy and/or the amount of data is very limited, this method often fails to converge unless the starting values are very close to the MLEs. Valid suggestions to implement solving algorithms are given in [18–20]. In [9] an efficient procedure is presented for computing the MLEs of the Weibull-generalized log-linear regression model (both with constant and nonconstant shape parameters).

It can be worth to remark that when $\Phi(\cdot) = \Phi_N(\cdot)$ and observations are all uncensored, closed-form MLEs can be readily obtained using simple linear regression theory.

Example 1 Fit the Weibull–Arrhenius regression model to the temperature ALT data displayed in Table 1 (data from [17, p. 270]) and estimate the Weibull scale parameter, α_u (i.e. the 63.2th quantile), at the use temperature of 308.15 K (35 °C).

Solution The Weibull–Arrhenius log-likelihood function for data in Table 1 can be rewritten as follows:

$$l(\gamma_0, \gamma_1, \beta) = \ln[L(\gamma_0, \gamma_1, \beta)] = r_T \cdot \ln(\beta) - r_T \cdot \beta \times \gamma_0 - \beta \cdot \gamma_1 \cdot \sum_{i=1}^m r_i \cdot z_i + (\beta - 1) \times \sum_{i=1}^m \sum_{j=1}^{r_i} \ln x_{i,j} - \exp(-\beta \cdot \gamma_0) \times \sum_{i=1}^m \exp(-\beta \cdot \gamma_1 \cdot z_i) \cdot \sum_{j=1}^{n_i} x_{i,j}^\beta \quad (26)$$

with

$$r_1 = 8, r_2 = 7, r_3 = 10, r_T = \sum_{i=1}^3 r_i = 25, n_1 = 12,$$

Table 1 Temperature-accelerated life test data of a certain type of electric module^(a)

Temperature (absolute kelvin, K)	$T_1 = 373.15$	$T_2 = 393.15$	$T_3 = 423.15$
	1138	1121	420
	1944	1572	650
	2764	2329	703
	2846	2573	838
	3246	2702	1086
Life data	3803	3702	1125
(h)	5046	4277	1387
	5139	(4500)	1673
	(5500)		1896
	(5500)		2037
	(5500)		
	(5500)		

^(a)Parentheses indicate time-censored data

$$n_2 = 8, n_3 = 10, n_T = \sum_{i=1}^3 n_i = 30, m = 3,$$

$$z_1 = \frac{1}{373.15}, z_2 = \frac{1}{393.15}, \text{ and } z_3 = \frac{1}{423.15}$$

where r_i (n_i) indicates the number of modules failed (tested) at temperature T_i , z_i is the transformed stress, $x_{i,1}, x_{i,2}, \dots, x_{i,r_i}$ are failure times and $x_{i,r_i+1}, \dots, x_{i,n_i}$ are censored times.

To maximize the likelihood it is necessary to use numerical procedures. MLEs, $\hat{\gamma}_0$, $\hat{\gamma}_1$, and $\hat{\beta}$, can be obtained by directly maximizing the log-likelihood function (26).

This maximization can be simplified by adopting the following more effective approach [19]:

- obtain $\hat{\gamma}_1$ and $\hat{\beta}$ by maximizing the profile log-likelihood:

$$l^*(\gamma_1, \beta)$$

$$= -r_T \cdot \ln \left[\left(\sum_{i=1}^m \exp(-\beta \cdot \gamma_1 \cdot z_i) \cdot \sum_{j=1}^{n_i} x_{i,j}^\beta \right) / \beta \right]$$

$$- \beta \cdot \gamma_1 \cdot \sum_{i=1}^m r_i \cdot z_i + (\beta - 1) \cdot \sum_{i=1}^m \sum_{j=1}^{r_i} \ln x_{i,j}$$

$$+ r_T \cdot [\ln(r_T) - 1] \quad (27)$$

- calculate $\hat{\gamma}_0$ as

$$\hat{\gamma}_0 = \ln \left[\frac{1}{r_T} \sum_{i=1}^m \exp(-\hat{\beta} \cdot \hat{\gamma}_1 \cdot z_i) \sum_{j=1}^{n_i} x_{i,j}^{\hat{\beta}} \right]^{\frac{1}{\hat{\beta}}} \quad (28)$$

Resulting MLEs of γ_0 , γ_1 , and β are the following:

$$\hat{\gamma}_0 = -3.156, \hat{\gamma}_1 = 4930, \text{ and } \hat{\beta} = 2.27.$$

The MLE, $\hat{\alpha}_u$, of the Weibull scale parameter at 308.15 K, α_u , can be obtained as

$$\hat{\alpha}_u = \exp \left(\hat{\gamma}_0 + \hat{\gamma}_1 \cdot \frac{1}{308.15} \right) = 65532.97 \text{ h} \quad (29)$$

Normal Approximation Confidence Intervals

Approximate $(1 - \alpha)\%$ **confidence intervals** for the parameters γ_0 , γ_1 , and σ can be computed by

adopting the following relationships:

$$\left[\hat{\theta}_i + \Phi_N^{-1}(\alpha/2) \cdot \sqrt{\widehat{\text{Asvar}}(\hat{\theta}_i)}, \hat{\theta}_i + \Phi_N^{-1}(1 - \alpha/2) \right. \\ \left. \times \sqrt{\widehat{\text{Asvar}}(\hat{\theta}_i)} \right]; \quad \theta_1 = \gamma_0, \theta_2 = \gamma_1, \theta_3 = \sigma \quad (30)$$

that are based on the assumption that the quantities

$$\frac{\hat{\theta}_i - \theta_i}{\sqrt{\widehat{\text{Asvar}}(\hat{\theta}_i)}}; \quad \theta_1 = \gamma_0, \theta_2 = \gamma_1, \theta_3 = \sigma \quad (31)$$

can be approximated by standard normal, $N(0, 1)$, random variables.

These intervals make use of the elements of the local estimate, $\hat{\Sigma}_{\hat{\gamma}_0, \hat{\gamma}_1, \hat{\sigma}}$, of the asymptotic variance–**covariance** matrix, $\Sigma_{\hat{\gamma}_0, \hat{\gamma}_1, \hat{\sigma}}$, of $\hat{\gamma}_0$, $\hat{\gamma}_1$, and $\hat{\sigma}$:

$$\hat{\Sigma}_{\hat{\gamma}_0, \hat{\gamma}_1, \hat{\sigma}}$$

$$= \begin{bmatrix} \widehat{\text{Asvar}}(\hat{\gamma}_0) & \widehat{\text{Ascov}}(\hat{\gamma}_0, \hat{\gamma}_1) & \widehat{\text{Ascov}}(\hat{\gamma}_0, \hat{\sigma}) \\ \text{Symmetric} & \widehat{\text{Asvar}}(\hat{\gamma}_1) & \widehat{\text{Ascov}}(\hat{\gamma}_1, \hat{\sigma}) \\ & & \widehat{\text{Asvar}}(\hat{\sigma}) \end{bmatrix} = I_0^{-1}$$

$$= \begin{bmatrix} -\frac{\partial^2 \ln L}{\partial \gamma_0^2} & -\frac{\partial^2 \ln L}{\partial \gamma_0 \cdot \partial \gamma_1} & -\frac{\partial^2 \ln L}{\partial \gamma_0 \cdot \partial \sigma} \\ & -\frac{\partial^2 \ln L}{\partial \gamma_1^2} & -\frac{\partial^2 \ln L}{\partial \gamma_1 \cdot \partial \sigma} \\ \text{Symmetric} & & -\frac{\partial^2 \ln L}{\partial \sigma^2} \end{bmatrix}_{(\hat{\gamma}_0, \hat{\gamma}_1, \hat{\sigma})}^{-1} \quad (32)$$

where I_0 is the observed information matrix, evaluated at $(\hat{\gamma}_0, \hat{\gamma}_1, \hat{\sigma})$, and I_0^{-1} is its inverse.

A better estimate, $\hat{\Sigma}_{\hat{\gamma}_0, \hat{\gamma}_1, \hat{\sigma}}$ (i.e., the so-called **MLE**), of $\Sigma_{\hat{\gamma}_0, \hat{\gamma}_1, \hat{\sigma}}$ can be obtained evaluating at $(\hat{\gamma}_0, \hat{\gamma}_1, \hat{\sigma})$ the inverse of the information matrix I :

$$I =$$

$$\begin{bmatrix} -E \left(\frac{\partial^2 \ln L}{\partial \gamma_0^2} \right) & -E \left(\frac{\partial^2 \ln L}{\partial \gamma_0 \cdot \partial \gamma_1} \right) & -E \left(\frac{\partial^2 \ln L}{\partial \gamma_0 \cdot \partial \sigma} \right) \\ & -E \left(\frac{\partial^2 \ln L}{\partial \gamma_1^2} \right) & -E \left(\frac{\partial^2 \ln L}{\partial \gamma_1 \cdot \partial \sigma} \right) \\ \text{Symmetric} & & -E \left(\frac{\partial^2 \ln L}{\partial \sigma^2} \right) \end{bmatrix} \quad (33)$$

For most models $\hat{\Sigma}_{\hat{\gamma}_0, \hat{\gamma}_1, \hat{\sigma}}$ cannot be expressed in closed form. Moreover the use of $\hat{\Sigma}_{\hat{\gamma}_0, \hat{\gamma}_1, \hat{\sigma}}$ does not give clear theoretical advantages with respect to $\hat{\Sigma}_{\hat{\gamma}_0, \hat{\gamma}_1, \hat{\sigma}}$ [3, p. 369; 12, p. 301].

This normal approximation is usually poor for small samples. When θ is a non negative parameter a *positive* confidence interval can be obtained by adopting the following approximation [1, p. 163; 17, p. 277]:

$$\frac{\ln(\hat{\theta}) - \ln(\theta)}{\sqrt{\widehat{\text{Asvar}}[\ln(\hat{\theta})]}} \sim N(0, 1) \quad (34)$$

the resulting $(1 - \alpha)\%$ two-sided approximate confidence interval is:

$$\left\{ \hat{\theta} \cdot \exp \left[\Phi_N^{-1}(\alpha/2) \cdot \frac{\sqrt{\widehat{\text{Asvar}}(\hat{\theta})}}{\hat{\theta}} \right], \right. \\ \left. \times \hat{\theta} \cdot \exp \left[\Phi_N^{-1}(1 - \alpha/2) \cdot \frac{\sqrt{\widehat{\text{Asvar}}(\hat{\theta})}}{\hat{\theta}} \right] \right\}$$

where the approximate identity

$$\widehat{\text{Asvar}}[\ln(\hat{\theta})] = \frac{\widehat{\text{Asvar}}(\hat{\theta})}{\hat{\theta}^2} \quad (35)$$

has been obtained adopting the δ method [1, p. 619].

More accurate confidence intervals can be obtained by adopting the likelihood ratio or other nonstandard methods [1, p. 438; 3, Section 5.8; 12, Chapter 6; 21]. Exact confidence intervals have been developed in the case of singly failure censored data for the exponential [22] and the Weibull regression models [23].

Checking Model Assumptions Using Graphical Methods

Adequacy of the simple LLS model can be assessed by analyzing the standardized **residuals**. For the model (8) these can be expressed as follows [1, p. 443]:

$$\hat{e}_{i,j} = \exp \left[\frac{\ln x_{i,j} - (\hat{\gamma}_0 + \hat{\gamma}_1 \cdot z_i)}{\hat{\sigma}} \right] \\ = \left(\frac{x_{i,j}}{\exp(\hat{\gamma}_0) \cdot \exp(\hat{\gamma}_1 \cdot z_i)} \right)^{\hat{\beta}} \quad (36)$$

where $\hat{\gamma}_0$, $\hat{\gamma}_1$, $\hat{\sigma}$ and $\hat{\beta} = 1/\hat{\sigma}$ are MLEs. Residuals, $e_{i,j}$, is censored if corresponding observation, $x_{i,j}$, is censored. In the case of censored data, the single sample of residuals is usually multiply censored.

The overall adequacy of the fitted model can be tested making a **probability plot** of the residuals. For example, a Weibull probability chart can be used with $\Phi(\cdot) = \Phi_{\text{SEV}}(\cdot)$ and a lognormal probability chart can be used for $\Phi(\cdot) = \Phi_N(\cdot)$. If the departure from linearity is not “strong” the model can be considered adequate.

A valid alternative to use residuals (36) consists in plotting on a probability chart the following *translated* data [3, p. 251]:

$$\hat{y}_{i,j} = \frac{x_{i,j}}{\exp[\hat{\gamma}_1 \cdot (z_i - z_u)]}; \quad i = 1, 2, \dots, m; \\ j = 1, 2, \dots, n_i \quad (37)$$

which can be thought of as their *equivalent* at the use stress level z_u .

In this case, the ML fit of the cdf at use stress level z_u can be plotted on the probability chart along with the *equivalent* data $\hat{y}_{i,j}$, to investigate potential deficiency of the overall model. This plot can also be used to perform graphical estimates of characteristics of the lifetime distribution at use condition. Values, $\hat{y}_{i,j}$, corresponding to censored observations $x_{i,j}$ have to be considered censored. As in the case of residuals (36), in the presence of censoring, the simple sample of *equivalent* data is usually multiply censored.

Specific causes of departure from the fitted model can be identified by making a multiple probability plot of the data, $x_{i,j}$, obtained at each stress level. The fitted model can be considered adequate if:

- data obtained at different stress levels depict approximately parallel straight lines on a “proper” probability chart;
- estimates of the characteristic life (e.g. the scale parameter), obtained at different stress levels, z_i , depict a straight line (see equation (10)) on a paper which has a log scale for characteristic life and a linear scale for the, possibly transformed, stress, z , where, for example, it is $z = 1/T$ for the Arrhenius model and $z = \ln V$ for the inverse power law model (see Appendix).

Finally, the ML fits of the “individual” cdfs obtained at each of the stress levels, z_i , can be plotted on the probability chart, along with the data, $x_{i,j}$, to identify potential deficiencies of the overall model at specific stress levels. Details about the application of

8 Accelerated Life Tests: Classical Methods

graphical methods in ALT are given in [1, Chapters 18, 19; 3, Section 3].

Example 2 Check adequacy of the Weibull–Arrhenius regression model for data of Example 1.

Solution The overall adequacy of model can be checked by plotting on a Weibull probability chart the single sample of equivalent data displayed in Table 2.

Equivalent data have been calculated adopting equation (37) with:

$$z_1 = \frac{1}{373.15}, \quad z_2 = \frac{1}{393.15}, \quad z_3 = \frac{1}{423.15},$$

$$z_u = \frac{1}{308.15}, \quad \text{and } \hat{\gamma}_1 = 4930$$

where $\hat{\gamma}_1$ is the MLE of γ_1 (see Example 1).

The plotting position, F_i , for the i th-ordered failure has been obtained *via* the Herd–Johnson method [24, p. 149]:

$$F_i = \prod_{j=1}^i \left(\frac{r_j}{r_j + 1} \right)^{\delta_j} \quad (38)$$

where r_j (see Table 2) is the reverse rank of the j th ordered observation, and δ_j is an indicator that assumes the value 1 if the j th observation is uncensored and 0 if it is censored. Only failure data are plotted on the probability chart. The Weibull probability plot of the equivalent data along with the ML fit of the cdf (see Example 1) at use stress level, z_u , are displayed in Figure 1.

Table 2 Equivalent data at $T = 308.15 \text{ K}^{(a)}$

r_i	Time (h)	r_i	Time (h)	r_i	Time (h)	r_i	Time (h)	r_i	Time (h)
30	13610.96	24	33762.09	18	50679.73	12	61464.59	6	80347.05
29	20170.81	23	34039.35	17	52155.94	11	(65782.30)	5	80556.62
28	23251.05	22	34207.19	16	54028.95	10	(65782.30)	4	91056.79
28	24393.29	21	38823.52	15	55989.24	9	(65782.30)	3	93068.79
26	31216.73	20	40245.56	14	58796.32	8	(65782.30)	2	97828.42
25	33058.60	19	45485.47	13	60352.27	7	66611.69	1	(97921.34)

^(a)Parentheses indicate time-censored data

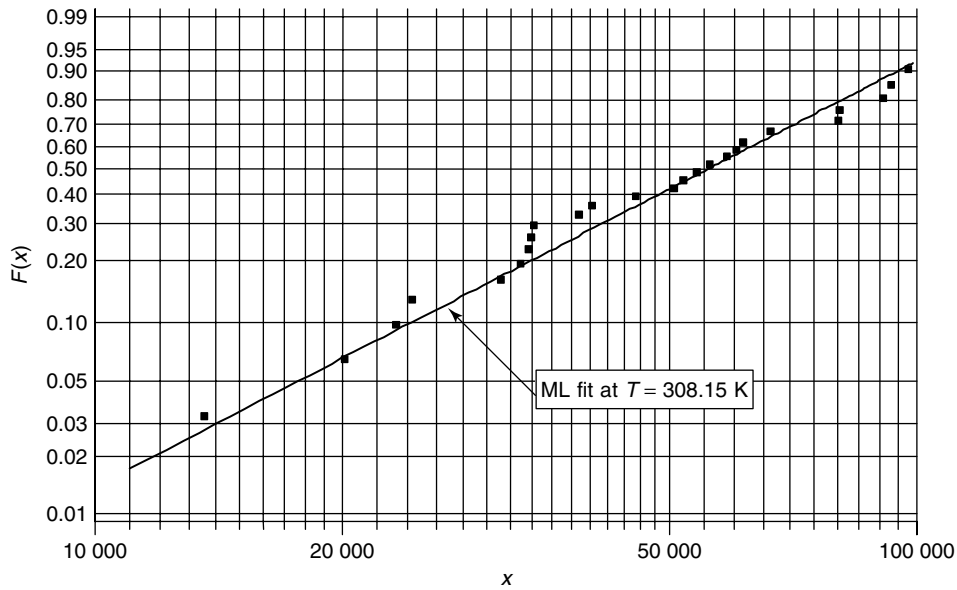


Figure 1 Weibull plot of equivalent data at use temperature, $T = 308.15 \text{ K}$

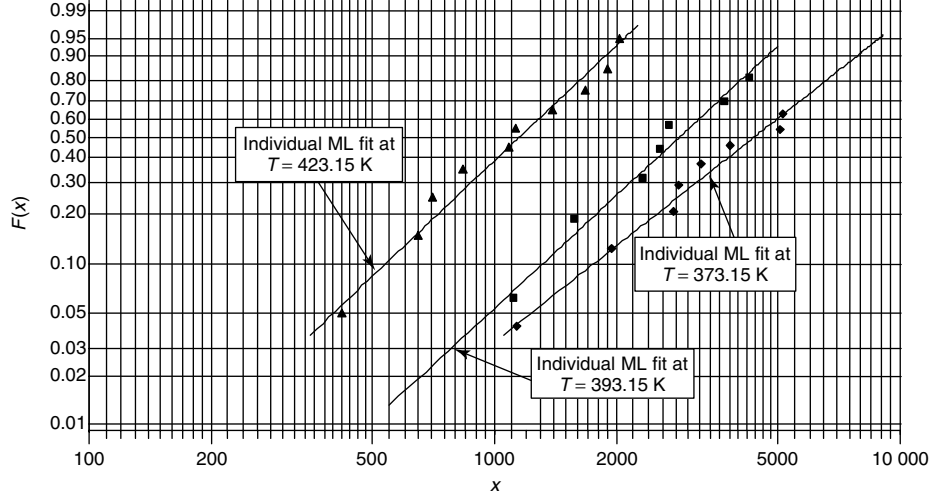


Figure 2 Weibull multiple probability plot with individual Weibull ML fits at each temperature

Departure from linearity in Figure 1 is not “strong”, so the model can be considered adequate.

Specific causes of departures from the fitted model can be investigated making a multiple probability plot of individual samples, $x_{i,j}$, obtained at each stress level.

The plotting position, $F_{i,j}$, for the j th-ordered failure at temperature T_i has been calculated as:

$$F_{i,j} = \frac{j - 0.5}{n} \quad (39)$$

Censored data are not reported on the plot. Figure 2 shows the multiple probability plot along with the individual ML fits of the Weibull cdf at each temperature.

Individual MLEs of $\hat{\alpha}_i$ and $\hat{\beta}_i$ at each temperature have been obtained by solving, *via* numerical methods, the following system of equations:

$$\frac{\sum_{j=1}^{n_i} x_{i,j}^{\hat{\beta}_i} \ln(x_{i,j})}{\sum_{j=1}^{n_i} x_{i,j}^{\hat{\beta}_i}} - \frac{1}{\hat{\beta}_i} - \frac{1}{r_i} \cdot \sum_{j=1}^{r_i} \ln(x_{i,j}) = 0 \quad (40)$$

$$\hat{\alpha}_i = \left[\frac{1}{r_i} \sum_{j=1}^{n_i} x_{i,j}^{\hat{\beta}_i} \right]^{1/\hat{\beta}_i} \quad (41)$$

The resulting MLEs are reported in Table 3.

Table 3 Individual ML estimates of the Weibull parameters at each temperature

Individual MLE	$T_1 = 373.15 \text{ K}$	$T_2 = 393.15 \text{ K}$	$T_3 = 423.15 \text{ K}$
$\hat{\alpha}_i$	5215.54	3326.84	1337.54
$\hat{\beta}_i$	2.05	2.41	2.45

Data relative to different stress levels depict approximately straight lines on the Weibull probability chart, hence the lifetime at each different stress level can be considered Weibull distributed. Moreover these straight lines are roughly parallel. So it is possible to assume that these Weibull models have the same shape parameter.

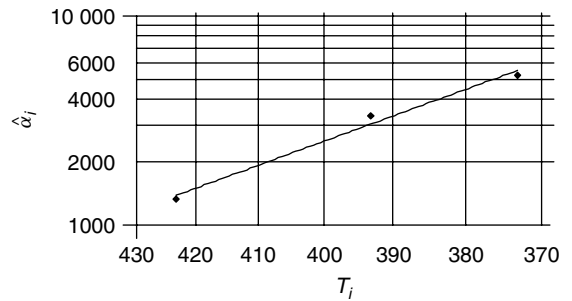


Figure 3 Arrhenius plot of individual MLE of α_i versus T_i

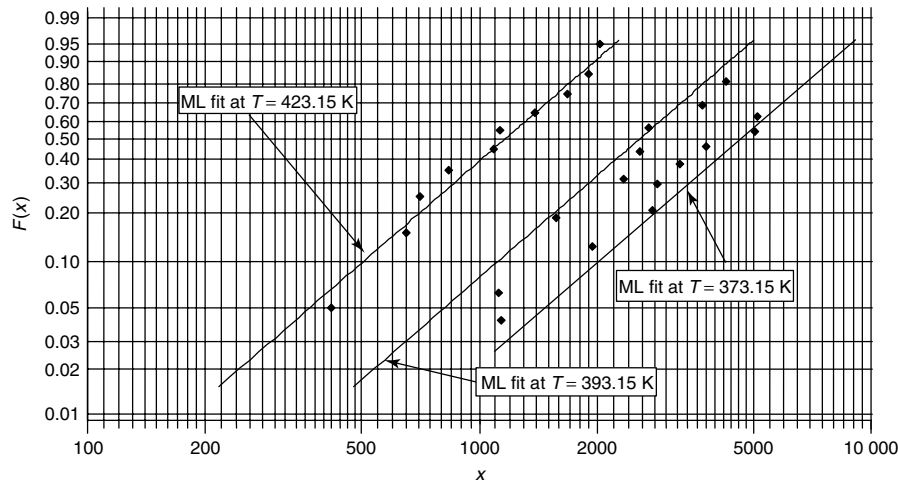


Figure 4 Weibull multiple probability plot with ML fits of Weibull–Arrhenius model at each temperature

Validity of the Arrhenius relationship can be assessed by plotting each individual MLE (see Table 3) of the scale parameter, α_i , against its temperature, T_i , on a paper that has a log scale for α_i and a reciprocal scale for T_i . The resulting plot is displayed in Figure 3. The straight line on this plot is fitted by eye. Departure from linearity is not strong, so the Arrhenius model can be considered adequate.

Finally, Figure 4 shows a multiple probability plot of original data along with ML fits at each temperature. This plot shows no strong evidence of **lack of fit** of the overall model at any stress level.

Planning Test for the Simple Log-Location Scale Regression Model

Planning ALTs for the model (8) consists in choosing: a set of constant values of the accelerating stress; the number of units to test at each stress level; an optimization criterion and a type of censoring, if any. The purpose of developing ALT plans is that of obtaining good estimates respecting cost and time constraints. Traditional test plans, such as those used for analysis of variance (ANOVA) and classical regression analysis, normally consider equal number of units at each stress level. Moreover test levels are usually equally spaced over the test region. These kind of plans that perform well when the aim is to estimate product response at levels that fall within the experimental

region are not appropriate for ALTs. In fact, estimates from ALT data are more accurate when more units are tested at low rather than at high stress levels [1, p. 537]. The most-used optimization criteria for developing an ALT plan is the minimization of the asymptotic variance of the estimator of a quantile of interest, $x_p(z_u)$, at normal use conditions, z_u . As an alternative, when the interest is less focused, a useful *general purpose* criterion is to minimize the determinant $|\Sigma_{\hat{\gamma}_0, \hat{\gamma}_1, \hat{\sigma}}|$ of the asymptotic variance–covariance matrix, $\Sigma_{\hat{\gamma}_0, \hat{\gamma}_1, \hat{\sigma}}$. Most plans consider type I censoring. Some consider complete samples and type II censoring. Test plans for complete samples are given in [3, Section 6.1]. In the next section the basic principles for developing optimum censored ALT plans for the simple LLS model are briefly described. Details about theory for optimum type I censored ALT plans are given in [25, 26]. The main problem with this kind of plan is that it uses only two stress levels. This does not allow the validity of the assumed life–stress relationship to be checked (see Example 2). Moreover, they do not allow estimation of the model parameters when there are no failures at one of these two stress levels. Hence, these plans are often not of practical utility. Nevertheless, optimum plans have minimum variance and can be used as a reference to evaluate performances of alternative solutions. To overcome limitations of optimum plans, various authors have developed *compromise* test plans that are robust and have good performances. Compromise plans and plans with more than one

accelerating variable may be found, for example, in [1, Chapter 20; 3, Section 6; 17, Chapter 7]. Plans for the simple LLS model with nonconstant spread are described in [7]. An updated bibliography of ALT plans is provided in [27, 28].

Theoretical Derivation of an Optimum ALT Plan

The optimum experimental plan considered in this section consists of

- two stress levels: the low level, z_1 , and the high level, z_2 , both higher than the normal one, z_u ;
- $n \cdot \pi_1$ units, randomly chosen among n , to test at level z_1 and the remaining $n \cdot \pi_2 = n \cdot (1 - \pi_1)$ to test at level z_2 ;
- a common censoring time, η .

It is assumed that

- censoring time and sample size are prefixed to respect cost and time constraints.
- the high stress level z_2 is given. The accuracy of estimates increases with z . Thus, in practice, z_2 is chosen to be the higher value of z under which the assumed stress–life relation holds.

The aim of the plan considered in this section is to determine the value of the design factors z_1 and π_1 that minimize the asymptotic variance, $\text{Asvar}[\ln(\widehat{x_p(z_u)})]$, of the MLE, $\ln(\widehat{x_p(z_u)})$, of the quantile $\ln(x_p(z_u))$.

We have:

$$\ln[\widehat{x_p(z_u)}] = \hat{\gamma}_0 + \hat{\sigma} \cdot \Phi^{-1}(p) + \hat{\gamma}_1 \cdot z_u \quad (42)$$

given the asymptotic variance–covariance matrix:

$$\begin{aligned} \sum_{\hat{\gamma}_0, \hat{\gamma}_1, \hat{\sigma}} &= \begin{bmatrix} \text{Asvar}(\hat{\gamma}_0) & \text{Ascov}(\hat{\gamma}_0, \hat{\gamma}_1) & \text{Ascov}(\hat{\gamma}_0, \hat{\sigma}) \\ & \text{Asvar}(\hat{\gamma}_1) & \text{Ascov}(\hat{\gamma}_1, \hat{\sigma}) \\ \text{Symmetric} & & \text{Asvar}(\hat{\sigma}) \end{bmatrix} \\ &= I^{-1} \end{aligned} \quad (43)$$

the following expression for the $\text{Asvar}[\ln(\widehat{x_p(z_u)})]$ is readily obtained:

$$\begin{aligned} \text{Asvar}[\ln(\widehat{x_p(z_u)})] &= \text{Asvar}(\hat{\gamma}_0) + [\Phi^{-1}(p)]^2 \\ &\quad \times \text{Asvar}(\hat{\sigma}) + z_u^2 \cdot \text{Asvar}(\hat{\gamma}_1) \\ &\quad + 2 \cdot \Phi^{-1}(p) \cdot \text{Ascov}(\hat{\gamma}_0, \hat{\sigma}) \\ &\quad + 2 \cdot z_u \cdot \text{Ascov}(\hat{\gamma}_0, \hat{\gamma}_1) + 2 \\ &\quad \times \Phi^{-1}(p) \cdot z_u \text{Ascov}(\hat{\gamma}_1, \hat{\sigma}) \end{aligned} \quad (44)$$

Given that $\sum_{\hat{\gamma}_0, \hat{\gamma}_1, \hat{\sigma}}$ cannot be obtained in closed form, the values of z_1 and π_1 that minimize the $\text{Asvar}[\ln(\widehat{x_p(z_u)})]$ must be obtained by adopting appropriate numerical procedures.

It is important to remark that in the case of censored data, $\text{Asvar}[\ln(\widehat{x_p(z_u)})]$ depends on the unknown parameters, γ_0 , γ_1 , and σ , that one wants to estimate [3, p. 331]. So, to calculate optimum plans one needs to use in their place values obtained from experience, or from preliminary tests. An optimum plan that uses tentative values is, in general, better than plans devised by other means. Nevertheless, it is only locally optimum, hence it is strongly recommended to test the robustness of the procedure by performing a sensitivity analysis over the region in which the true values plausibly fall.

Practical charts for calculating optimum plans are given in [29, 30]. These charts give the optima extrapolation factor, ξ^* , and proportion, π_1^* as a function of the standardized censoring time, a , and slope, b , where:

$$\xi^* = \frac{z_2 - z_u}{z_2 - z_1^*} \quad (45)$$

$$a = \frac{\eta - (\gamma_0 + \gamma_1 \cdot z_2)}{\sigma} \quad (46)$$

$$b = \frac{\gamma_1 \cdot (z_u - z_2)}{\sigma} \quad (47)$$

Calculation of a and b requires tentative values for γ_0 , γ_1 , and σ (censoring time η and high stress level z_2 are given). The optimum low level z_1^* can be calculated as follows:

$$z_1^* = z_2 - \xi^*(z_2 - z_u) \quad (48)$$

Concluding Remarks

Inference from ALT data involves extrapolation in stress. One of the key problems in ALTs is to determine whether the model that fits data obtained under accelerated experimental conditions is able to fit data under normal use condition or not. The use of models that have a sound physical or chemical justification can help to mitigate the risk of making dangerous errors. Nevertheless, many serious pitfalls (see [31] for an extensive analysis of this topic) can be caused by ALT.

The use of sensitive analysis can help to develop robust procedures and evaluate the consequences of model misspecifications. Nevertheless, not exaggerating the overstress level is always the best solution [32], even if this leads again to a reduction of failure data.

In particular, with regard to the models presented in this paper, it is important to remark that they are for single-failure modes and simple failure-causing processes. Thus, their use can be considered adequate, for the aforementioned extrapolation, only if overstress accelerates the mechanism of failure under study without altering its nature or inducing other failure modes.

Analyses of ALT data in presence of more, detected or undetected, causes of failure or failure masking, have to be faced by adopting specific methods and models (*see Accelerated Life Tests: Analysis with Competing Failure Modes; Masked Failure Data*).

Appendix

Life–Stress Relationships

The models presented in this appendix are three of the most widely used to express the dependence of lifetime on a single accelerating variable. They are for constant-stresses and single-failure modes.

Their adequacy has been empirically demonstrated for many products; nevertheless, their use can often be theoretically justified.

On the other hand, for the same theoretical arguments, they can be considered adequate for extrapolation only if the separation between test and use conditions is limited.

Arrhenius Model. The Arrhenius life–stress relationship is widely employed to model the effect of temperature on lifetime of degrading product that fail when degradation overcomes a critical threshold. It can be satisfactorily applied in the case of simple failure-causing chemical or physical reactions (i.e., a single rate-controlling step exists, and that step is either a diffusion or an n th order chemical reaction [33]). According to this model, the dependence between the product characteristic life, τ , (i.e. the median or a given quantile of the lifetime) and the test temperature, T (absolute kelvin), can be expressed as

follows:

$$\tau = A \cdot \exp\left(\frac{E_a}{k \cdot T}\right) \quad (\text{A1})$$

where E_a is the activation energy (eV), k is the Boltzmann's constant (eV K⁻¹), and A is a constant that depends on product geometry, test methods, and other factors.

The acceleration factor between life, τ_u , at use temperature, T_u , and life τ at a given reference temperature, T , can be expressed as:

$$\varphi(T_u, T) = \frac{\tau_u}{\tau} = \exp\left[\frac{E_a}{k} \left(\frac{1}{T_u} - \frac{1}{T}\right)\right] \quad (\text{A2})$$

Appropriateness of the model has been confirmed experimentally for many product and materials [3, p. 75].

Inverse Power Law Model. The inverse power law model is an empirical relationship that is widely used to model the life–stress relationship in many different ALTs. As an instance, it has been successfully applied in accelerated voltage tests of insulators, accelerated thermal fatigue tests of metals, and in accelerated mechanical load tests of rolling bearings [3, p. 85]. Other popular models such as the Coffin–Manson relationship and the Palmgren equation can be considered special characterization of inverse power law model.

This life–stress model expresses the dependence between the product characteristic life, τ , and the applied stress as follows:

$$\tau = A \cdot S^{-c} \quad (\text{A3})$$

where the parameters A and c (usually called the *power*) are constants whose values depend on product geometry, test methods, and other factors.

The acceleration factor between life τ_u , at normal use stress level S_u and life, τ , at reference stress S and can be expressed as:

$$\varphi(S_u, S) = \frac{\tau_u}{\tau} = \left(\frac{S}{S_u}\right)^c \quad (\text{A4})$$

A justification for this model in the case of voltage-stress acceleration of insulators is given in [1, p. 482].

Exponential Relationship. According to this model, the relationship between life and a possibly transformed stress, S , can be expressed as follows:

$$\tau = \exp(a_1 - a_2 \cdot S) \quad (\text{A5})$$

The acceleration factor between life, τ_u , at use stress, S_u , and life τ at a given reference stress, S , is:

$$\varphi(S_u, S) = \frac{\tau_u}{\tau} = \exp[a_2(S - S_u)] \quad (\text{A6})$$

The use of this model is considered adequate for many electronic components [3, p. 96].

References

- [1] Meeker, W.Q. & Escobar, L.A. (1998). *Statistical Methods for Reliability Data*, John Wiley & Sons, New York.
- [2] Meeker, W.Q. & Escobar, L.A. (1993). Review of recent research in accelerated testing, *International Statistical Review* **61**(1), 147–168.
- [3] Nelson, W.B. (1990). *Accelerated Testing*, John Wiley & Sons, New York.
- [4] MIL-HDBK-217F. (1991). *Reliability Prediction of Electronic Equipment*, 2 December, 1991 (superseding MIL-HDBK-217E, Notice 1, 2 January 1990).
- [5] SKF. (1989). *General Catalogue, Short Version*, (in Italian), 1335-I Reg. 47 · 20000 1990 – 05, Stamperia Artistica Nazionale.
- [6] Cacciari, M. & Montanari, G. (1989). Modelli statistici per l'invecchiamento elettrico dei materiali isolanti solidi (in Italian), *Statistica* **49**(1), 73–81.
- [7] Meeter, C.A. & Meeker, W.Q. (1994). Optimum accelerated life tests with a non-constant scale parameter, *Technometrics* **6**(1), 71–83.
- [8] Swartz, J.A. (1987). Effect of temperature on the variance of the log-normal distribution of failure time due to electromigration damage, *Journal of Applied Physics* **61**(2), 801–803.
- [9] Wang, W. & Kececioglu, D.B. (2000). Fitting the Weibull log-linear model to accelerated life-test data, *IEEE Transactions on Reliability* **49**(2), 217–223.
- [10] Chan, C.K. (1991). Temperature-dependent standard deviation of log(failure time) distribution, *IEEE Transactions on Reliability* **40**(2), 157–160.
- [11] Cox, D.R. (1972). Regression model and life table (with discussion), *Journal of the Royal Statistical Society. Series B* **26**, 103–110.
- [12] Lawless, J.F. (1982). *Statistical Models and Methods for Lifetime Data*, John Wiley & Sons.
- [13] Dale, C.J. (1985). Application of the proportional hazard model in the reliability field, *Reliability Engineering* **10**, 1–14.
- [14] Mazzuchi, T.A., Soyer, R. & Spring, R.V. (1989). The proportional hazard model in reliability, *Proc. of the Annual Reliability and Maintainability Symposium*, Atlanta, Georgia USA, January 24–26 252–256.
- [15] Mann, N.R., Shafer, R.E. & Singpurwalla, N.D. (1974). *Methods for Statistical Analysis of Reliability and Life Time Data*, John Wiley & Sons.
- [16] Viertl, R. (1988). *Statistical Methods in Accelerated Life Testing*, Vandenhoeck & Ruprecht, Göttingen.
- [17] Yang, G. (2007). *Life Cycle Reliability Engineering*, John Wiley & Sons.
- [18] Bugaighis, M.M. (1988b). A note on convergence problems in numerical techniques for accelerated life-testing analysis, *IEEE Transactions on Reliability* **37**(3), 348–349.
- [19] Watkins, A.J. (1991). On the analysis of accelerated life testing experiment, *IEEE Transactions on Reliability* **40**(1), 98–101.
- [20] Watkins, A.J. (1994). Review: likelihood method for fitting to accelerated Weibull log-linear models life-test data, *IEEE Transactions on Reliability* **43**(3), 361–365.
- [21] Vander Wiel, S.A. & Meeker, W.Q. (1990). Accuracy of approx confidence bounds using censored Weibull regression data from accelerated life tests, *IEEE Transactions on Reliability* **39**(3), 346–351.
- [22] Lawless, J.F. (1976). Confidence interval estimation in the inverse power law model, *Applied Statistics* **25**(2), 128–138.
- [23] McCool, J.I. (1980). Confidence, limits for Weibull regression with censored data, *IEEE Transactions on Reliability* **R-29**(2), 145–150.
- [24] Nelson, W.B. (1982). *Applied Life Data Analysis*, John Wiley & Sons, New York.
- [25] Nelson, W.B. & Kielpinski, T.J. (1976). Theory for optimum censored accelerated life tests for normal and lognormal life distributions, *Technometrics* **18**(1), 105–114.
- [26] Nelson, W.B. & Meeker, W.Q. (1978). Theory for optimum accelerated censored life tests for Weibull and extreme value distributions, *Technometrics* **20**(2), 171–177.
- [27] Nelson, W.B. (2005). A bibliography of accelerated test plans, *IEEE Transactions on Reliability* **54**(2), 194–197.
- [28] Nelson, W.B. (2005). A bibliography of accelerated test plans, Part II – references, *IEEE Transactions on Reliability* **54**(3), 370–373.
- [29] Kielpinski, T.J. & Nelson, W.B. (1975). Optimum accelerated life-tests for the normal and lognormal life distributions, *IEEE Transactions on Reliability* **R-24**(5), 310–320.
- [30] Meeker, W.Q. & Nelson, W.B. (1975). Optimum accelerated life tests for the Weibull and extreme value distributions, *IEEE Transactions on Reliability* **R-24**(5), 321–332.
- [31] Meeker, W.Q. & Escobar, L.A. (1998). Pitfalls of accelerated testing, *IEEE Transactions on Reliability* **47**(2), 114–118.
- [32] Evans, R.A. (1991). Accelerated testing, *IEEE Transactions on Reliability* **40**(5), 497.

- [33] LuValle, M.J. (1999). An approximate kinetic theory for accelerated testing, *IIE Transactions* **31**, 1147–1156.

Related Articles

Accelerated Life Models; Accelerated Life Tests: Bayesian Design; Accelerated Life Tests: Bayesian

Models; Accelerated Life Tests: Designs Comparison within a Bayesian Framework; Accelerated Life Tests: Nonparametric Approach; Masked Failure Data; Prediction of Expected Fatigue Lives of Fiber Reinforced Plastic Joints; The Cox Proportional Hazard Model.

MASSIMILIANO GIORGIO

Accelerated Life Tests: Bayesian Design

Introduction and Overview

Accelerated life tests (ALTs) involve testing systems in an environment that is more severe than the use environment and using the data collected in the accelerated environment for inferring failure behavior in the use environment. The design problem in accelerated life testing is concerned with specification of the number and magnitude of the accelerated stress levels, and the number of items to be tested at these stress levels.

Most of the Bayesian literature in ALTs have focused on developing inference methods; see for example, Mazzuchi and Soyer [1], van Dorp *et al.* [2], Mazzuchi *et al.* [3], and van Dorp and Mazzuchi [4, 5]. The majority of the work published on design of ALTs relied on sample theoretic methods; see for example, the text by Nelson [6] and the recent bibliography of accelerated testing plans by Nelson [7, 8]. Exceptions to these are the earlier Bayesian papers by Martz and Waterman [9], DeGroot and Goel [10], and more recent Bayesian approaches by Menzefricke [11], Chaloner and Larntz [12], Polson [13], Verdinelli *et al.* [14], Soyer and Vopatek [15], Erkanli and Soyer [16], and Zhang and Meeker [17].

Most of these Bayesian approaches are based on the theory of optimal Bayesian designs for linear models as in Chaloner [18]. Thus, the results are applicable to ALT designs when the **life model** is normal or lognormal. For example, Chaloner and Larntz [12] consider Bayesian designs for type I censored tests when there is uncertainty about whether the underlying life model is lognormal or Weibull, and several fractiles of the life length distribution at the use stress are of interest. The optimality criterion considered by authors is proportional to the expected asymptotic variance of the fractiles of interest. Designs satisfying this criterion are identified using numerical methods. This approach can be considered as a Bayesian version of the sample theoretic methods considered in Nelson [6]. Menzefricke [11] formulates a Bayesian approach to the **optimal design** of type II censored ALTs for use when the life length model is lognormal. Verdinelli *et al.* [14] identify a design for a complete ALT

that maximizes utility, as represented by Shannon information, given the usual linear model assumptions. More recent work by Zhang and Meeker [17] also considers some large sample results for Bayesian ALT designs as well as simulation-based methods.

As noted by Vopatek [19] and Soyer and Vopatek [15], if the underlying life models are exponential or Weibull, nonlinearities arise in the analysis, and then optimal designs can be obtained by use of either numerical methods or special techniques such as linear Bayesian methods. Recent advances in statistical computing, nonparametric surface estimation, and implementation of **Markov chain Monte Carlo** (MCMC) methods have contributed to the development of computationally efficient methods for optimal designs. For example, recent simulation-based approach of Muller and Parmigiani [20] and its extensions presented by Muller [21] have been considered in ALT designs. Erkanli and Soyer [16] have used this approach for fixed (nonsequential) designs and developed an extension to sequential designs.

In this article, we some recent Bayesian approaches in ALT Designs are reviewed. By doing so, the Bayesian decision theoretic setup for the optimal ALT design problem following Polson [13] and Erkanli and Soyer [16] is presented. This setup is adopted to a single-point design problem and difficulties involved in evaluation of preposterior losses in the implementation of the Bayesian paradigm are illustrated. Linear Bayesian methods and Monte Carlo-based approaches to alleviate some of these difficulties are presented. Finally the implementation of a simulation-based design algorithm using the approach of Muller and Parmigiani [20] is dealt with.

Bayesian Decision Theoretic Setup for the Optimal Design Problem

In Bayesian paradigm, the optimal design problem can be viewed as a decision problem in the sense of Lindley [22]. Thus, the optimal designs are chosen by maximizing expected utility. As noted by Polson [13], this provides a formal approach to the design of experiments. A comprehensive review of Bayesian experimental design can be found in Chaloner and Verdinelli [23].

Erkanli and Soyer [16] point out that the Bayesian approach to the optimal design problem requires specification of three components:

2 Accelerated Life Tests: Bayesian Design

1. a utility (loss) function that reflects the consequences of selecting a specific design;
2. a probability model;
3. a prior distribution reflecting designer's *a priori* beliefs about all unknown quantities.

Let d denote a specific design, that is, the decision variable, λ_u , the failure characteristic of interest (such as the use environment failure rate) and, a , an action based on data (for example, a prediction for λ_u). We denote the utility function as $U(\lambda_u, d, a)$ and the probability model as $p(D|\lambda_u, d)$ where D is the data observed from experiment d . Prior beliefs about λ_u are described by the prior distribution $p(\lambda_u)$. It is possible that λ_u may be a vector or a function of several parameters. Once the design d is specified and the ALT is performed, data D is observed and uncertainty about λ_u is revised according to the laws of probability. Then, the optimal action a is chosen on the basis of the data. This process can be depicted by the decision tree in Figure 1.

In Figure 1, the decision node D_2 represents selection of the Bayes rule a given data. For the case of squared error loss $L(\lambda_u, d, a)$, the utility function is $U(\lambda_u, d, a) = -(\lambda_u - a)^2$ and the optimal Bayes rule a^* is the posterior mean. Then the last two nodes can be replaced by $U(\lambda_u, d, a^*) = -V(\lambda_u|D, d)$, which is the posterior variance of λ_u and the optimal design is chosen by minimizing the preposterior variance

$$E_{D|d}[V(\lambda_u|D, d)] = \int V(\lambda_u|D, d) p(D|\lambda_u, d) \times p(\lambda_u) d\lambda_u dD \quad (1)$$

In other words, the optimal design is given by $d^* = \operatorname{argmin}\{E_{D|d}[V(\lambda_u|D, d)]\}$.

As pointed out by Erkanli and Soyer [16], this solution is obtained as a result of “folding back” the decision tree through taking expectations at random nodes and maximizing the expected utility at the decision nodes of Figure 1. The above setup can be adopted for any form of the utility function

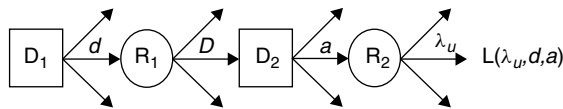


Figure 1 Decision tree for the design problem

$U(\lambda_u, d, a)$. For example, Verdinelli *et al.* [14] used the Shannon's information as the utility function. Thus, in general the optimal design is obtained as

$$d^* = \operatorname{argmax}_d \{E_{D|d}[U(\lambda_u, d, a^*)|D, d]\} \quad (2)$$

where $a^* = \operatorname{argmax}\{E_{\lambda_u|D, d}[U(\lambda_u, d, a)|D, d]\}$. In what follows, an example of the decision theoretic setup assuming exponential lifetimes is presented.

Example: Single Stress ALT Designs

An important component of the ALTs is the time transformation function (or the acceleration function) that describes the relationship between the failure characteristic of interest and the applied level of the stress variable.

Let λ_i denote the failure rate at the accelerated stress environment S_i , assuming that life length X_i is exponential with λ_i and the time transformation function is given by the power law model

$$\lambda_i = \alpha S_i^\beta \quad (3)$$

Let us consider the design of an ALT where n items will be tested until failure at a single accelerated stress environment S_i with the purpose of making inference about $\lambda_u = \alpha S_u^\beta$, the failure rate at the use stress environment. Such Single-point designs have been considered by Martz and Waterman [9].

In the above case, the design problem is to find the stress level S_i minimizing the preposterior variance of λ_u . In other words, in our setup, we assume that the design variable $d = S_i$ and the loss function is given by $L(\lambda_u, d, a) = (\lambda_u - a)^2$. Thus, the design problem requires computation of

$$S_i^* = d^* = \operatorname{argmin}\{E_{D|d}[V(\lambda_u|D, d)]\} \quad (4)$$

In other words, our decision problem in Figure 1 reduces to a single-stage problem. For illustrative purposes, we assume that β is known and we specify a gamma prior distribution for α with shape parameter a and scale b . Given n failures $D = (x_1, x_2, \dots, x_n)$ the posterior distribution of α is given by a gamma density denoted as $[G(a + n), (b + S_i^\beta T_i)]$, where T_i is the **total time on test** at stress environment S_i . Using the power law model,

the posterior variance of λ_u is

$$V(\lambda_u|D, S_i) = V(\alpha S_u^\beta|D, S_i) = \frac{(a+n)(S_u/S_i)^{2\beta}}{\left(T_i + \frac{b}{S_i^\beta}\right)^2} \quad (5)$$

It can be shown that the preposterior variance $E_{D|d}[V(\lambda_u|D, S_i)]$ is given by

$$V(\lambda_u|S_i) = \frac{a(a+1)S_u^{2\beta}}{(a+n+1)b^2} \quad (6)$$

where the expectation is with respect to the predictive distribution of T_i given S_i (Erkanli and Soyer 16). Since equation (6) is not a function of the stress S_i , it does not matter at what stress level the items are tested. This is due to the fact that β is assumed to be known in the power law model. This result was also noted by Vopatek [19] and Verdinelli *et al.* [14] for the special case of the power law with $\beta = 1$.

In general, except in a few special cases, the optimal designs in equations (1) and (2) cannot be obtained analytically. In ALT models where the underlying failure distribution is exponential or Weibull, posterior and/or predictive inferences are not in closed form. Thus, an evaluation of the designs in equations (1) and (2) requires use of either approximation techniques such as linear Bayesian methods as in Soyer and Vopatek [15] or Monte Carlo methods as in Erkanli and Soyer [16] and Zhang and Meeker [17].

Linear Bayesian Designs for ALTs

In the power law model example in the section titled “Example: Single Stress ALT Designs”, if the coefficient β is unknown, then the posterior variance $V(\lambda_u|D, S_i)$ cannot be obtained in closed form for any prior specification for β in the case of exponential lifetimes. One way to deal with this problem is to consider approximate Bayesian designs using linear Bayesian methods.

Soyer and Vopatek [15] consider the power law model in equation (3) with exponential lifetimes where both α and β are treated as unknown. The authors describe their uncertainty about these unknown quantities partially by specifying only the first- and second-order **moments**. More specifically,

they consider the log transformation of the power law model as

$$\eta_i = \log \lambda_i = \log \alpha + \beta \log S_i = \mathbf{F}_i' \boldsymbol{\theta} \quad (7)$$

where $\mathbf{F}_i' = (1 \log S_i)$ and $\boldsymbol{\theta} = (\log \alpha \beta)$. Prior to testing at S_i distribution of $\boldsymbol{\theta}$ is partially specified as $\boldsymbol{\theta} \sim (\mathbf{m}_0, \mathbf{C}_0)$. Following Mazzuchi and Soyer [1], the prior moments can be used in equation (7) to specify a log-gamma distribution as the prior for η_i with parameters a_i and b_i selected such that

$$\Psi(a_i) - \log b_i = \mathbf{F}_i' \mathbf{m}_0 \text{ and } \Psi'(a_i) = \mathbf{F}_i' \mathbf{C}_0 \mathbf{F}_i \quad (8)$$

where $\Psi(\bullet)$ and $\Psi'(\bullet)$ are the digamma and trigamma functions, respectively. Note that complete specification of the prior for η_i (and thus for λ_i) enables us to obtain the posterior distribution for η_i after testing at S_i . Under the assumption of no censoring, standard Bayesian conjugate analysis shows that the posterior distribution of η_i is a log-gamma density with parameters $(a_i + n)$ and $(b_i + T_i)$, where n is the total number of items tested under environment S_i and T_i is the total time on test.

To obtain the posterior variance $V(\eta_u|D, S_i)$ or $V(\lambda_u|D, S_i)$, it is necessary to update $\boldsymbol{\theta} = (\log \alpha \beta)$. Since prior distribution of $\boldsymbol{\theta}$ is only partially specified through moments, one can only update the moments given the data. This updating is done *via* using the linear Bayesian methods. Soyer and Vopatek [15] show that the posterior moments of $\boldsymbol{\theta}$ are given by

$$\mathbf{m}_i = \mathbf{m}_0 + \mathbf{h}_i \frac{E(\eta_i|D) - E(\eta_i)}{V(\eta_i|D)} \quad (9)$$

$$\mathbf{C}_i = \mathbf{C}_0 - \mathbf{h}_i \mathbf{h}_i' \frac{1 - V(\eta_i|D)/V(\eta_i)}{V(\eta_i|D)} \quad (10)$$

where $\mathbf{h}_i = \mathbf{C}_0 \mathbf{F}_i$. Thus, the posterior distribution of $\boldsymbol{\theta}$ is partially specified by the two moments as $(\boldsymbol{\theta}|D) \sim (\mathbf{m}_i, \mathbf{C}_i)$.

Once such posterior specification is available, we can use the fact that at the use stress we have $\eta_u = \mathbf{F}_u' \boldsymbol{\theta}$, where $\mathbf{F}_u' = (1 \log S_u)$ and obtain the posterior variance of η_u . We can show that $V(\eta_u|D, S_i)$ is given by

$$\mathbf{F}_u' \mathbf{C}_0 \mathbf{F}_u - (\mathbf{F}_u' \mathbf{C}_0 \mathbf{F}_i)^2 \left[\frac{1 - \Psi'(a_i + n)/\mathbf{F}_i' \mathbf{C}_0 \mathbf{F}_i}{\mathbf{F}_i' \mathbf{C}_0 \mathbf{F}_i} \right] \quad (11)$$

which implies that as n increases the posterior variance decreases for any level of S_i as expected. We

note that the posterior variance is not a function of the prior mean $\mathbf{m}_0$ but it depends on prior **covariance** matrix $\mathbf{C}_0$. Furthermore, the above is not a function of the total time on test T_i . As a result of this, the posterior variance is the same as the preposterior variance, that is, $V(\eta_u|D, S_i) = V(\eta_u|S_i)$. In other words, one can find the single-point optimal design by simply minimizing equation (11) with respect to S_i .

Different forms of the prior covariance matrix $\mathbf{C}_0$ were considered by Soyer and Vopatek [24] and the corresponding optimal designs were presented. For example, in the special case of the power law model where α is known, say $\alpha = 1$, it can be shown that the optimal design is to test all the items at the highest possible stress level S_H , that is, $S_i^* = S_H$. It was shown by Vopatek [19] that if $\log \alpha$ and β are assumed to be independent *a priori*, that is, if $\mathbf{C}_0$ is a diagonal matrix, then for large n an optimal design can be obtained as

$$S_i^* = S_u^{[1+1/[nV(\log \alpha)]]} \quad (12)$$

This implies that as the number of items to be tested is large, the optimal stress is closer to the use stress S_u . Similarly, as prior uncertainty about $\log \alpha$ increases, the optimal stress level moves closer to S_u .

If the loss function is specified in terms of the posterior variance of λ_u rather than of η_u then it can be shown that $V(\lambda_u|D, S_i)$ is a function of the total time on test. More Specifically, in order to obtain $V(\lambda_u|D, S_i)$, it is necessary to have a distributional form for η_u and thus for λ_u . This is done after updating by specifying a log-gamma distribution for η_u with parameters a_u and b_u selected such that

$$\Psi(a_u) - \log b_u = \mathbf{F}'_u \mathbf{m}_i \text{ and } \Psi'(a_u) = \mathbf{F}'_u \mathbf{C}_i \mathbf{F}_u \quad (13)$$

The above implies a gamma posterior for λ_u , implying that $V(\lambda_u|D, S_i) = a_u/(b_u)^2$. Using this form, Vopatek [19] obtained the preposterior variance $V(\lambda_u|S_i)$.

The linear Bayesian setup can be generalized to multiple point ALT designs in which one considers m different stress levels where $n_1, n_2, \dots, n_m$ are tested at stress environments $S_1, S_2, \dots, S_m$. This can be done either in an adaptive manner by solving m sequential one-point design problems or solving a fixed design problem. Some of these extensions as well as design of censored tests are considered in Vopatek [19] using linear Bayesian methods.

Simulation Methods for Bayesian Designs

In the power law model of the section titled “Example: Single Stress ALT Designs”, if β is unknown, then the preposterior variance of λ_u cannot be obtained in closed form. One way to compute the posterior variance is to use a MCMC method such as the Gibbs sampler; see Gilks *et al.* [24] for an excellent introduction to MCMC methods.

The Gibbs sampler requires the full conditional distributions $p(\alpha|\beta, D)$ and $p(\beta|\alpha, D)$. Using a gamma prior for α as before, the full conditional of α is given by the gamma distribution $\Gamma[(a+n), (b+S_i^\beta T_i)]$. On the other hand, any prior distribution on β does not yield a standard form for the full conditional $p(\beta|\alpha, D)$, which is given by

$$p(\beta|\alpha, D) \propto S_i^{\beta n} \exp(-\alpha S_i^\beta T_i) p(\beta) \quad (14)$$

Thus, a method such as the rejection sampling needs to be used to draw samples from $p(\beta|\alpha, D)$. In this particular case, the standard rejection sampling can be easily implemented since the maximum of the conditional likelihood $\mathcal{L}(\beta; \alpha, D)$ is analytically available. We can show that the maximum of the conditional likelihood is given by

$$\hat{\beta} = \frac{\log(n) - \log(\alpha T_i)}{\log(S_i)} \quad (15)$$

where $\log(n) - \log(\alpha T_i) > 0$. Thus, we can design a rejection sampling algorithm to generate from $p(\beta|\alpha, D)$ by using the prior $p(\beta)$ as the importance function as in Smith and Gelfand [25]. In other words, we can draw β from the prior and accept it with probability

$$\frac{\mathcal{L}(\beta; \alpha, D)}{\mathcal{L}(\hat{\beta}; \alpha, D)}$$

Once a sample is obtained from the joint posterior $p(\alpha, \beta|D)$, the posterior distribution $p(\lambda_u|D, S_i)$ and the variance $V(\lambda_u|D, S_i)$ can be computed from this sample. In finding the optimal design, we need to obtain the preposterior variance

$$E_D[V(\lambda_u|D, S_i)] = \int V(\lambda_u|D, S_i) p(D|S_i) dD \quad (16)$$

which can be written as

$$\begin{aligned} E_D[V(\lambda_u|D, S_i)] &= \iint V(\lambda_u|D, S_i) \\ &\times p(D|\lambda_u, S_i) \\ &\times p(\lambda_u) d\lambda_u dD \end{aligned} \quad (17)$$

We note that in the above $V(\lambda_u|D, S_i) = V(\lambda_u|T_i, S_i)$ is evaluated *via* the Gibbs sampler for a given design $d = S_i$. Since the integral in equation (17) is not available in closed form, we can use a Monte Carlo average to evaluate it. More specifically, for each generated value of (α_r, β_r) , $r = 1, \dots, R$, from the prior we can generate λ_i and T_i . On the basis of each generated value of T_i , using the Gibbs sampler we can obtain the posterior variance $V[\lambda_u|T_i^r, S_i]$ and compute preposterior variance for S_i using the Monte Carlo average

$$\frac{1}{R} \sum_{r=1}^R V[\lambda_u|T_i^r, S_i] \quad (18)$$

Note that an optimal one-point design can be obtained by minimizing equation (18) with respect to S_i . This Monte Carlo setup can be modified for the general ALT design problem as discussed next.

Standard Monte Carlo Approach for Bayesian ALT Designs

Erkanli and Soyer [16] consider a general fixed design for ALTs, where m distinct stress levels are used and n_i items are tested at the stress level S_i such that $n = \sum_{i=1}^m n_i$ is a predetermined number. The design problem is then to select

- $m \leq n$, the number of distinct stress levels
- the accelerated stress levels S_i , $i = 1, \dots, m$; and
- the number of items tested at each stress level, n_i , $i = 1, \dots, m$

in such a way that the expected utility is maximized. In the above, a specific design is given by $d = \{m, S_i, n_i, i = 1, \dots, m\}$. As in the section titled “Example: Single Stress ALT Designs”, the authors consider the power law model where one is interested in making inference about λ_u , the failure rate at the use stress environment.

As before, we assume that there is no censoring in the ALT and we choose the design minimizing the preposterior variance of λ_u . In other words, the evaluation of the optimal design requires the computation of posterior variance $V(\lambda_u|D, d)$, where $D = \{T_1, \dots, T_m\}$ and $T_i = \sum_{j=1}^{n_i} x_j$ is the total time on test at stress environment S_i . As in the previous case, evaluation of $V(\lambda_u|D, d)$ requires use of MCMC methods.

The following algorithm has been suggested by Erkanli and Soyer [16] for standard Monte Carlo evaluation of the preposterior variance $E_D[V(\lambda_u|D, d)]$:

- Step 1. Choose $d = \{m, S_i, n_i, i = 1, \dots, m\}$.
- Step 2. Generate (α_r, β_r) from the prior $p(\alpha, \beta)$, $r = 1, \dots, R$.
- Step 3. For $i = 1, \dots, m$, generate T_{ir} from $p(T_i|\lambda_i)$ using the power law.
- Step 4. For each $D_r = \{T_{1r}, \dots, T_{mr}\}$ evaluate $V[\lambda_u|D_r, d]$ using MCMC.
- Step 5. Compute preposterior variance for d using Monte Carlo average

$$\frac{1}{R} \sum_{r=1}^R V[\lambda_u|D_r, d] \quad (19)$$

- Step 6. Go to step 1 and repeat the steps 2–5 for a different design d .

The optimal design is selected as the d with the minimum value of equation (19).

A Surface Fitting Algorithm for Bayesian ALT Designs

As noted by Erkanli and Soyer [16], the implementation of the above Monte Carlo approach is not computationally efficient. The approach requires, for each level of the design variable, R draws from the statistical model and, for each draw, evaluation of the posterior variance using MCMC methods that typically requires large number of iterations. Thus, the approximation in equation (19) may require a large-scale computational effort. Especially, this will be inefficient for the case of multiple stress designs where we need to specify m optimal stress environments. To avoid the potential computational burden, the authors suggested a curve fitting approach proposed by Muller and Parmigiani [20] for finding optimal ALT designs. This approach facilitates preposterior analysis by replacing the expectation step with a smoothing step. The posterior variance is evaluated as $v_r = V(\lambda_{ur}|D_r, d_r)$ for each simulated experiment $\{d_r, \alpha_r, \beta_r, D_r\}$ and a smooth surface $L(d)$ is fitted to the points (d_r, v_r) and the optimal design is found by minimizing the fitted surface $L(d)$.

The proposed approach by Erkanli and Soyer [16] for the ALT design problem with $d = \{m, S_i, n_i, i = 1, \dots, m\}$ and $\theta = (\alpha, \beta)$ is as follows:

- Step 1. Select designs $d_r, r = 1, \dots, R$.
- Step 2. Draw R points (D_r, θ_r) from the density $p(D, \theta|d_r)$. This is done by independently generating $(T_{1r}, T_{2r}, \dots, T_{mr})$ from $p(T_i | \lambda_i, d_r)$ for $i = 1, \dots, m$. Use MCMC to evaluate $v_r = V(\lambda_{ur} | D_r, d_r)$ and record sample points (d_r, D_r, v_r) .
- Step 3. Fit a surface $L(d)$ to the points (d_r, v_r) .
- Step 4. Find the minimum over d of $L(d)$.

The above setup assumes that a total number of n items will be tested in an ALT. It is important to note that d consists of both discrete and continuous components. Given n , the discrete components of d , $\{m, n_i, i = 1, \dots, m\}$ are constrained as $n = \sum_{i=1}^m n_i$ and $m \leq n$. Erkanli and Soyer [16] considered the case of $m = 2$ point designs which are shown to be optimal for relationships such as power law model with two parameters (see for example, Chaloner and Larntz [12]).

For the case of two-point designs, an alternate design strategy is to consider a sequential design. The general m -stage sequential design problem for ALT is shown in Figure 2. The surface-fitting approach conceptually can be adopted to the sequential problem, but as the number of stages increases, the dimension of the surfaces that we fit at each stage also increases and this makes the implementation quite difficult. Thus, the proposed approach was extended in Erkanli and Soyer [16] for only two-stage design problems using a setup similar to what is considered in Erkanli *et al.* [26] for prevalence estimation.

An Illustration of the Curve Fitting Method

In this section, we will illustrate the implementation of the curve fitting algorithm introduced in the last section to a one-point design problem using exponential lifetimes and a power law model. In the illustration, the local regression model, Loess, of Chambers and Hastie [27] has been used in

step 3. Any one-dimensional smoothing method such as splines can be used for this purpose.

For illustrative purposes, we will assume independent gamma priors for α and β such that $\alpha \sim \Gamma(20, 1000)$ and $\beta \sim \Gamma(3, 1)$. We will assume that at the optimal stress level, $n = 2$ items will be tested and the design space consists of the stress range $S \in (1.05, 11)$ where $S_u = 1.05$ is the use stress. As previously discussed, our goal is to find the optimal stress level S_i in such a way that the preposterior variance of λ_u is minimized. In what follows, as an alternative, we are interested in finding the optimal stress level minimizing the preposterior variance of $\eta_u = \log \lambda_u$ so that our results are comparable to those of Soyer and Vopatek [15] where linear Bayesian designs are obtained.

We note that under the specified form of priors, we still need to use the Gibbs-rejection algorithm to evaluate the posterior variance $V(\eta_u | T_{ir}, d_r)$ in step 2 of the curve fitting algorithm for each design point $d_r = S_i, r = 1, \dots, R$. To find the optimal stress level, we choose $R = 1000$ design points in the range $(1.05, 11)$ in step 1 of the algorithm. The next 1000 random vectors $(\alpha_r, \beta_r, T_{ir})$ are simulated using the joint distribution $p(\alpha, \beta) p(T_{ir} | \alpha, \beta)$, where $(T_{ir} | \alpha, \beta) \sim \Gamma(n, \lambda_{ir})$ and $\lambda_{ir} = \alpha S_i^\beta$. After implementing the Gibbs sampler for each generated data point and evaluating the posterior variance $V(\eta_u | T_{ir}, d_r)$ for $r = 1, \dots, 1000$, we can fit a nonparametric regression curve to the points $(d_r = S_i, v_r = V(\eta_u | T_{ir}, d_r))$. As discussed before, this is equivalent to taking the expectation of the posterior variance with respect to T_{ir} . In Figure 3 we have the nonparametric approximation to preposterior variance of η_u , that is, $V(\eta_u | S_i)$. As can be seen from the figure the minimum of the curve is around 2.4–2.5. We note the fitted curve is flat around the minimum. This is owing to the fact that the loss function does not consider any costs associated with testing. Also, we note the variation in the preposterior variance, which is due to the large variance of β . These results are very similar to the

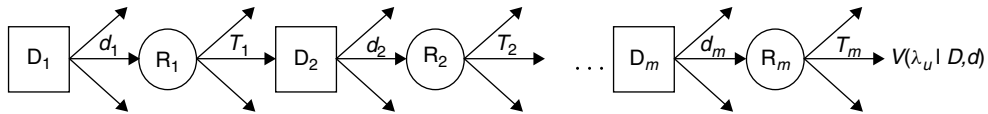


Figure 2 Decision tree for the m -stage design problem

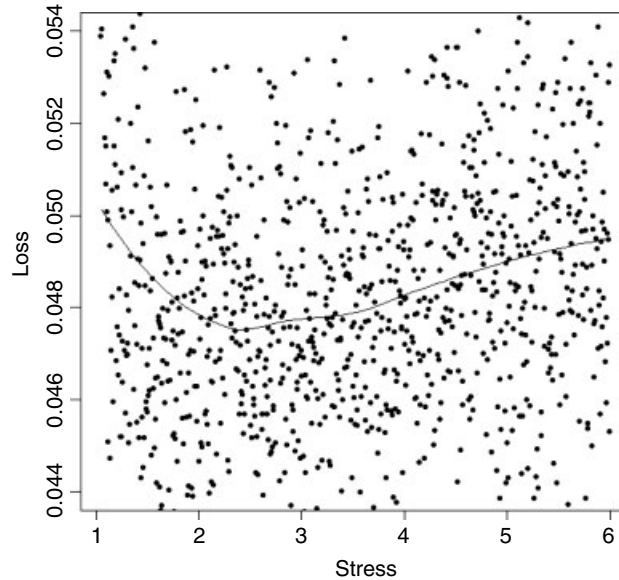


Figure 3 Nonparametric approximation to $V(\eta_u|S_i)$ with $n = 2$

findings of Vopatek [19], which are based on linear Bayesian methods. The results are also pretty robust to the choice of R as long as $R > 500$.

Concluding Remarks

In this article, some of the recent results in the Bayesian design of ALTs have been reviewed. Following Polson [13] and Erkanli and Soyer [16], the design problem has been presented as a decision problem in the sense of Lindley [22] and some of the difficulties that arise in the development of Bayesian designs have been discussed. There is an illustration of how linear Bayesian methods and Monte Carlo based approaches can be used in alleviating some of these difficulties. Some of the author's experiences in using the simulation-based design algorithm of Muller and Parmigiani [20] have been presented. These findings suggest that this approach can be successfully implemented to fixed design problems of moderate dimensions. However, extension of the approach to more general ALT designs and sequential problems of more than two stages requires more experience.

References

- [1] Mazzuchi, T.A. & Soyer, R. (1992). A dynamic general linear model for inference from accelerated life testing, *Naval Research Logistics Quarterly* **39**, 757–773.
- [2] van Dorp, R., Mazzuchi, T.A., Fornell, G. & Pollock, L.R. (1996). A Bayes approach to step-stress accelerated life testing, *IEEE Transactions on Reliability*, **R-45** 491–498.
- [3] Mazzuchi, T.A., Soyer, R. & Vopatek, A.L. (1997). Linear Bayesian inference for accelerated Weibull model, *LifeTime Data Analysis* **3**, 1–13.
- [4] van Dorp, R. & Mazzuchi, T.A. (2004a). A general Bayes Weibull inference model for accelerated life testing, *Reliability Engineering and System Safety* **90**, 140–147.
- [5] van Dorp, R. & Mazzuchi, T.A. (2004b). A general Bayes exponential inference model for accelerated life testing, *Journal of Statistical Planning and Inference* **119**, 55–74.
- [6] Nelson, W.B. (1990). *Accelerated Testing*, John Wiley & Sons.
- [7] Nelson, W.B. (2005a). A bibliography of accelerated test plans, *IEEE Transactions on Reliability* **54**, 194–197.
- [8] Nelson, W.B. (2005b). A bibliography of accelerated test plans part II-references, *IEEE Transactions on Reliability* **54**, 370–373.
- [9] Martz, H.F. & Waterman, M.S. (1978). A Bayesian model for determining the optimal test stress for a single test unit, *Technometrics* **20**, 179–185.

- [10] DeGroot, M.H. & Goel, P. (1979). Bayesian estimation and optimal designs in partially accelerated life testing, *Naval Research Logistics Quarterly* **26**, 223–235.
- [11] Menzefricke, U. (1991). Designing accelerated life tests when there is type II censoring, Unpublished Technical report, University of Toronto.
- [12] Chaloner, K. & Larntz, K. (1992). Bayesian design for accelerated life testing, *Journal of Statistical Planning and Inference* **33**, 245–259.
- [13] Polson, N.G. (1993). A Bayesian perspective on the design of accelerated life tests, in *Advances in Reliability*, A.P. Basu, ed, North Holland, pp. 321–330.
- [14] Verdinelli, I., Polson, N.G. & Singpurwalla, N.D. (1993). Shannon information and Bayesian design for prediction in accelerated life testing, in *Reliability and Decision Making*, R.E. Barlow, C.A. Clariotti & F. Spizzichino, eds, Chapman & Hall, pp. 247–256.
- [15] Soyer, R. & Vopatek, A.L. (1995). Adaptive Bayesian designs for accelerated life testing, in *Adaptive Designs, Lecture Notes, Vol. 25*, N. Flournoy & W.F. Rosenberger, eds, Institute of Mathematical Statistics, pp. 263–275.
- [16] Erkanli, A. & Soyer, R. (2000). Simulation-based designs for accelerated life tests, *Journal of Statistical Planning and Inference* **90**, 335–348.
- [17] Zhang, Y. & Meeker, W.Q. (2006). Bayesian methods for planning accelerated life tests, *Technometrics* **48**, 49–60.
- [18] Chaloner, K. (1984). Optimal Bayesian experimental designs for linear models, *Annals of Statistics* **12**, 283–300.
- [19] Vopatek, A.L. (1992). *Design of accelerated life tests: a Bayesian approach*, Doctoral Dissertation, The George Washington University.
- [20] Muller, P. & Parmigiani, G. (1995). Optimal design via curve fitting of Monte Carlo experiments, *Journal of the American Statistical Association* **90**, 503–510.
- [21] Muller, P. (1999). Simulation-based optimal design, in *Bayesian Statistics*, J.O. Berger, J.M. Bernardo, A.P. Dawid & A.F.M. Smith, eds, Oxford University Press, Vol. 6, pp. 459–474.
- [22] Lindley, D.V. (1985). *Making Decisions*, John Wiley & Sons, London.
- [23] Chaloner, K. & Verdinelli, I. (1995). Bayesian experimental design: a review, *Statistical Science* **10**, 273–304.
- [24] Gilks, W.R., Richardson, S. & Spiegelhalter, D.J. (1996). *Markov Chain Monte Carlo in Practice*, Chapman & Hall, London.
- [25] Smith, A.F.M. & Gelfand, A.E. (1992). Bayesian statistics without tears: a sampling-resampling perspective, *The American Statistician* **46**, 84–88.
- [26] Erkanli, A., Soyer, R. & Angold, A. (1998). Optimal Bayesian two-phase designs, *Journal of Statistical Planning and Inference* **66**, 171–191.
- [27] Chambers, J.M. & Hastie, T.J. (1992). *Statistical Models in S*, Wadsworth.

Related Articles

Accelerated Life Models; Accelerated Life Tests: Classical Methods for Design and Analysis; Accelerated Life Tests: Bayesian Models; Accelerated Life Tests: Designs Comparison within a Bayesian Framework.

REFIK SOYER

Accelerated Life Tests: Stress Step (Classical Methods)

Introduction to Accelerated Life Testing

The failure time data are used for product reliability estimation. Failures of highly reliable units are rare and information other than traditional censored failure time data should be used. One way of obtaining this complementary reliability information is to apply the methods of the so-called accelerated life testing (ALT) for improvements in reliability life testing, to obtain the desired reliability results for the products more quickly than under normal operating conditions, given in terms of step stresses. Step-stress ALT provides an interesting approach to study the possibility to take into account the cumulative effect of the applied stresses on aging, fatigue, and **degradation** of testing items, systems, and so on. The idea of using step stresses in ALT was proposed at first by Miner [1], Pieruschka [2], Sedyakin [3], and Singpurwalla [4]. Later this idea was actively developed in publications of many statisticians, see, for example [5–28], and so on, where one can find a list of references on applications and estimation in ALT.

In ALT any *accelerated life model* is defined on a set $E = \{x(\cdot)\}$ of all possible or admissible stresses for considered problem, where any stress (explanatory variable) $x(\cdot) \in E$ is a deterministic (or random) time function:

$$x(\cdot) = (x_1(\cdot), \dots, x_m(\cdot))^T : [0, \infty) \rightarrow \mathbf{R}^m \quad (1)$$

$x_i(\cdot)$ is univariate explanatory variable. If $x(\cdot)$ is constant in time, $x(\cdot) \equiv x = \text{const}$, we shall write x instead of $x(\cdot)$ in all formulas.

Let E_1 denote a set of constant-in-time stresses, $E_1 \subset E$. Using a set E_1 one may construct a set $E_2, E_2 \subset E$ of step stresses of the form

$$x(u) = \begin{cases} x_1, & 0 \leq u < t_1 \\ x_2, & u \geq t_1 \end{cases} \quad (2)$$

for any $t_1 > 0$, where $x_1, x_2 \in E_1$. Similarly for any $k, k = 2, 3, \dots$ one may consider a class $E_k, E_k \subset E$,

of step stresses of the form

$$x(u) = \begin{cases} x_1, & 0 \leq u < t_1 \\ x_2, & t_1 \leq u < t_2 \\ \dots & \dots \\ x_k, & t_{k-1} \leq u < t_k \leq \infty \end{cases} \quad (3)$$

where $x_1, \dots, x_k$ are constant-in-time stresses from E_1 . It is clear that $E_1 \subset E_2 \subset \dots \subset E_k \subset E$, where E is a set of all admissible stresses.

The most-used stresses in ALT are step stresses from the classes E_2, E_3 , or E_4 , and so on, which we have just defined. The units are placed on test at an initial low constant stress x_1 and if they do not fail in a predetermined time t_1 , the stress is increased. If they do not fail in a predetermined time $t_2 > t_1$, the constant stress x_2 is increased once more, and so on.

Let $T_{x(\cdot)}$ be the failure time under $x(\cdot) \in E$. Then the cumulative distribution function, the hazard rate function, and the cumulative hazard function given $x(\cdot) \in E$ are as follows:

$$\begin{aligned} F_{x(\cdot)}(t) &= \mathbf{P}\{T_{x(\cdot)} \leq t \mid x(s) : 0 \leq s \leq t\}, \\ \lambda_{x(\cdot)}(t) &= \frac{F'_{x(\cdot)}(t)}{1 - F_{x(\cdot)}(t)}, \quad \Lambda_{x(\cdot)}(t) = \int_0^t \lambda_{x(\cdot)}(s) ds \end{aligned} \quad (4)$$

from which one can see their dependence on the *life history up to time t*. Denote by $\lambda_0(t)$ a *baseline rate function*. It corresponds to the failure rate of given items when hypothetically all of them are testing in the same “standard”, “normal”, or “usual” conditions, given by a stress $x_0(\cdot), x_0(\cdot) \in E$.

The two most-known *general* models of ALT with time-varying stresses are the *accelerated failure time* (AFT) model and the *proportional hazards* (PH) or **Cox model**.

The *AFT model* on E is

$$F_{x(\cdot)}(t) = F_0 \left(\int_0^t r[x(\tau)] d\tau \right), \quad x(\cdot) \in E \quad (5)$$

where $r(\cdot)$ is a positive function on E .

In practice, often a baseline cumulative distribution function F_0 is taken from some class of simple parametric distributions, such as Weibull, lognormal, log-logistic, and so on, and in this case we obtain the parametric AFT model. We have an interesting family when F_0 belongs to the power generalized Weibull (PGW) family of distributions (see [21]). In terms of

2 Accelerated Life Tests: Stress Step

the cumulative distribution functions the PGW family is given by the next formula:

$$F(t, \sigma, \nu, \gamma) = 1 - \exp \left\{ 1 - \left[1 + \left(\frac{t}{\sigma} \right)^\nu \right]^{\frac{1}{\gamma}} \right\},$$

$$t > 0, \gamma > 0, \nu > 0, \sigma > 0 \quad (6)$$

If $\gamma = 1$ we have the Weibull family of distributions. If $\gamma = 1$ and $\nu = 1$, we have the exponential family of distributions. The PGW family of distributions has very nice probability properties. All **moments** of this distribution are finite. Depending on the parameter values the hazard rate can be *constant*, *monotone* (increasing or decreasing), *unimodal* or $\cap$ shaped, and bathtub or $\cup$ shaped. Another interesting family, the *exponentiated Weibull family* of distributions, was proposed in [29].

The PH model on E is

$$\lambda_{x(\cdot)}(t) = r[x(t)]\lambda_0(t), \quad x(\cdot) \in E \quad (7)$$

The function r in the AFT model (as in the PH model) can be parameterized in the following way:

$$r[x(\tau)] = e^{\beta\varphi(x(\tau))} \quad (8)$$

where φ is some known function of the stress. For example, if $\varphi(x) = x, \ln x, 1/x$, we have generalizations of the log-linear, power rule, and Arrhenius models, respectively (see [9, 15]).

The AFT and PH models are rather restrictive. In the case of the AFT model, the stress changes locally only the scale. Suppose that the PH model holds on E , x_0 is a *design stress*, x_1 is an *accelerated stress* with respect to the stress x_0 , $x_0, x_1 \in E_1$, so that $F_{x_0}(t) \leq F_{x_1}(t)$ for all $t \geq 0$, and let

$$x(t) = \begin{cases} x_1, & 0 \leq t < t_1 \\ x_0, & t \geq t_1 \end{cases} \quad (9)$$

be a step stress. Then in the case of the PH model, $\lambda_{x(\cdot)}(t) = \lambda_{x_0}(t)$ for all $t > t_1$. So if one group of items is tested under the design stress x_0 and the second group is tested at first under an accelerated stress x_1 until the moment t_1 and after this moment is tested under the design stress x_0 , so that both groups are observed under the same design stress x_0 after the moment t_1 , the failure rate after the moment t_1 is the same for both groups. If items are aging, it is not natural to apply the PH model.

AFT and PH Models for Constant and Step Stresses

In terms of the cumulative distribution function, the PH model (7) is written as

$$F_{x(\cdot)}(t) = 1 - \exp \left\{ - \int_0^t r[x(\tau)] d\Lambda_0(\tau) \right\},$$

$$x(\cdot) \in E \quad (10)$$

In the particular case of constant-in-time stress x_j and the step stress (3), the AFT model (5) implies

$$F_{x_j}(t) = F_0(r_j t) \quad (11)$$

and for $t \in [t_{i-1}, t_i)$, ($i = 1, 2, \dots, k$)

$$F_{x(\cdot)}(t) = F_0 \left(\sum_{j=1}^{i-1} r_j \Delta t_j + r_i(t - t_{i-1}) \right) \quad (12)$$

In terms of F_{x_j} ,

$$F_{x(\cdot)}(t) = F_{x_i}(s_{i-1} + t - t_{i-1}) \quad (13)$$

where s_i can be determined recurrently from the equations

$$F_{x_{i+1}}(s_i) = F_{x_i}(s_{i-1} + t_i - t_{i-1}) \quad (14)$$

For the step stress (9) and $k = 2$ the PH model (10) implies

$$\lambda_{x(\cdot)}(\tau) = \begin{cases} \lambda_{x_1}(\tau), & 0 \leq \tau \leq t \\ \rho \lambda_{x_1}(\tau), & \tau > t \end{cases} \quad (15)$$

where $\rho = r(x_1)/r(x_0)$.

In this particular case of step stresses the Cox model is sometimes called the *tapered failure rate* (TFR) model, see [25]. If $k \geq 2$, equation (10) implies $F_{x_j}(t) = 1 - e^{-r_j \Lambda_0(t)}$ and for $t \in [t_{i-1}, t_i)$

$$F_{x(\cdot)}(t) = 1 - \exp \left\{ - \sum_{j=1}^{i-1} r_j (\Lambda_0(t_j) - \Lambda_0(t_{j-1})) \right. \\ \left. - r_i (\Lambda_0(t) - \Lambda_0(t_{i-1})) \right\} \quad (16)$$

In the case when the distribution under the design stress x_0 is Weibull:

$$F_{x_0}(t) = 1 - e^{-t^\delta/\alpha_0}, \quad \Lambda_0(t) = \frac{t^\delta}{r_0 \alpha_0} \quad (17)$$

we have

$$F_{x_i}(t) = 1 - e^{-t^\delta/\theta_i} \quad (18)$$

and for $t \in [t_{i-1}, t_i]$

$$F_{x(\cdot)}(t) = 1 - \exp \left\{ - \sum_{j=1}^{i-1} \frac{t_j^\delta - t_{j-1}^\delta}{\theta_j} - \frac{t^\delta - t_{i-1}^\delta}{\theta_i} \right\} \quad (19)$$

This particular case of the PH model was considered in [30].

Models with Finite Supports

Following [17], we consider generalization of both the AFT and the PH models: generalized proportional hazards (GPH) model. For this purpose we formulate these models in other terms. For any $x(\cdot)$ the random variable $R = \Lambda_{x(\cdot)}(T_{x(\cdot)})$ has the standard exponential distribution with the survival function $S_R(t) = e^{-t}$, $t \geq 0$, and is called the *exponential resource*, the number $\Lambda_{x(\cdot)}(t)$ being the *exponential resource used until the moment t under the stress $x(\cdot)$* . So the failure rate $\lambda_{x(\cdot)}(t)$ is the *rate of exponential resource usage*.

Under the PH model we have

$$\Lambda'_{x(\cdot)}(t) = r[x(t)]\lambda_0(t) \quad (20)$$

i.e., the rate of resource usage is proportional to some baseline rate and the constant of proportionality is a function of a stress.

Let us consider now the GPH model, proposed in [18]. GPH model holds on E if

$$\lambda_{x(\cdot)}(t) = r\{x(t)\}q\{\Lambda_{x(\cdot)}(t)\}\lambda_0(t), \quad x(\cdot) \in E \quad (21)$$

The rate of resource usage at the moment t depends not only on a stress applied at this moment but also on the resources used until t .

The PH model ($q(u) \equiv 1$) and the AFT model ($\lambda_0(t) \equiv \lambda_0 = \text{const}$) are particular cases of the GPH model.

Consider the case when the rate of exponential resource usage with respect to the baseline rate is proportional to a monotone function of the used resource. One of the most natural parameterizations of the function q in this case would be

$$q(u) = e^{\gamma u}, \quad \gamma \in R^1 \quad (22)$$

So we consider the so-called generalized linear proportional hazards (GLPH) model on E :

$$\lambda_{x(\cdot)}(t) = e^{\beta x(t) + \gamma \Lambda_{x(\cdot)}(t)} \lambda_0(t), \quad x(\cdot) \in E \quad (23)$$

This model implies, that if $\gamma \leq 0$, the support of the distribution is $[0, \infty)$, and if $\gamma > 0$, we have the model with *finite support* $[0, sp_{x(\cdot)})$, where $sp_{x(\cdot)}$ verifies the equation

$$\int_0^{sp_{x(\cdot)}} e^{\beta x(u)} d\Lambda_0(u) = 1 \quad (24)$$

In the particular cases of the step stress (equation 3) we have the equation (16) when $\gamma = 0$. In the case $\gamma \neq 0$ and $x(\cdot) \in E_k$, we have for all $t \in [t_{i-1}, t_i]$

$$F_{x(\cdot)}(t) = 1 - \left\{ 1 - \gamma \sum_{j=1}^{i-1} e^{\beta x_j} (\Lambda_0(t_j) - \Lambda_0(t_{j-1})) - \gamma e^{\beta x_i} (\Lambda_0(t) - \Lambda_0(t_{i-1})) \right\}^{1/\gamma} \quad (25)$$

and for the constant stress x_j

$$F_{x_j}(t) = 1 - \{1 - \gamma e^{\beta x_j} \Lambda_0(t)\}^{1/\gamma} \quad (26)$$

If the distribution of time to failure under the design stress x_0 is Weibull, i.e.

$$F_{x_0}(t) = 1 - e^{-t^\delta/\alpha_0}, \quad \lambda_0(t) = \frac{1}{\gamma} e^{-\beta x_0} \left(1 - e^{-\frac{\gamma}{\alpha_0} + \delta} \right) \quad (27)$$

then under the GLPH model (23) and the step stress (equation 3),

$$F_{x(\cdot)}(t) = 1 - \left\{ 1 + \sum_{j=1}^{i-1} e^{\beta(x_j - x_0)} [e^{-\gamma t_j^\delta/\alpha_0} - e^{-\gamma t_{j-1}^\delta/\alpha_0}] + e^{\beta(x_i - x_0)} [e^{-\gamma t^\delta/\alpha_0} - e^{-\gamma t_{i-1}^\delta/\alpha_0}] \right\}^{1/\gamma} \quad (28)$$

Note that if the rate of resource usage depends on used resource (i.e. $\gamma \neq 0$), the distribution of the time to failure under constant-in-time stresses $x_j \neq x_0$ is not Weibull:

$$F_{x_j}(t) = 1 - \{1 - e^{\beta(x_j - x_0)} [1 - e^{-\gamma t^\delta/\alpha_0}]\}^{1/\gamma} \quad (29)$$

Under the GLPH model, distributions of times to failure under constant-in-time stresses are from the same interesting family of distributions if we take

4 Accelerated Life Tests: Stress Step

one of the following two families of distributions K_1 and K_2 . The family K_1 :

$$F_{x_0}(t) = 1 - (1 + \alpha_0 t^\delta)^{1/\gamma}, \quad t \geq 0, \\ (\alpha_0 > 0, \delta > 0, \gamma < 0) \quad (30)$$

If $0 < \delta \leq 1$, the failure rate $\lambda_{x_0}(t)$ is *decreasing*. If $\delta > 1$, then the failure rate is increasing until the moment $t = \left(\frac{\delta-1}{\alpha_0}\right)^{1/\delta}$ and decreasing after this moment, i.e. it has the $\cap$ shape.

The family K_2 :

$$F_{x_0}(t) = 1 - (1 + \alpha_0 t^\delta)^{1/\gamma}, \quad t \in \left[0, \left(-\frac{1}{\alpha_0}\right)^{1/\delta}\right), \\ (\alpha_0 < 0, \delta > 0, \gamma > 0) \quad (31)$$

If $0 < \delta < 1$, the hazard rate is decreasing in the interval $\left[0, \left(\frac{\delta-1}{\alpha_0}\right)^{1/\delta}\right)$ and increasing in the interval $\left(\left(\frac{\delta-1}{\alpha_0}\right)^{1/\delta}, \left(-\frac{1}{\alpha_0}\right)^{1/\delta}\right)$. The hazard rate has the $\cup$ shape. Therefore this family is very appealing. If $\delta > 1$, the hazard rate is *increasing* on $\left(0, \left(-\frac{1}{\alpha_0}\right)^{1/\delta}\right)$.

Suppose that the time-to-failure distribution is from one of these families. If $F_{x_0}(t) \in K_i (i = 1, 2)$ and the GLPH model holds, then

$$\Lambda_0(t) = -\frac{\alpha_0}{\gamma} e^{-\beta x_0} t^\delta \quad \text{and} \\ F_{x_j}(t) = 1 - (1 + \alpha_j t^\delta)^{1/\gamma} \quad (32)$$

where $\alpha_j = \alpha_0 e^{\beta(x_j - x_0)}$. If $F_{x_0} \in K_2$, the support of the distribution is finite: $\left[0, \left(-\frac{1}{\alpha_j}\right)^{1/\delta}\right)$.

If $x(\cdot)$ is a step stress of the form (3) then

$$F_{x(\cdot)}(t) = 1 - \{1 + \alpha_0 \varphi(t, \beta, \delta)\}^{1/\gamma}, \\ 0 < t < sp_{x(\cdot)} \quad (33)$$

where

$$\varphi(t, \beta, \delta) = \begin{cases} \sum_{j=1}^{k^*-1} e^{\beta(x_j - x_0)} (t_j^\delta - t_{j-1}^\delta) + e^{\beta(x_{k^*} - x_0)} (t^\delta - t_{k^*-1}^\delta), & k^* > 1 \\ e^{\beta(x_1 - x_0)} t^\delta, & k^* = 1 \end{cases} \quad (34)$$

where

$$k^* = \left\{ k : \sum_{j=1}^{k-1} e^{\beta(x_j - x_0)} (t_j^\delta - t_{j-1}^\delta) \leq -\frac{1}{\alpha_0}, \right. \\ \left. \sum_{j=1}^k e^{\beta(x_j - x_0)} (t_j^\delta - t_{j-1}^\delta) > -\frac{1}{\alpha_0} \right\} \quad (35)$$

If $\alpha_0 > 0, \gamma < 0$, then $sp_{x(\cdot)} = +\infty$. If $\alpha_0 < 0, \gamma > 0$, the right limit of the support $(0, sp_{x(\cdot)})$ verifies the equation $\varphi(sp_{x(\cdot)}, \beta, \delta) = -1/\alpha_0$, i.e.

$$sp_{x(\cdot)} = \left\{ t_{k^*-1}^\delta - e^{-\beta(x_{k^*} - x_0)} \left[\frac{1}{\alpha_0} + \sum_{j=1}^{k^*-1} e^{\beta(x_j - x_0)} (t_j^\delta - t_{j-1}^\delta) \right] \right\}^{1/\delta} \quad (36)$$

and the failure rate

$$\lambda_{x(\cdot)}(t) = -\alpha_0 \frac{\delta}{\gamma} e^{\beta(x_{k^*} - x_0)} \frac{t^{\delta-1}}{(1 + \alpha_0 \varphi(t, \beta, \delta))}, \\ 0 < t < sp_{x(\cdot)} \quad (37)$$

has the same shape as in the case of constant-in-time stresses (decreasing or $\cap$ shape for the family K_1 , increasing or $\cup$ shape for the family K_2).

Remarks on Estimation

Suppose that all units are tested under stresses from E_k , i.e. they are initially placed on test at x_1 and run until t_1 when the stress is changed to $x_2 > x_1$. At x_2 , testing continues until t_2 when stress is changed to $x_3 > x_2$, and so on, until x_k . Under x_k , testing continues until all remaining units fail or until t_f , whichever comes first.

Denote by D the number of failures during the experiment, $T_{(1)} \leq \dots \leq T_{(D)}$ the observed moments of failures. In the case of a finite support, the following reparameterization can be introduced: $\sigma = sp_{x(\cdot)}, \beta, \gamma, \delta$. Then $\alpha_0 = -1/\varphi(\sigma, \beta, \delta)$. The likelihood function for the family K_1 is given by

the formula

$$L(\alpha_0, \beta, \gamma, \delta) = \left(-\alpha_0 \frac{\delta}{\gamma}\right)^D e^{\beta \sum_{i=1}^D (x_{k(T_{(i)})} - x_0)} \times \prod_{i=1}^D T_{(i)}^{\delta-1} \{1 + \alpha_0 \varphi(T_{(i)}; \beta, \delta)\}^{\frac{1}{\gamma}-1} \times \{1 + \alpha_0 \varphi(t_f; \beta, \delta)\}^{\frac{n-D}{\gamma}} \quad (38)$$

and for the family K_2 by the next one

$$L(\sigma, \beta, \gamma, \delta) = \left(\frac{\delta}{\gamma}\right)^D e^{\beta \sum_{i=1}^D (x_{k(T_{(i)})} - x_0)} \left(\frac{1}{\varphi(\sigma, \beta, \delta)}\right)^D \times \prod_{i=1}^n T_{(i)}^{\delta-1} \left(1 - \frac{\varphi(T_{(i)}, \beta, \delta)}{\varphi(\sigma, \beta, \delta)}\right)^{\frac{1}{\gamma}-1} \times \left(1 - \frac{\varphi(t_f, \beta, \delta)}{\varphi(\sigma, \beta, \delta)}\right)^{\frac{n-D}{\gamma}} \mathbf{1}_{\{T_{(D)} < \sigma\}} \quad (39)$$

The estimator for the cumulative distribution function $F_{x_0}(t)$ is

$$\hat{F}_{x_0}(t) = 1 - (1 + \hat{\alpha}_0 t^{\hat{\delta}})^{\frac{1}{\hat{\gamma}}} \quad (40)$$

In the case of the class K_1 , approximate **confidence intervals** for $F_{x_0}(t)$ and other reliability characteristics under the design stress x_0 are obtained by standard ML (maximum-likelihood) methods using the estimated Fisher information matrix. In the case of the class K_2 when the support is finite, rigorous asymptotic theory is to be established. If the time to failure under the design stress x_0 has the Weibull (or another) distribution, then estimation under the GLPH model can be done similarly taking equation (28) (or equation (25) with respective Λ_0) for $F_{x(\cdot)}$ instead of equation (33) in all equations.

References

- [1] Miner, M. (1945). Cumulative damage in fatigue, *Journal of Applied Mechanics* **12**, A159–A164.
- [2] Pieruschka, E. (1961). *Relation between Lifetime Distribution and the Stress Level Causing Failures*, LMSD-800440, Lockheed Missiles and Space Division, Sunnyvale.
- [3] Sedyakin, N. (1966). On one physical principle in reliability theory, *Technical Cybernetics* **3**, 80–87.
- [4] Singpurwalla, N. (1971). Inference from accelerated life tests when observations are obtained from censored data, *Technometrics* **13**, 161–170.
- [5] Mann, N.R., Schafer, R.E. & Singpurwalla, N. (1974). *Methods for Statistical Analysis of Reliability and Life Data*, John Wiley & Sons, New York.
- [6] Sethuraman, J. & Singpurwalla, N. (1982). Testing of hypotheses for distributions in accelerated life testing, *Journal of the American Statistical Association* **77**, 204–208.
- [7] Klein, J. & Basu, A.P. (1981). Weibull accelerated life tests when there are competing cause of failure, *Communications in Statistics: Methods and Theory* **A10**, 2073–2100.
- [8] Shaked, M. & Singpurwalla, N. (1983). Inference for step-stress accelerated life tests, *Journal Statistical Planning and Inference* **7**, 295–306.
- [9] Nelson, W. (1990). *Accelerated Life Testing. Test Plan and Data Analysis*, John Wiley & Sons, New York.
- [10] Nelson, W. (2005). A bibliography of accelerated test plan, *IEEE Transactions on Reliability* **54**, 194–197.
- [11] Nelson, W. (2005). A bibliography of accelerated test plan part II, *IEEE Transactions on Reliability* **54**, 370–373.
- [12] Viertl, R. (1988). *Statistical Methods in Accelerated Life Testing*, Goettingen, Vandenhoeck and Ruprecht.
- [13] Schmoyer, R. (1991). Nonparametric analysis for two-level single-stress accelerated life tests, *Technometrics* **33**, 175–186.
- [14] Singpurwalla, N. (1995). Survival in dynamic environments, *Statistical Science* **1**, 86–103.
- [15] Meeker, W. & Escobar, L. (1998). *Statistical Methods for Reliability Data*, John Wiley & Sons, New York.
- [16] Bagdonavičius, V. & Nikulin, M. (1995). Semiparametric models in accelerated life testing, *Queen's Papers in Pure and Applied Mathematics*, Queen's University, Kingston, Vol. 98.
- [17] Bagdonavičius, V., Gerville-Réache, L. & Nikulin, M. (2002). On parametric inference for step-stresses models, *IEEE Transactions on Reliability* **1**, 27–31.
- [18] Bagdonavičius, V. & Nikulin, M. (1999). Generalized proportional hazards model based on modified partial likelihood, *Lifetime Data Analysis* **5**, 329–350.
- [19] Bagdonavičius, V. (1978). Testing the hypothesis of the additive accumulation of damage, *Probability Theory and its Applications* **23**, 403–408.
- [20] Bagdonavičius, V., Gerville-Réache, L., Nikoulina, V. & Nikulin, M. (2000). Expériences accélérées: analyse statistique du modèle standard de vie accélérée, *Revue de Statistique Appliquée* **3**, 5–38.
- [21] Bagdonavičius, V. & Nikulin, M. (2002). *Accelerated Life Models*, Chapman & Hall/CRC, Boca Raton.

- [22] Bagdonavičius, V., Cheminade, O. & Nikulin, M. (2004). Statistical planning and inference in accelerated life testing using the CHSS model, *Journal of Statistical Planning and Inference* **126**, 535–551.
- [23] Bagdonavičius, V. & Nikulin, M. (1998). On semi-parametric estimation of reliability from accelerated life data, in *Statistical and Probabilistic Models in Reliability*, D. Ionescu & N. Limnios, eds, Birkhauser, Boston, pp. 75–89.
- [24] Bagdonavičius, V. & Nikulin, M. (1999). On nonparametric estimation in accelerated experiments with step-stress, *Statistics* **33**, 349–365.
- [25] Bhattacharyya, G. & Stoejoeti, Z. (1989). A tampered failure rate model for step-stress accelerated life test, *Communications in Statistics: Theory and Methods* **18**, 1627–1643.
- [26] Bolshhev, L. (1976). *Statistical Trials on Fatigue and Longevity*, Unpublished manuscript.
- [27] Gerville-Réache, L. & Nikulin, M. (2006). On statistical modelling in accelerated life testing, *Maintenance and Reliability* **30**, 48–51.
- [28] Nikulin, M., Gerville-Réach, L. & Couallier, V. (2007). *Statistique des essais accélérés*, Hermes, London.
- [29] Mudholkar, G. & Srivastava, D. (1995). The exponentiated Weibull family: a reanalysis of the bus-motor-failure data, *Technometrics* **37**, 436–445.
- [30] Khamis, I. & Higgins, J. (1998). A new model for step-stress testing, *IEEE Transactions on Reliability* **47**, 131–134.

Further Reading

- Lin, D.Y. & Ying, Z. (1995). Semiparametric inference for accelerated life model with time dependent covariates, *Journal of Statistical Planning and Inference* **44**, 47–63.
- Lawless, J.F. (2003). *Statistical Models and Methods for Lifetime Data*, John Wiley & Sons, New York.

Related Articles

Accelerated Life Models; Accelerated Life Tests: Classical Methods for Design and Analysis; Accelerated Life Tests: Bayesian Design; Accelerated Life Tests: Designs Comparison within a Bayesian Framework; Accelerated Life Tests: Nonparametric Approach; Accelerated Life Tests: Analysis with Competing Failure Modes; Prediction of Expected Fatigue Lives of Fiber Reinforced Plastic Joints; Accelerated Life Tests: Bayesian Models.

MIKHAIL NIKULIN

Accelerated Life Tests: Designs Comparison within a Bayesian Framework

Introduction

Accelerated life tests (ALT) are used to reduce the time needed to obtain information related to life characteristics of an item, material, or part of interest. In such tests, some characteristics of the use environment (e.g., temperature, voltage, vibration, etc.) are selected and applied to the items at higher (accelerated) levels with the intent of inducing failures in a more rapid fashion. The two main avenues of research in ALT are (a) the **optimal design** of the test and (b) the statistical inference from the failure data obtained. The importance of the design of these tests is paramount to the process of obtaining meaningful inferences about the life of the items in their use environment. Statistical inference is based on one of two paradigms, classical or Bayesian, and often makes use of a time transformation function, which relates failure times at accelerated stress levels to failure times at use stress levels through the scale parameter of the lifetime distribution. Examples of such functions are the power law model, the Eyring model, and the Arrhenius model (see for example Mann *et al.* [1]).

Although there is a host of literature on ALT design, there has been a significant increase in the topic of step-stress accelerated life tests (SSALT) and their comparison to regular ALTs from a classical viewpoint (see for example, Miller and Nelson [2] and more recently Bai and Chung [3], Khamis and Higgins [4], and Khamis [5]). Bayesian approaches to ALT design have been considered in DeGroot and Goel [6], who investigated the optimal design of a partially accelerated life test (PALT), and Erkanli and Soyer [7] who developed a Bayesian decision theoretic method for selecting optimal fixed stress ALTs based on minimization of the expected predictive variance of the parameter of interest. To the best of our knowledge Miller and Nelson [2] were among the earliest to claim similarity of asymptotic estimator variance when comparing step-stress testing ALT with fixed stress ALT. To date, however, no one has considered a comparison of SSALTs and constant

stress ALTs using a Bayesian approach. This is necessary, because in ALT the number of data items available is often not large, bringing into question the use of asymptotic values for comparison (which are the basis for most of the classical design comparisons). This notion is illustrated clearly in McSorley *et al.* [8].

In this article we consider a comparison of ALT designs using the Bayesian inference model developed by van Dorp and Mazzuchi [9]. The model assumes that the life length of the test item at any stress is exponentially distributed and that the failure rates of the life lengths are ordered with respect to the severity of the environment. This restricted ordering of the prior and posterior failure rate estimates acts as a nonparametric time transformation function. Two optimality criteria shall be used for ALT design comparisons: the usual preposterior variance and a measure that takes into account both preposterior mean and preposterior variance.

A general likelihood for the ALT design problem is presented in the section titled “A General Likelihood Function for Accelerated Life Tests”. In the section titled “Prior Distribution”, the joint prior distribution of the parameters is introduced and in the section titled “Posterior Moments after a Single Test Item” expressions for posterior characteristics are developed. An example analysis is presented in the section titled “A Comparison of ALT Step Patterns”. In the last section we provide some concluding remarks.

A General Likelihood Function for Accelerated Life Tests

Any statistical inference procedure involves developing a likelihood. The likelihood model developed in this section allows for a comprehensive representation of regular fixed-stress, progressive step-stress, regressive step-stress, and profile step-stress life testing (see for example, Figure 1).

Following the setup of van Dorp *et al.* [10], typically in ALT, an environment can be described by a collection of stress variables and their levels. In this model, it is assumed that a total of K environments $E_1, \dots, E_K$ are preselected as candidate test environments within minimum and maximum design ranges of stresses to ensure that the predominant failure modes at use environments and accelerated environments are the same. It will be

2 ALT: Designs Comparison within a Bayesian Framework

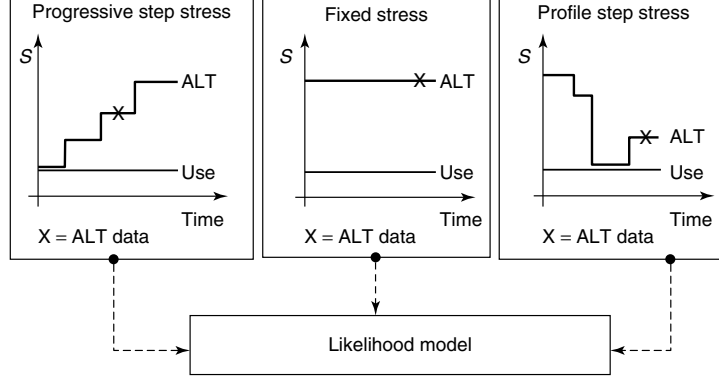


Figure 1 Different ALT scenarios

assumed that these candidate test environments may be rank ordered with respect to severity using engineering knowledge, i.e., E_1 (E_K) coincides with the least (most) severe environment.

The lifetime distribution in a constant stress environment E_e will be modeled as an exponential lifetime distribution with failure rate λ_e , $e = 1, \dots, K$. The ordering of the test environments in terms of severity induces the same ordering in the associated failure rates, i.e.,

$$0 \equiv \lambda_0 < \lambda_1 < \dots < \lambda_K < \lambda_{K+1} \equiv \infty \quad (1)$$

We shall consider a single test item subjected to an ALT with m test intervals $[t_{i-1}, t_i)$, $i = 1, \dots, m$, with $t_0 \equiv 0$, $t_{m+1} \equiv \infty$. If the item has not failed by time t_m , it is removed (censored) from testing. As testing may proceed in a variety of step patterns, the actual test environment in $[t_{i-1}, t_i)$ will be denoted by E_{a_i} , $a_i \in \{1, \dots, K\}$. With the above setup, we have for the ALT reliability function during the test at time $t \in [t_{i-1}, t_i)$ equation (2) and for the ALT failure density or likelihood of a failure at time t during the test equation (3).

$$R(t|\lambda_1, \dots, \lambda_K) = \begin{cases} \exp\{-\lambda_{a_1} t\} & t \in [0, t_1) \\ \exp\left\{-\sum_{j=1}^{i-1} \lambda_{a_j} (t_j - t_{j-1}) - \lambda_{a_i} (t - t_{i-1})\right\} & t \in [t_{i-1}, t_i), \\ \exp\left\{-\sum_{j=1}^m \lambda_{a_j} (t_j - t_{j-1})\right\} & t > t_m \end{cases} \quad (2)$$

Figure 2 presents a comparison of the ALT reliability and density functions for a fixed stress and profile ALT test. The solid lines in Figure 2(a) and (b) (denoted 213) display the ALT reliability and density functions for a profile ALT test as defined by equation (4).

$$\begin{aligned} f(t|\lambda_1, \dots, \lambda_K) &= -\frac{d}{dt} R(t|\lambda_1, \dots, \lambda_K) \\ &= \begin{cases} \lambda_{a_1} R(t|\lambda_1, \dots, \lambda_K) & t \in [0, t_1) \\ \lambda_{a_i} R(t|\lambda_1, \dots, \lambda_K) & t \in [t_{i-1}, t_i), i = 2, \dots, m \\ 0 & t > t_m \end{cases} \end{aligned} \quad (3)$$

$$a_1 = 2, a_2 = 1, a_3 = 3$$

$$t_1 = 1, t_2 = 2, t_3 = 3$$

$$\lambda_1 = \frac{1}{2}, \lambda_2 = 1, \lambda_3 = \frac{1}{2} \quad (4)$$

Note that the equation (4) coincides with a profile ALT test since the environment in the first step is 2, in the second step it equals 1, and finally in the third step it equals 3. The solid line ALT reliability function in Figure 2(a) is continuous but not differentiable at the endpoint of each test interval. The solid line ALT density in Figure 2(b) is discontinuous and exhibits here jumps up and down. The dashed line in Figure 2(a) and (b) (indicated as 222) depict the ALT reliability and density functions of a fixed stress ALT with $a_i = 2$, $i = 1, \dots, 3$ and thus $\lambda_{a_i} = 1$, $i = 1, \dots, 3$.

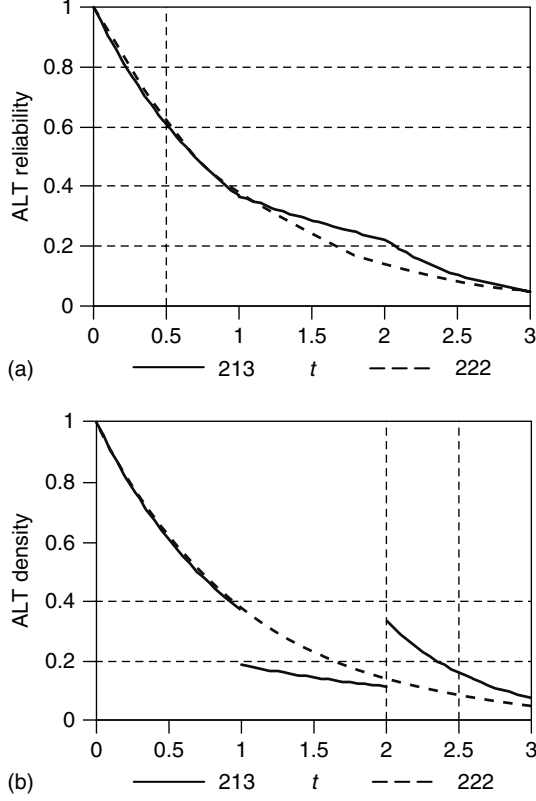


Figure 2 ALT reliability function and ALT density function for a profile ALT test defined by equation (4) as a function of time t .

From equation (2) and the unit step lengths we obtain for the reliability at the end of the ALT test (i.e., the survival probability) for both ALT setups in Figure 2:

$$R(3) = \exp\left(-\sum_{i=1}^3 \lambda_{a_i}\right) = \exp(-3) \approx 0.05 \quad (5)$$

Hence, both ALT setups in Figure 2 are comparable from this vantage point. The immediate question that arises is whether there is any benefit with profile stepping (or step-stress testing in general) compared to a fixed stress ALT setup. In this article we shall compare different ALT step patterns by sampling different values of λ_i , $i = 1, \dots, 3$ from a prior distribution to be defined in the next section. For consistency in comparison, we shall evaluate posterior use stress statistics keeping the

ALT survival probability constant under different values of λ_i . This is achieved by “stretching” or “shrinking” the total test time t_m keeping the length of each test interval equidistant. Introducing the transformation

$$\lambda_e = -\text{Ln}(u_e) \iff u_e = \exp(-\lambda_e) \quad (6)$$

equations (2) and (3) may be expressed in terms of u_e for mathematical convenience (see section titled “Prior Distribution”). Quantities u_e may be interpreted as the reliability after one time unit of exposure to environment e and will henceforth be referred to as *unit reliabilities in the different testing environments*. Substituting equation (6) into equation (2) (and equation (3)), links the reliability of the test item at time t (the likelihood of failure at time t) to the unit reliabilities u_e , $e = 1, \dots, K$, we have, respectively:

$$R(t|u_1, \dots, u_K) = \begin{cases} (u_{a_1})^t & t \in [0, t_1) \\ \left\{ \prod_{j=1}^{i-1} (u_{a_j})^{t_j - t_{j-1}} \right\} (u_{a_i})^{t - t_{i-1}} & t \in [t_{i-1}, t_i), \\ & i = 2, \dots, m \\ \prod_{j=1}^m (u_{a_j})^{t_j - t_{j-1}} & t > t_m \end{cases} \quad (7)$$

$$f(t|u_1, \dots, u_K) = \begin{cases} -\text{Ln}(u_{a_i}) R(t|u_1, \dots, u_K) & t \in [t_{i-1}, t_i), \\ & i = 1, \dots, m \\ 0 & t > t_m \end{cases} \quad (8)$$

The mathematical forms of expressions (7) and (8) are more conducive to posterior updating as compared to expressions (2) and (3) utilizing the prior distribution to be defined in the next section.

Prior Distribution

Denoting Λ_e to be the random variable associated with failure rate realization λ_e and using equations (6) and (1) it follows that

$$0 \equiv u_{K+1} < U_K < \dots < U_1 < u_0 \equiv 1 \quad (9)$$

Rather than defining a prior distribution for $\underline{\Lambda} = (\Lambda_1, \dots, \Lambda_K)$ exhibiting property (equation (1)) for all realizations $\underline{\lambda} = (\lambda_1, \dots, \lambda_K)$, one may equivalently define a prior for $\underline{U} = (U_1, \dots, U_K)$ exhibiting property (equation (9)). Concentrating on $\underline{U} = (U_1, \dots, U_K)$, a prior distribution that is (a) mathematically tractable, (b) defined over the region specified in equation (9), and (c) imposes no other restrictions on u , is the multivariate ordered Dirichlet distribution,

$$\Pi\{\underline{u}|\underline{\eta}, \underline{v}\} = \frac{\prod_{e=1}^{K+1} (u_{e-1} - u_e)^{\eta \cdot v_e - 1}}{\mathbb{D}(\underline{\eta}, \underline{v})} \quad (10)$$

where, $\eta > 0$, $v_e > 0$, $e = 1, \dots, K + 1$, $\sum_{i=1}^{K+1} v_i = 1$,

$$\mathbb{D}(\underline{\eta}, \underline{v}) = \frac{\left\{ \prod_{e=1}^{K+1} \Gamma(\eta v_e) \right\}}{\Gamma(\underline{\eta})} \quad (11)$$

and $\Gamma(\cdot)$ is the gamma function satisfying $\Gamma(n + 1) = n\Gamma(n)$.

Using the transformation given by the rhs of equation (6) allows one to take full advantage of the mathematical properties of the ordered Dirichlet Distribution. For a detailed discussion of these properties see van Dorp and Mazzuchi [9]. The joint moment expression for $\underline{U} \sim \Pi\{\underline{u}|\underline{\eta}, \underline{v}\}$ defined by equation (10) follows as

$$E_{\underline{U}} \left[\prod_{e=1}^K (U_e)^{k_e} \middle| \underline{\eta}, \underline{v} \right] = \frac{\Gamma(\underline{\eta})}{\Gamma(\eta v_{K+1})} \times \prod_{e=1}^K \frac{\Gamma(\eta v_{\geq e+1} + k_{\geq e})}{\Gamma(\eta v_{\geq e} + k_{\geq e})} \quad (12)$$

where

$$v_{\geq e+1} = \sum_{i=e+1}^{K+1} v_i \text{ and } k_{\geq e} = \sum_{i=e}^K k_i \quad (13)$$

We have added the subscript $\underline{U} = (U_1, \dots, U_K)$ to emphasize that the expectation in equation (12) is taken with respect to the prior distribution (10).

Posterior Moments after a Single Test item

In this section we shall derive a closed form expression for posterior use stress moments of order k , for $k \geq 1$, given the failure or survival of the ALT test item up to time t . These derivations employ (a) the joint moments expressions (12), (b) the mathematical form of $R(t|u_1, \dots, u_K)$ defined by equation (7), (c) the fact that

$$\begin{aligned} & \int_0^1 \int_0^{u_1} \dots \int_0^{u_{K-1}} u_1^m f(t|u_1, \dots, u_K) \\ & \times \Pi\{\underline{u}|\underline{\eta}, \underline{v}\} du_K \dots du_1 \\ & = -\frac{d}{dt} \int_0^1 \int_0^{u_1} \dots \int_0^{u_{K-1}} u_1^m R(t|u_1, \dots, u_K) \\ & \times \Pi\{\underline{u}|\underline{\eta}, \underline{v}\} du_K \dots du_1 \end{aligned} \quad (14)$$

and (d) that for $W(t) = \prod_{e=1}^K \{g_e(t)/h_e(t)\}$ we have

$$\frac{d}{dt} W(t) = W(t) \times \sum_{e=1}^K \left[\frac{\frac{d}{dt} g_e(t)}{g_e(t)} - \frac{\frac{d}{dt} h_e(t)}{h_e(t)} \right] \quad (15)$$

Posterior Moments after Test a Item Survived up to Time t

Applying Bayes theorem allows for derivation of closed form posterior moments of the unit reliability in the use stress failure rate environment, i.e., u_1 . We have

$$\begin{aligned} E_{\underline{U}}[U_1^m | \underline{\eta}, \underline{v}, t] &= \\ & \frac{\int_0^1 \int_0^{u_1} \dots \int_0^{u_{K-1}} u_1^m \Pi\{\underline{u}|t, \underline{\eta}, \underline{v}\} du_K \dots du_1}{\int_0^1 \int_0^{u_1} \dots \int_0^{u_{K-1}} R(t|u_1, \dots, u_K) \Pi\{\underline{u}|\underline{\eta}, \underline{v}\} du_K \dots du_1} \end{aligned} \quad (16)$$

Denoting the amount of time that the test item has been exposed to environment E_e prior to t with $\tau_e(t)$,

one obtains

$$\begin{aligned} \tau_e(t) = & \sum_{j=1}^m (t_j - t_{j-1}) 1_{a_j}(e) \times 1_{[0,t)}(t_j) \\ & + \sum_{j=1}^m (t - t_{j-1}) 1_{a_j}(e) \times 1_{[t_{j-1}, t_j)}(t) \end{aligned} \quad (17)$$

where

$$1_{a_j}(e) = \begin{cases} 1 & \text{if in test interval } j \text{ the test} \\ & \text{environment } a_j \text{ is } e \\ 0 & \text{else} \end{cases} \quad (18)$$

and

$$1_{[0,t)}(t_j) = \begin{cases} 1 & t_j \in [0, t) \\ 0 & \text{else} \end{cases} \quad (19)$$

Employing the mathematical definition of $R(t|u_1, \dots, u_K)$ equation (7) it follows that we may rewrite

$$\begin{aligned} & \int_0^1 \int_0^{u_1} \dots \int_0^{u_{K-1}} u_1^m R(t|u_1, \dots, u_K) \\ & \quad \times \Pi\{\underline{u}|\eta, \underline{v}\} du_K \dots du_1 \\ & = E_U \left[\prod_{e=1}^K (U_e)^{k_e(m,t)} \right] \end{aligned} \quad (20)$$

where

$$k_e(m, t) = \begin{cases} \tau_e(t) + m & e = 1 \\ \tau_e(t) & e > 1 \end{cases} \quad (21)$$

Substitution of equation (20) into equation (16) yields

$$E_U[U_1^m|\eta, \underline{v}, t] = \frac{E_U \left[\prod_{e=1}^K (U_e)^{k_e(m,t)} \middle| \eta, \underline{v} \right]}{E_U \left[\prod_{e=1}^K (U_e)^{k_e(0,t)} \middle| \eta, \underline{v} \right]} \quad (22)$$

where $k_e(m, t)$ is defined by equation (21), $\tau_e(t)$ by equation (17) and the prior joint moments are defined by equation (12).

Posterior Moments after a Test Item Failed at Time t

Analogous to equation (16) and employing equations (14) and (20) it immediately follows that

$$\begin{aligned} E_U[U_1^m|\eta, \underline{v}, t] = & \frac{\int_0^1 \dots \int_0^{u_{K-1}} u_1^m f(t|u_1, \dots, u_K) \Pi\{\underline{u}|\eta, \underline{v}\} du_K \dots du_1}{\int_0^1 \dots \int_0^{u_{K-1}} f(t|u_1, \dots, u_K) \Pi\{\underline{u}|\eta, \underline{v}\} du_K \dots du_1} = \\ & \frac{\frac{d}{dt} \int_0^1 \dots \int_0^{u_{K-1}} u_1^m R(t|u_1, \dots, u_K) \Pi\{\underline{u}|\eta, \underline{v}\} du_K \dots du_1}{\frac{d}{dt} \int_0^1 \dots \int_0^{u_{K-1}} R(t|u_1, \dots, u_K) \Pi\{\underline{u}|\eta, \underline{v}\} du_K \dots du_1} \\ & = \frac{\frac{d}{dt} E_U \left[\prod_{e=1}^K (U_e)^{k_e(m,t)} \middle| \eta, \underline{v} \right]}{\frac{d}{dt} E_U \left[\prod_{e=1}^K (U_e)^{k_e(0,t)} \middle| \eta, \underline{v} \right]} \end{aligned} \quad (23)$$

where $k_e(m, t)$ is defined by equation (21) and $\tau_e(t)$ by equation (17). Utilizing equation (15) we immediately obtain

$$\begin{aligned} & \frac{d}{dt} E_U \left[\prod_{e=1}^K (U_e)^{k_e(m,t)} \middle| \eta, \underline{v} \right] \\ & = E_U \left[\prod_{e=1}^K (U_e)^{k_e(m,t)} \middle| \eta, \underline{v} \right] \\ & \quad \times \sum_{e=1}^K \left[\frac{\frac{d}{dt} \Gamma\{\beta v_{\geq e+1} + k_{\geq e}(m, t)\}}{\Gamma\{\beta v_{\geq e+1} + k_{\geq e}(m, t)\}} \right. \\ & \quad \left. - \frac{\frac{d}{dt} \Gamma\{\beta v_{\geq e} + k_{\geq e}(m, t)\}}{\Gamma\{\beta v_{\geq e} + k_{\geq e}(m, t)\}} \right] \end{aligned} \quad (24)$$

$$k_{\geq e}(m, t) = \sum_{i=e}^K k_i(m, t) \quad (25)$$

To further reduce equation (24), we note that $\frac{d}{dt} \Gamma(t)/\Gamma(t)$ may be recognized as the psi or

digamma function $\psi(\cdot)$ and thus

$$\frac{\frac{d}{dt} \Gamma\{\beta v_{\geq e+1} + k_{\geq e}(m, t)\}}{\Gamma\{\beta v_{\geq e+1} + k_{\geq e}(m, t)\}} = \psi\{\beta v_{\geq e+1} + k_{\geq e}(m, t)\} \sum_{i=e}^K 1_{a_j(i)} \times 1_{[t_{i-1}, t_i)}(t) \quad (26)$$

since

$$\frac{d}{dt} \{\beta v_{\geq e+1} + k_{\geq e}(m, t)\} = \sum_{i=e}^K \frac{d}{dt} k_i(m, t) \quad (27)$$

and finally

$$\frac{d}{dt} k_i(m, t) = \frac{d}{dt} \tau_i(t) = 1_{a_j(i)} \times 1_{[t_{j-1}, t_j)}(t) \quad (28)$$

Example

Let us consider the prior distribution on the unit reliabilities $\underline{U}$ with prior parameters

$$\begin{aligned} \eta &= 5, v_1 = 0.39347, v_2 = 0.23865 \\ v_3 &= 0.14475, v_4 = 0.22313 \end{aligned} \quad (29)$$

These prior parameters settings are chosen for illustration purposes such that:

$$\begin{aligned} E_{\underline{U}}[U_e | \eta, v_{\geq e+1}] &= \exp(-\lambda_e) \\ \text{where } \lambda_1 &= \frac{1}{2}, \lambda_2 = 1, \lambda_3 = 1\frac{1}{2} \end{aligned} \quad (30)$$

These values coincide with the ALT setup in Figure 2. Figure 3 depicts the behavior of $E_{\underline{U}}[U_1 | \eta, \underline{v}, t]$ given by equation (23) for $t < 3$ (i.e., the test item failed before the end of the ALT test) and $E_{\underline{U}}[U_1 | \eta, \underline{v}, t]$ given by equation (22) for $t > 3$ (i.e., the test item survived the ALT test) for both the profile 213 ALT setup and the fixed stress ALT setup such that $\lambda_{a_i} = 2, i = 1, 2, 3$ (indicated by 222 in Figure 3).

Note that the up and down jumps of the solid line in Figure 3 mimic those of the ALT density function $f(t | \lambda_1, \dots, \lambda_3) \equiv f(t | u_1, \dots, u_3)$ in Figure 2(b). Hence, a test item exposed to the ALT 213 setup will have a higher posterior unit reliability at use stress when it fails shortly before $t = 1$ rather than

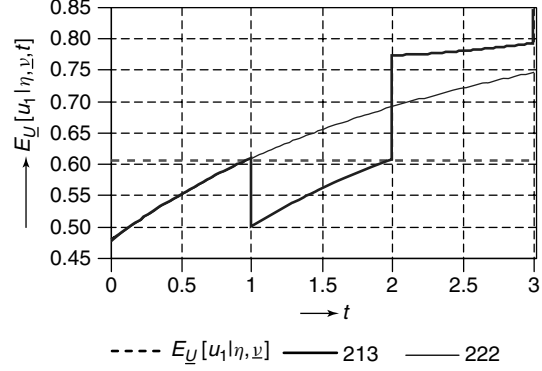


Figure 3 Posterior expected use stress reliability as a function of test failure time and total test time 3 for prior parameter setting equation (29) and the ALT density functions in Figure 2(b)

failing shortly after $t = 1$ due to the jump down in likelihood of failure at $t = 1$ of the ALT 213 density. A reverse behavior can be observed at $t = 2$. On the other hand, the posterior unit reliability always increases smoothly in the case of the fixed stress ALT setup indicated by the line with wide dashes in Figure 3. The horizontal line with dashes depicts the average value of the prior use stress unit reliability i.e.,

$$E_{\underline{U}}[U_1 | \eta, v_{\geq 2}] = \exp(-\lambda_1) = \exp\left(-\frac{1}{2}\right) \approx 0.6065 \quad (31)$$

A Comparison of ALT Step Patterns

In this section we shall first set the stage for what is called a *preposterior analysis* of ALT step patterns in the section titled “A Comparison Analysis of Step Patterns in ALT”, by providing a preliminary comparison of the ALT profile and ALT fixed stress test presented in Figure 2 in the section titled “A Preliminary Comparison of the ALT Setups in Figure 2” We utilize this example to motivate a single measure to compare different step patterns that rewards both calibration and precision.

A Preliminary Comparison of the ALT Setups in Figure 2

Owing to equation (30), via which the prior parameters equation (29) are chosen, it seems intuitive to

expect that given a random failure time T_{λ^*} , where T_{λ^*} follows either of the ALT step patterns 213 and 222 defined in Figure 2(b) with $\lambda^* = (\frac{1}{2}, 1, 1\frac{1}{2})$, average posterior unit reliability should be equal or close to prior unit reliability, i.e.,

$$E_{T_{\lambda^*}}[E_U[U_1|\eta, \underline{v}, T_{\lambda^*}]] \approx E_U[U_1|\eta, v_{\geq 2}] \quad (32)$$

where the value of $E_U[U_1|\eta, v_{\geq 2}]$ is given by equation (31). (It is important here to emphasize the distinction between $E_U[U_1|\eta, \underline{v}, T_{\lambda^*}]$ —which is a posterior mean—and $E_U[U_1, T_{\lambda^*}|\eta, \underline{v}]$ —which is a joint mean. If we were to substitute $E_U[U_1, T_{\lambda^*}|\eta, \underline{v}]$ for $E_U[U_1|\eta, \underline{v}, T_{\lambda^*}]$ in equation (31), equality would hold owing to the law of total probability.) An additional measure to monitor agreement between the prior and the posterior distribution is the posterior variance. An increase in posterior variance compared to prior variance is typically associated, in a Bayesian analysis, with a disagreement between prior information and the observed data. To test the above assertion and observe prior–posterior disagreement, we generate samples $t_i^\circ, i = 1, \dots, 5000$ from both ALT density functions in Figure 2(b) and evaluate:

$$\hat{E}_{T_{\lambda^*}}[E_U[U_1|\eta, \underline{v}, T_{\lambda^*}]] = \frac{1}{5000} \sum_{i=1}^{5000} E_U[U_1|\eta, \underline{v}, t_i^\circ] \quad (33)$$

$$\begin{aligned} & \hat{E}_{T_{\lambda^*}}[\text{Var}_U[U_1|\eta, \underline{v}, T_{\lambda^*}]] \\ &= \frac{1}{4999} \sum_{i=1}^{5000} \text{Var}_U[U_1|\eta, \underline{v}, t_i^\circ] \end{aligned} \quad (34)$$

The results as well as the prior mean and variance $E_U[U_1|\eta, v_{\geq 2}]$, $\text{Var}_U[U_1|\eta, v_{\geq 2}]$ of use stress unit reliability are presented in Table 1. Note that the results for the ALT 213 profile stress test indicate slightly better calibration (since its values in the second row are closer to those in the first row compared to the ALT 222 setup) and higher precision (since the variance in its third row is less than that of the ALT 222 setup). Both ALT scenarios indicate agreement between prior distribution and observed data (on average) since the fourth row entries are reduced by more than 50% of the third row entries.

Table 1 A preliminary comparison of the ALT step patterns in Figure 2 utilizing prior parameter values $\eta, \underline{v}$ provided in equation (29)

	ALT 213	ALT 222
$E_U[U_1 \eta, v_{\geq 2}]$	0.6065	0.6065
$\hat{E}_{T_{\lambda^*}}[E_U[u_1 \eta, \underline{v}, T_{\lambda^*}]]$	0.6025	0.6024
$\text{Var}_U[U_1 \eta, v_{\geq 2}]$	0.0795	0.0795
$\hat{E}_{T_{\lambda^*}}[\text{Var}_U[u_1 \eta, \underline{v}, T_{\lambda^*}]]$	0.0354	0.0364

A test evaluation measure that rewards both calibration and precision may be defined to be

$$\begin{aligned} \mathcal{S}(T_{\lambda^*}, \eta, \underline{v}) = & \left[\left\{ \hat{E}_{T_{\lambda^*}}[E_U[U_1|\eta, \underline{v}, T_{\lambda^*}]] \right. \right. \\ & \left. \left. - E_U[U_1|\eta, v_{\geq 2}] \right\}^2 + \hat{E}_{T_{\lambda^*}}[\text{Var}_U[U_1|\eta, \underline{v}, T_{\lambda^*}]] \right]^{-\frac{1}{2}} \end{aligned} \quad (35)$$

The second term (first term) in equation (34) addresses precision (calibration). Recalling that precision is defined as the *reciprocal of the variance* it follows that higher values of $\mathcal{S}(T_{\lambda^*}, \eta, \underline{v})$ are preferred over lower values. For the results in Table 1, we have

$$\mathcal{S}_{213}(T_{\lambda^*}, \eta, \underline{v}) \approx 5.31 \text{ and } \mathcal{S}_{222}(T_{\lambda^*}, \eta, \underline{v}) \approx 5.28 \quad (36)$$

and on the basis of the evaluation measure defined by equation (36) we prefer the profile step-stress over the fixed stress one.

However, the comparison in equation (36) can only be considered a preliminary one (and not a meaningful one from Bayesian perspective) since it assumed perfect knowledge about failure rates $\lambda^* = (\frac{1}{2}, 1, 1\frac{1}{2})$ (and thus indirectly also about unit reliabilities u). This allowed to sample from the random variable T_{λ^*} and subsequently compare the known value of use stress unit reliability $u_1 = \exp(-\lambda_1) = \exp(-\frac{1}{2}) \approx 0.6065$ (with a value here equal to $E_U[U_1|\eta, v_{\geq 2}]$ in equation (31) due to the choice of prior parameters equation (29)) to posterior estimates $E_U[u_1|\eta, \underline{v}, T_{\lambda^*}]$.

A comparison Analysis of Step Patterns in ALT

Instead of assuming a fixed value $\lambda^* = (\frac{1}{2}, 1, 1\frac{1}{2})$ we shall sample failure rate vectors $\lambda^i, i = 1, \dots, 1000$

indirectly, by sampling vectors of unit reliabilities $\underline{u}^i, i = 1, \dots, 1000$ from the prior distribution equation (10) with the same prior parameters equation (29) and utilizing the transformations equation (6). For each realization $\underline{\lambda}^i$, we “stretch” the length of time t_m of the ALT test, to ensure the same ALT survival probability of ≈ 0.05 (see equation (5)) of the ALT densities in Figure 2 (which terminate at $t_m = 3$). Next, we sample a 1000 failure times of the random variable $T_{\underline{\lambda}^i}$ with a predefined ALT step pattern and evaluate one realization of the performance measures

$$\hat{E}_{T_{\underline{\lambda}^i}}[E_U[U_1 | \eta, \underline{v}, T_{\underline{\lambda}^i}]], \hat{E}_{T_{\underline{\lambda}^i}}[\text{Var}_U[U_1 | \eta, \underline{v}, T_{\underline{\lambda}^i}]] \text{ and } S(T_{\underline{\lambda}^i}, \eta, \underline{v}) \quad (37)$$

defined by equations (33–35). (We reduce the sample size from $T_{\underline{\lambda}^i}$ from 5000 to 1000 to reduce calculation time.) Summarizing, we shall obtain a sample of 1000 values of the statistics indicated by equation (37) associated with a predefined step pattern. Figure 4 summarizes the algorithm to generate these observations.

Figure 5 summarizes the analysis results for two fixed stress ALTs (indicated by 111 and 222), a progressive ALT (indicated by 123), and a profile ALT indicated by 213). The horizontal lines on each subplot indicate the prior values of each performance measure, whereas the values indicate the average value for the iteration series plotted in each subfigure. By comparing Figures A:111 and A:222, we immediately observe a larger sensitivity in the iteration series $i = 1, \dots, 1000$ to failure time data obtained in the use stress environment 1 as compared to the accelerated environment 3.

This larger sensitivity can be explained immediately by the correlations ϱ_{ef} between unit reliabilities U_e and U_f defined by equation (12). We obtain from the prior parameters equation (29) that

$$\varrho_{12} \approx 0.61; \varrho_{13} \approx 0.43; \varrho_{23} \approx 0.70 \quad (38)$$

The **correlation** $\varrho_{12} \approx 0.61$ in equation (38) thus acts as a “dampening” effect of data observed in the environment 2 on posterior updating of use stress unit reliability in environment 1. As a result, the numbers indicated in Figures A:111 and A:222 indicated higher calibration for the higher stress (environment 3) with prior use stress unit reliability ≈ 0.6065 (see equation (31)). This observation is amplified

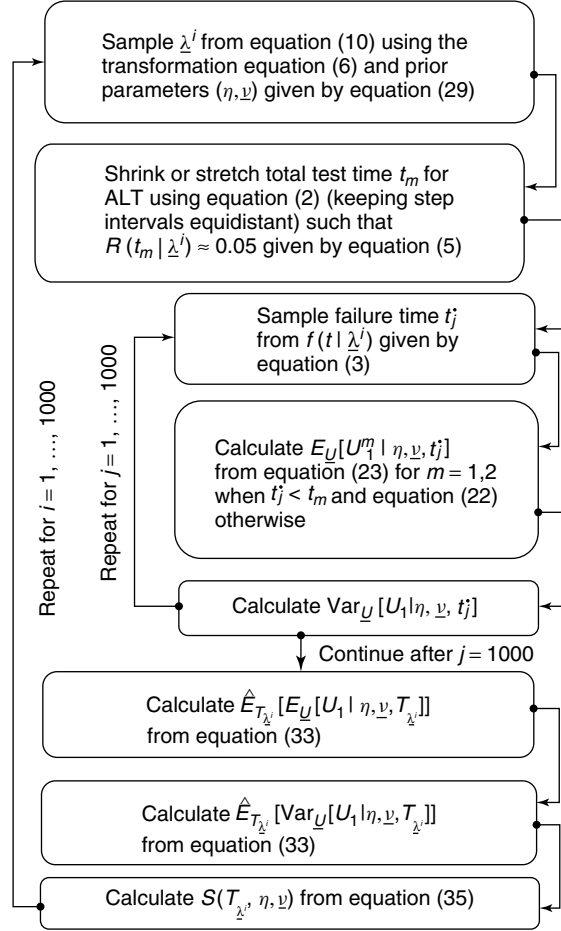


Figure 4 Preposterior analysis algorithm to obtain 1000 values of the test statistics defined by equation (37)

when comparing the Figures B:222 and B:111. A lesser departure from prior variance 0.0795 (see Table 1) is observed in Figure B:222, whereas in Figure B:111 this value is reduced on average by more than 50%. This reduction implies a larger agreement between prior information and observed data in Figure B:111 than in Figure B:222. Both comparisons are combined in Figures C:111 and C:222, where the higher evaluation score 4.207 shows a preference for obtaining data in the use stress environment than in the higher stress environment 2. (Again, this also follows immediately from the correlation coefficient $\varrho_{13} \approx 0.61$ in equation (38).) Keeping ALT survivability constant, however, implies excessive test time at the use stress levels, which is

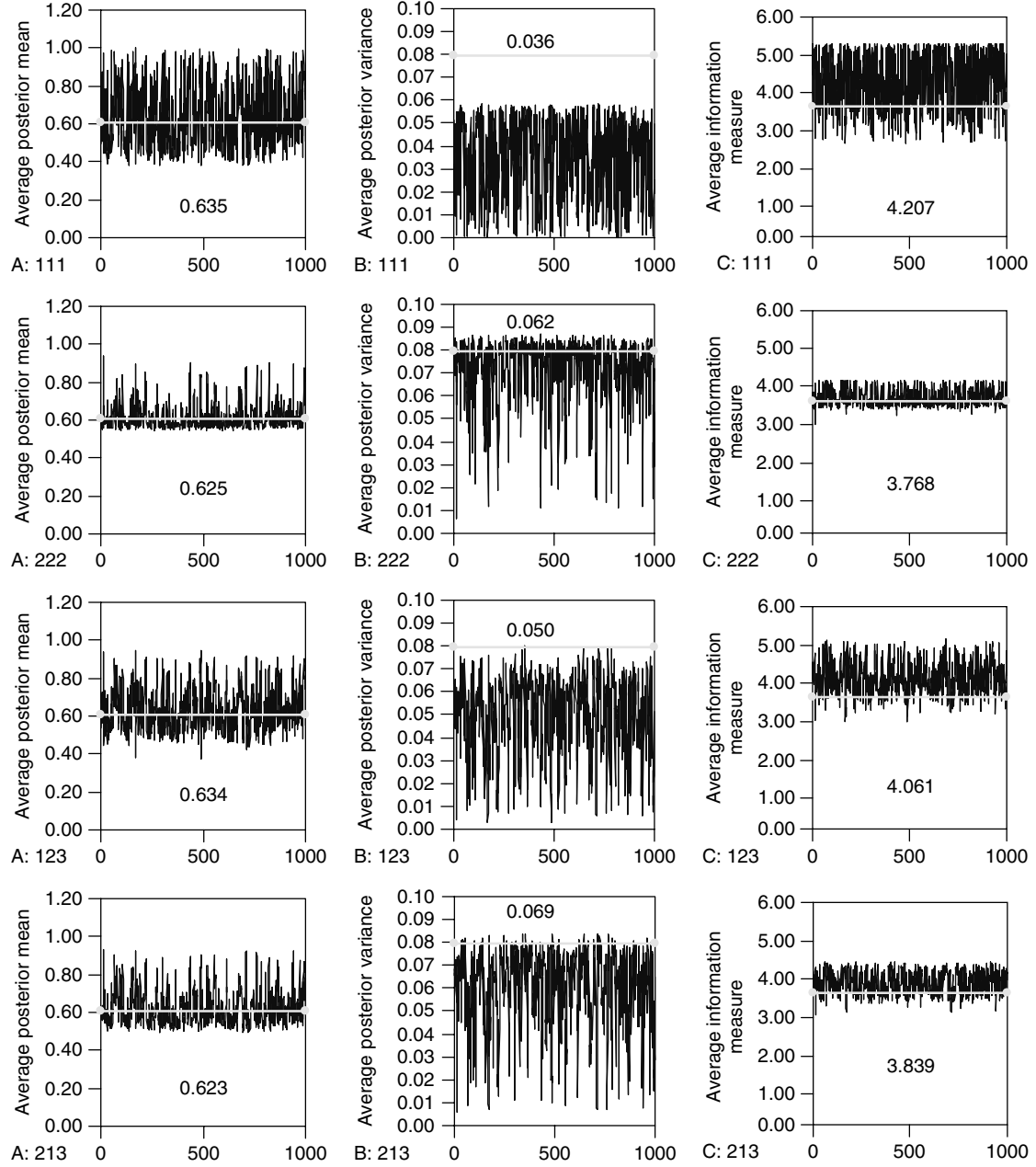


Figure 5 Preposterior analysis of different ALT test patterns

why ALT is being considered in the first place. On the other hand, while testing in a higher stress environment reduces test time, more test results may be needed to confirm or contest specified prior parameters. Differences observed above are similar but

exacerbated for a fixed stress ALT at environment 3 (but are not displayed) owing to the lower correlation $\varrho_{13} \approx 0.43$.

Comparing the progressive step-stress results in Figures A:123, B:123, and C:123 with those in their

“111 and 222” counterparts, it is interesting to note a lesser sensitivity in prior updating than observed in the “111” figures, but a slightly larger one than observed in the “222” figures. In Figure B:123 we still observe a relative large reduction in posterior variance on average (from 0.795 to 0.50) and an average evaluation score of value 4.061 in Figure C:123 (close to the value of use stress testing in Figure C:111). Similar comparisons of the figures associated with the profile ALT setup 213 reveal an improvement of results when compared to fixed stress testing at environment 2 except for the value 0.069 in Figure B:222, but less favorable results when compared to the progressive ALT 123 setup. (Results associated with other profiles are not displayed here, but are worse than the “213” ones and are available upon request.)

Summarizing our results seem to confirm within a Bayesian paradigm a favorable conclusion earlier made by Miller and Nelson [2] in a classical paradigm towards progressive ALT step-stress testing compared to fixed stress ALT. Moreover, our results seem to indicate a preference from an inference point of view toward progressive step-stress testing when compared to any other profile or regressive step patterns.

A Concluding Remark

It is important to recall here that comparisons in the section titled “A Comparison of ALT Step Patterns” are conducted in a context where the ALT survival probability is fixed at approximately 5% regardless of the step patterns and values of the sample failure rates λ^i (which requires “stretching” or “shrinking” the total test time). The conclusion of comparison between fixed stress testing at use stress and at a higher stress level may thus be considered “kicking in an open door”. Of course, we would prefer to test in use stress than in a higher stress from an inference point of view if we had the time to do so. The important observation is, however, that the larger the difference between the higher stress environment and the use stress environment, the larger the “dampening effect” of high-stress environment data on use stress updating (due to an ever smaller correlation between the two). Consequently, a larger amount of data at higher stress levels may be needed to observe similar departures from specified prior use stress information when compared to use stress testing. In

our judgment, this dampening effect is intuitive and realistic.

Such a dampening effect is not present when one utilizes a one-to-one transformation function (such as, for example, the power law) between a higher stress environment failure rate and that of use stress. Essentially, the one-to-one time transformation function “translates” a failure in a high-stress environment to that of use stress (and thus assumes no loss of information by testing at a higher stress level). “Correctness” of any statistical inference procedure that utilizes such a one-to-one time transformation function therefore “stands and falls” with the correctness of the time transformation function used. We are of the opinion that neither the functional form of the transformation function nor its parameters can be known exactly (especially when more than one stress parameter is varied at one time), which is why we propose instead the use of a relaxed ranking assumption of environments in terms of failure rates (see equation (1)) and a subsequent Bayesian analysis as presented in this article.

References

- [1] Mann, N.R., Schafer, R.E. & Singpurwalla, N.D. (1974). *Methods for Statistical Analysis of Reliability and Life Data*, John Wiley & Sons.
- [2] Miller, R.W. & Nelson, W.B. (1983). Optimum simple step-stress plans for accelerated life testing, *IEEE Transactions on Reliability* **R-32**(1), 59–65.
- [3] Bai, D.S. & Chung, S.W. (1992). Optimal design of partially accelerated life tests for the exponential distribution under type-I censoring, *IEEE Transactions on Reliability* **41**(3), 400–406.
- [4] Khamis, I.H. & Higgins, J.J. (1996). Optimum 3-step step-stress tests, *IEEE Transactions on Reliability* **45**(2), 341–345.
- [5] Khamis, I.H. (1997). Comparison between constant and step-stress tests for Weibull models, *International Journal of Quality & Reliability Management* **14**(1), 74–81.
- [6] DeGroot, M.H. & Goel, P.K. (1979). Bayesian estimation and optimal designs in partially accelerated life testing, *Naval Research Logistics Quarterly* **26**, 223–235.
- [7] Erkanli, A. & Soyer, R. (2000). Simulation-based designs for accelerated life tests, *Journal of Statistical Planning and Inference* **90**(2), 335–348.
- [8] McSorley, E.O., Lu, J.C. & Li, C.S. (2002). Performance of parameter-estimates in step-stress accelerated life-test with various sample-sizes, *IEEE Transactions on Reliability* **51**(3), 271–277.

- [9] van Dorp, J.R. & Mazzuchi, T.A. (2004). A general Bayes exponential inference model for accelerated life testing, *Journal of Statistical Planning and Inference* **119**(1), 55–74.
- [10] van Dorp, J.R., Mazzuchi, T.A., Fornell, G.E. & Pollock, L.R. (1996). A Bayes approach to step-stress accelerated life testing, *IEEE Transactions on Reliability* **45**(3), 491–498.

Models; Accelerated Life Tests: Classical Methods for Design and Analysis.

JOHAN RENÉ VAN DORP, THOMAS A.
MAZZUCHI AND JORGE E. GARCIDUEÑAS

Related Articles

Accelerated Life Models; Accelerated Life Tests: Bayesian Design; Accelerated Life Tests: Bayesian

Accelerated Life Tests: Nonparametric Approach

Introduction

Accelerated life testing (ALT) is a method that exposes items to higher stress levels than they expect to receive during a normal use to induce early failures. Typical accelerating stresses include temperature, humidity, voltage, and pressure. These reduced times to failure are beneficial in many situations. For example, to test an item with a mean failure time that is measured in decades, such as a turbine rotor in a hydroelectric dam, under normal conditions would require waiting many years to determine whether the item is sufficiently reliable to be utilized. Even for an item with a much shorter failure time, sometimes it is necessary to accelerate life by testing at stress levels that are higher than normal. For example, for a computer memory chip, the changes in current technologies can be such that it may become obsolete before it is proved reliable. Thus, ALT is now increasingly used in manufacturing industries to reduce the test time and costs. The objective of ALT is either for reliability prediction or to compare one design with another. It is also used to generate failures for failure-mode analysis. Nelson [1] and Bagdonavicius and Nikulin [2] provide a comprehensive review of the theory and practice of ALT.

Parametric models have been introduced to describe the relationship between the distribution functions of lifetimes for different stress levels. The parametric models assume a particular form of the distribution functions at each stress level, and the relationship between them. Nonparametric models for ALT assume a physically reasonable functional relationship between the life distribution functions at the accelerated and nonaccelerated stress levels. However, there are no assumptions on the life distribution function. Parameter transformations are then used to obtain the lifetime distribution under normal stress level. Here, our focus is strictly on nonparametric models. For more see also Basu and Ebrahimi [3], Devarajan and Ebrahimi [4], Jin *et al.* [5], Park and Wei [6], Liero and Liero [7], and references cited there.

Nonparametric Approach to ALT under Multiple Stress Variables

In this section we describe methods for obtaining estimate of life distribution function under normal stress level. Consider a k -dimensional stress vector $\mathbf{S} = (V_1, \dots, V_k)$ applied to an item. Let $F(t|\mathbf{S}_i)$ and $F(t|\mathbf{S}_0)$ be the distribution functions at the stress level $\mathbf{S}_i$ and at normal stress level $\mathbf{S}_0$, respectively. Here $\mathbf{S}_j = (V_{1j}, \dots, V_{kj})$. Our objective is to estimate $F(t|\mathbf{S}_0)$ using observations from accelerated stress levels alone.

Suppose that the relationship between $F(t|\mathbf{S})$ and $F(t|\mathbf{S}_0)$ is

$$F(t|\mathbf{S}) = F(g_{\mathbf{S}_0, \mathbf{S}}(t)|\mathbf{S}_0), \text{ for all } t \geq 0 \quad (1)$$

where $g_{\mathbf{S}_0, \mathbf{S}}(t)$ is called an *acceleration function* and $g_{\mathbf{S}_0, \mathbf{S}}(t) \geq t$ with $g_{\mathbf{S}_0, \mathbf{S}_0}(t) = t$. An acceleration function $g_{\mathbf{S}_0, \mathbf{S}}(t)$ is said to be linear if

$$g_{\mathbf{S}_0, \mathbf{S}}(t) = \alpha(\mathbf{S}_0, \mathbf{S})t, \text{ for all } t \geq 0 \quad (2)$$

In the equation (2), $\alpha(\mathbf{S}_0, \mathbf{S})$ is the acceleration coefficient between $\mathbf{S}$ and $\mathbf{S}_0$. It is important to note that using equations (1) and (2) the acceleration coefficient at the stress level $\mathbf{S}_0$ is 1, that is, $\alpha(\mathbf{S}_0, \mathbf{S}_0) = 1$.

If equation (2) holds, then it is clear that for any two acceleration levels $\mathbf{S}_i$ and $\mathbf{S}_j$ we have

$$F(t|\mathbf{S}_j) = F(\theta_{ij}t|\mathbf{S}_i), \text{ for all } t \geq 0 \quad (3)$$

where $\theta_{ij} = \alpha(\mathbf{S}_0, \mathbf{S}_j)/\alpha(\mathbf{S}_0, \mathbf{S}_i)$ is the relative acceleration coefficient.

In the light of equation (3) one can show that

$$F(t|\mathbf{S}_0) = F\left(\frac{t}{\alpha(\mathbf{S}_0, \mathbf{S})} \middle| \mathbf{S}\right) \quad (4)$$

Throughout this section $\mathbf{S}_1 < \mathbf{S}_2$ means that $\mathbf{S}_1$ is less severe than $\mathbf{S}_2$. That is, $F(t/\alpha(\mathbf{S}_0, \mathbf{S}_1)) \leq F(t/\alpha(\mathbf{S}_0, \mathbf{S}_2))$ for all $t \geq 0$. Also, in our discussion throughout for simplicity, we shall limit ourselves only to two-dimensional stress vectors $\mathbf{S} = (V_1, V_2)$.

We assume that the relation between $F(t|\mathbf{S}_i)$ and $F(t|\mathbf{S}_0)$ is given by the following acceleration function:

$$g_{\mathbf{S}_0, \mathbf{S}_i}(t) = \left(\frac{V_{1i}}{V_{10}}\right)^a \left(\frac{V_{2i}}{V_{20}}\right)^b t \quad (5)$$

2 Accelerated Life Tests: Nonparametric Approach

From the equation (3) the model is thus given by

$$F(t|S_i) = F\left(\left(\frac{V_{1i}}{V_{10}}\right)^a \left(\frac{V_{2i}}{V_{20}}\right)^b t | S_0\right) \quad (6)$$

One way to find a and b in the equation (6) is to generate a linear system of two equations by choosing two different levels (i, j) , (i, k) , $i \neq j \neq k$ and solve the system for a and b . From the equation (2), one can easily show that

$$\theta_{ij} = \frac{\alpha(S_0, S_j)}{\alpha(S_0, S_i)} = \left(\frac{V_{1j}}{V_{1i}}\right)^a \left(\frac{V_{2j}}{V_{2i}}\right)^b \quad (7)$$

and

$$\theta_{ik} = \frac{\alpha(S_0, S_k)}{\alpha(S_0, S_i)} = \left(\frac{V_{1k}}{V_{1i}}\right)^a \left(\frac{V_{2k}}{V_{2i}}\right)^b \quad (8)$$

Taking log on both sides, the two equations become

$$\log \theta_{ij} = a \log \frac{V_{1j}}{V_{1i}} + b \log \frac{V_{2j}}{V_{2i}} \quad (9)$$

$$\log \theta_{ik} = a \log \frac{V_{1k}}{V_{1i}} + b \log \frac{V_{2k}}{V_{2i}} \quad (10)$$

Estimation Procedure When Observations Are Not Censored

An ALT is conducted by selecting $S_1, \dots, S_\ell$. At S_j n_j items are put on test. For item p at the stress level S_j its failure time T_{jp} , $j = 1, \dots, \ell$, $p = 1, \dots, n_j$ is observed. Using this information the point estimates of θ_{ij} and θ_{ik} are

$$\hat{\theta}_{ij} = \frac{\bar{T}_i}{\bar{T}_j}, \quad i \neq j \quad (11)$$

$$\hat{\theta}_{ik} = \frac{\bar{T}_i}{\bar{T}_k}, \quad i \neq k \quad (12)$$

respectively, where $\bar{T}_i = \sum_{p=1}^{n_i} T_{ip}$. Now, using equations (12), we rewrite equations (10) as follows:

$$\log \left(\frac{\bar{T}_i}{\bar{T}_j} \right) = a \cdot \log \left(\frac{V_{1j}}{V_{1i}} \right) + b \cdot \log \left(\frac{V_{2j}}{V_{2i}} \right) \quad (13)$$

$$\log \left(\frac{\bar{T}_i}{\bar{T}_k} \right) = a \cdot \log \left(\frac{V_{1k}}{V_{1i}} \right) + b \cdot \log \left(\frac{V_{2k}}{V_{2i}} \right) \quad (14)$$

Solving the above system for a and b , we obtain

$$\hat{a}_{ijk} = \frac{\log \hat{\theta}_{ij} \cdot \log(V_{2k}/V_{2i}) - \log \hat{\theta}_{ik} \cdot \log(V_{2j}/V_{2i})}{\log(V_{1j}/V_{1i}) \cdot \log(V_{2k}/V_{2i}) - \log(V_{1k}/V_{1i}) \cdot \log(V_{2j}/V_{2i})} \quad (15)$$

$$\hat{b}_{ijk} = \frac{\log \hat{\theta}_{ik} \cdot \log(V_{1j}/V_{1i}) - \log \hat{\theta}_{ij} \cdot \log(V_{1k}/V_{1i})}{\log(V_{1j}/V_{1i}) \cdot \log(V_{2k}/V_{2i}) - \log(V_{1k}/V_{1i}) \cdot \log(V_{2j}/V_{2i})} \quad (16)$$

where $\hat{a}_{ijk}$ and $\hat{b}_{ijk}$, $i, j, k = 1, \dots, \ell$, are estimators of a and b .

These $\ell(\ell-1)(\ell-2)/6$ estimators of a and b from equations (16) are used to form a weighted average that gives the estimators $\hat{a}$ and $\hat{b}$ by

$$\hat{a} = \frac{\sum_{i=1}^{\ell} \sum_{j=i+1}^{\ell} \sum_{k=j+1}^{\ell} w_{ijk} \cdot \left[\log \hat{\theta}_{ij} \cdot \log(V_{2k}/V_{2i}) - \log \hat{\theta}_{ik} \cdot \log(V_{2j}/V_{2i}) \right]}{\sum_{i=1}^{\ell} \sum_{j=i+1}^{\ell} \sum_{k=j+1}^{\ell} w_{ijk}^2} \quad (17)$$

and

$$\hat{b} = \frac{\sum_{i=1}^{\ell} \sum_{j=i+1}^{\ell} \sum_{k=j+1}^{\ell} w_{ijk} \cdot \left[\log \hat{\theta}_{ik} \cdot \log(V_{1j}/V_{1i}) - \log \hat{\theta}_{ij} \cdot \log(V_{1k}/V_{1i}) \right]}{\sum_{i=1}^{\ell} \sum_{j=i+1}^{\ell} \sum_{k=j+1}^{\ell} w_{ijk}^2} \quad (18)$$

respectively, where $w_{ijk} = \log(V_{1j}/V_{1i}) \cdot \log(V_{2k}/V_{2i}) - \log(V_{1k}/V_{1i}) \cdot \log(V_{2j}/V_{2i})$. It is important to note that the set of weights used in equations (18) is one of many different ways that one can assign weights.

Using the estimates $\hat{a}$ and $\hat{b}$, we could transform the failure times from any level to failures at the normal stress level. Observing that T_{ip} under the stress level S_i is stochastically equivalent to observing $(V_{1i}/V_{10})^a \cdot (V_{2i}/V_{20})^b \times T_{ip}$ under the stress level S_0 for $i = 1, \dots, \ell$, we define

$$T'_{ip} = \left(\frac{V_{1i}}{V_{10}}\right)^{\hat{a}} \cdot \left(\frac{V_{2i}}{V_{20}}\right)^{\hat{b}} \cdot T_{ip} \quad (19)$$

Using equation (19) and the fact that $(V_{2i}/V_{20})^b \cdot (V_{1i}/V_{10})^a \times T_{ip}$ and $T - T$ is the lifetime of an item

at the normal stress level—have identical distribution, a consistent estimate of average lifetime at S_0 is given by

$$\begin{aligned}\hat{\mu}_0 &= \frac{1}{N} \sum_{i=1}^{\ell} \sum_{p=1}^{n_i} T'_{ip} \\ &= \frac{1}{N} \sum_{i=1}^{\ell} \sum_{p=1}^{n_i} \left(\frac{V_{1i}}{V_{10}} \right)^{\hat{a}} \cdot \left(\frac{V_{2i}}{V_{20}} \right)^{\hat{b}} \cdot T_{ip} \\ &= \frac{1}{N} \sum_{i=1}^{\ell} n_i \left(\frac{V_{1i}}{V_{10}} \right)^{\hat{a}} \cdot \left(\frac{V_{2i}}{V_{20}} \right)^{\hat{b}} \bar{T}_i\end{aligned}\quad (20)$$

Also a consistent estimate of the life distribution function at S_0 is given by

$$\hat{F}(t|S_0) = \frac{1}{N} \sum_{i=1}^{\ell} \sum_{p=1}^{n_i} I(T'_{ip} \leq t) \quad (21)$$

where

$$\begin{aligned}I(T'_{ip} \leq t) &= 1, \text{ if } T'_{ip} \leq t \\ &= 0, \text{ otherwise}\end{aligned}\quad (22)$$

It should be noted that $N = \sum_{i=1}^{\ell} n_i$.

Estimation Procedure When Observations are Censored

Suppose an ALT is conducted at $S_1, \dots, S_{\ell}$. However, at S_j instead of observing T_{jp} , $p = 1, \dots, n_j$ we observe

$$X_{jp} = \min(T_{jp}, T_{jp}^*) \text{ and } \delta_{jp} = \begin{cases} 1, & T_{jp} \leq T_{jp}^* \\ 0, & \text{otherwise} \end{cases} \quad (23)$$

$p = 1, \dots, n_j$, $j = 1, \dots, \ell$. The random variable T_{jp}^* is referred to as the *censoring time* of item p at S_j , $p = 1, \dots, n_j$, $j = 1, \dots, \ell$. Here T_{jp} and T_{jp}^* are assumed to be independent.

In developing the estimators of a and b , we consider the transformation $W_{jp} = \log X_{jp}$, $j = 1, \dots, \ell$, $p = 1, \dots, n_j$. Then the observations are (W_{jp}, δ_{jp}) , $j = 1, \dots, \ell$, $p = 1, \dots, n_j$. In the presence of the above information, an estimator of survival function of $\log T_j$ at S_j , $j = 1, \dots, \ell$, is given by the

Kaplan–Meier [8] and it is

$$\hat{S}_j^{(1)}(w) = \begin{cases} 1, & w \leq w_{j(1)} \\ 0, & w > w_{j(n_j)} \\ \prod_{p=1}^{c-1} \left(\frac{n_j - p}{j - p + 1} \right)^{\delta_{jp}}, & \text{otherwise,} \end{cases} \quad (24)$$

for $c = 2, \dots, n_j$

Following arguments similar to equations (16) and (18), Basu and Ebrahimi [3] and Klein and Moeschberger [9], one can estimate θ_{ij} and θ_{ik} by

$$\hat{\theta}_{ij}(C) = \exp(\eta_{ij}), \quad i \neq j \quad (25)$$

and

$$\hat{\theta}_{ik}(C) = \exp(\eta_{ik}), \quad i \neq k \quad (26)$$

where η_{ij} and η_{ik} are values of η which minimize

$$g_{ij}(\eta) = \min_{\eta} \int_0^{\infty} [\hat{S}_j^{(1)}(w - \eta) - \hat{S}_i^{(1)}(w)]^2 dw \quad (27)$$

and

$$g_{ik}(\eta) = \min_{\eta} \int_0^{\infty} [\hat{S}_k^{(1)}(w - \eta) - \hat{S}_i^{(1)}(w)]^2 dw \quad (28)$$

respectively. One can show that η_{ij} , which minimizes equation (27) is the median of a discrete random variable with **pdf**:

$$h_{n_i n_j}(v) = \begin{cases} a_c b_d, & v = w_{j(d)} - w_{i(c)}, \text{ for all } c, d \\ 0, & \text{otherwise,} \end{cases} \quad (29)$$

$$a_c \equiv \begin{cases} \hat{S}_i^{(1)}(w_{i(c)}) - \hat{S}_i^{(1)}(w_{i(c+1)}), \\ c = 1, \dots, n_i - 1 \\ \hat{S}_i^{(1)}(w_{i(n_i)}), & c = n_i \end{cases} \quad (30)$$

$$b_d \equiv \begin{cases} \hat{S}_j^{(1)}(w_{j(d)}) - \hat{S}_j^{(1)}(w_{j(d+1)}), \\ d = 1, \dots, n_j - 1 \\ \hat{S}_j^{(1)}(w_{j(n_j)}), & d = n_j \end{cases} \quad (31)$$

Now our estimators of $\hat{a}$ and $\hat{b}$ can be simply obtained by replacing $\hat{\theta}_{ij}$ and $\hat{\theta}_{ik}$ in equation (18) by $\hat{\theta}_{ij}(C)$ and $\hat{\theta}_{ik}(C)$. Here C stands for censored. η_{ik} can be defined similarly.

Now consider the transformation

$$Z_{jp} = \left(\frac{V_{1j}}{V_{10}} \right)^{\hat{a}} \left(\frac{V_{2j}}{V_{20}} \right)^{\hat{b}} X_{jp} \quad (32)$$

4 Accelerated Life Tests: Nonparametric Approach

Table 1 Stress–rupture data of Kevlar/epoxy strands at elevated levels of pressure and temperature

Stress level i	Pressure (psi) V_{1i}	Temperature (°C) V_{2i}	Average lifetime $\bar{T}_i$	Sample size n_i
1	80	100	0.6902	47
2	80	110	0.2537	39
3	85	100	0.1472	47
4	73	120	1.2824	52
5	80	120	0.1639	75
6	85	110	0.0909	84
7	85	120	0.0343	128
8	73	100	4.99	75
9	73	110	1.9800	78
10	67	100	49.8500	55
11	67	110	20.5500	53
12	67	120	7.1100	73

Then, the observations are (Z_{jp}, δ_{jp}) , $j = 1, \dots, \ell$, $p = 1, \dots, n_j$. To estimate distribution function at S_0 , combine all observations, i.e., $U_1 = Z_{11}, U_2 = Z_{12}, \dots, U_N = Z_{\ell n_\ell}$, $\gamma_1 = \delta_{11}, \dots, \gamma_N = \delta_{\ell n_\ell}$ and then use the Kaplan–Meier estimator:

$$\hat{F}(u|S_0) = \begin{cases} 0, & u \leq u(1) \\ 1 - \left[\prod_{j=1}^{m-1} \left(\frac{N-j}{N-j+1} \right)^{\gamma_j} \right], & u \in [u(m-1), u(m)], m = 2, \dots, N \\ 1, & u > u(N) \end{cases} \quad (33)$$

Also, a consistent estimate of the average lifetime at S_0 is $\hat{\mu}_0 = \int_0^\infty (1 - \hat{F}(u|S_0)) du$.

Example

In this section, we illustrate the nonparametric **estimation** described in the previous section using the dataset presented in Table 1.

For this data the normal condition stress level was assumed to be $S_0 = (V_{10}, V_{20}) = (35, 70)$. We use equation (18) to obtain $\hat{a} = 23.0999$ and $\hat{b} = 8.4737$ and then, using the complete dataset, the $N = 806$ rescaled variables T'_{ip} , $i = 1, \dots, 12$, $p = 1, \dots, n_i$, were obtained applying equation (19). A consistent estimate of the mean lifetime μ_0 at the normal stress level S_0 was then computed using equation (20). In this case, $\hat{\mu}_0 = 267.658$.

References

- [1] Nelson, W. (1990). *Accelerated Testing: Statistical Models, Test plans, and Data Analyses*, John Wiley & Sons, New York.
- [2] Bagdonavicius, V. & Nikulin, M. (2001). *Accelerated Life Models*, Chapman & Hall, Boca Raton.
- [3] Basu, A.P. & Ebrahimi, N. (1982). Non-parametric accelerated life testing, *Ieee Transactions on Reliability* **R31**, 432–435.
- [4] Devarajan, K. & Ebrahimi, N. (1998). A non-parametric approach to accelerated life testing under multiple stresses, *Naval Research Logistics* **45**, 629–644.
- [5] Jin, Z., Liu, D.Y., Wei, L.J. & Ying, Z. (2003). Rank based inference for the accelerated failure time model, *Biometrika* **90**, 341–353.
- [6] Park, Y. & Wei, L.J. (2003). *Estimating Subject Specific Survival Functions Under the Accelerated Failure Time Model*, *Biometrika* **90**, 717–723.
- [7] Liero, H. & Liero, M. (2008). *Testing the influence function in accelerated life time models*, Statistical Models and Methods for Biomedical and Technical Systems. f. Vonta, M. Nikulin, N. Limnios and C. Huber eds. Birkhauser, Boston, (Expected release 2008).
- [8] Kaplan, E. L. & Meier, P. (1958). Nonparametric estimation from incomplete observations, *Journal of American Statistical Association* **53**, 457–481.
- [9] Klein, J & Moeshberger, M.L. (2003). *Survival Analysis: Statistical Methods for Censored and Truncated Data*, 2nd edition, Springer-Verlag, New York.

Related Articles

Accelerated Life Models; Accelerated Life Tests; Classical Methods for Design and Analysis; Accelerated Life Tests; Bayesian Design; Accelerated

Life Tests: Stress Step (Classical Methods); Accelerated Life Tests: Designs Comparison within a Bayesian Framework; Accelerated Life Tests: Analysis with Competing Failure Modes; Predic-

tion of Expected Fatigue Lives of Fiber Reinforced Plastic Joints; Accelerated Life Tests: Bayesian Models.

NADER EBRAHIMI

Accelerated Life Tests: Analysis with Competing Failure Modes

Introduction

Accelerated life tests (ALT) are applied on long-life components to obtain failure data in a reasonable time interval, and are performed at higher stress levels that exceed the normal use conditions. The goal is to make inference about the life distribution of the component under use conditions using failure data obtained under more severe environment. To be valid, an accelerated life test must not alter the basic modes and/or mechanism of failure or their relative prevalence. In real experiments, all these problems may occur if the stresses become too high. Many authors took into account the change in failure mechanism, and proposed distributional or Bayesian models that allow all data to be taken into account. However, when more failure modes are present, their relative prevalence, as the stress increases, has not been yet considered, even if this is indicated by accelerated life data.

Multiple failure modes or competing failure modes are often present in industrial accelerated life experiments. Classical and Bayesian models available in the literature usually assume independent failure modes. Little work has been done for dependent competing failure modes, and it resumes to multivariate models with positively correlated lives of failure modes, or to upper and lower bounds for the component life distribution. All these models do not take into account the change in the relative prevalence of the failure modes and the change in the dependence structure as the stress increases.

The purpose of this article is to study the influence of the dependence structure of competing failure modes, as their relative dominance changes with the increase of the applied stress. This is obtained by integrating the results of **competing risk** theory into the ALT procedure, and by using graphical techniques in model selection. Nelson's motorettes data [1] presents a strong change of the relative prevalence of competing failure modes and it will be used as a support application throughout the article.

The section titled "Overview of ALT and Competing Risks" presents several step-stress ALT models, with independent and dependent competing failure modes. Empirical distribution function based on competing risk data are used in the section titled "Graphical Competing Risks Analysis of Motorettes Data" to illustrate the change of the relative dominance of the competing failure modes, and to select the appropriate model to interpret ALT data. An ALT model with dependent failure modes is presented in the section titled "A Copula-Dependent ALT – Competing Risks Model". The dependence structure is modeled by a copula given the rank correlation coefficient, for each stress level. The impact of the independence assumption on the estimated lifetime of the component under use conditions is studied in the section titled "Data Application", with application to two ALT datasets.

Overview of ALT and Competing Risks

We consider an m -constant stress level ALT. At each stress level l , $l = 1, \dots, m$, a number of items are tested until a failure or a censored time occurs. The failure is assumed to occur owing to k competing failure modes, $X_1, X_2, \dots, X_k$. In a competing risks context, we observe the shortest of X_i , $i = 1, \dots, k$, and observe which failure mode it is. In many cases we can reduce the problem to the analysis of two competing risks classes, described by two random variables X_1 and X_2 , and we call X_2 the censoring variable. Competing risk data will only allow us to estimate the subsurvival functions,

$$S_{X_1}^*(t) = P(X_1 > t, X_1 < X_2) \quad (1)$$

and

$$S_{X_2}^*(t) = P(X_2 > t, X_2 < X_1) \quad (2)$$

but not the true survival functions of X_1 and X_2 . Hence, we are not able to estimate the underlying failure distributions without making additional model assumptions.

The conditional subsurvival function is the subsurvival function conditioned on the event that the failure mode in question is manifested. Assuming continuity of $S_{X_1}^*$ and $S_{X_2}^*$ at zero,

$$\begin{aligned} CS_{X_1}^*(t) &= P(X_1 > t, X_1 < X_2 | X_1 < X_2) \\ &= S_{X_1}^*(t) / S_{X_1}^*(0) \end{aligned} \quad (3)$$

2 ALT: Analysis with Competing Failure Modes

$$\begin{aligned} CS_{X_2}^*(t) &= P(X_2 > t, X_2 < X_1 | X_2 < X_1) \\ &= S_{X_2}^*(t)/S_{X_2}^*(0) \end{aligned} \quad (4)$$

Closely related to the notion of the subsurvival functions is the probability of censoring beyond time t ,

$$\begin{aligned} \phi(t) &= P(X_2 < X_1 | \min(X_1, X_2) > t) \\ &= S_{X_2}^*(t)/(S_{X_1}^*(t) + S_{X_2}^*(t)) \end{aligned} \quad (5)$$

This function seems to have some diagnostic value, enabling us to choose the competing risk model that fits the data.

ALT and Independent Competing Risks

The presence of independent competing risks in ALT has been widely studied in the literature. McCool [2] presented a technique for calculating estimate intervals for Weibull parameters of a primary failure mode when a secondary failure mode having the same (but unknown) Weibull shape parameter is acting. Klein and Basu [3, 4] presented the analysis of ALT when more than one failure mode is acting. Assuming independence among competing failure modes for each stress level, the authors obtained maximum likelihood estimators when the lifetimes are exponentially or Weibull – with common- or different-shape parameter – distributed.

Nelson [1] presented graphical and analytical (maximum likelihood) methods to analyze data on a failure mode, estimate a product life distribution when multiple failure modes act, and estimate a product life distribution with certain failure modes eliminated. Nelson indicated that complete datasets are usually analyzed with standard least-squares regression analysis. Such analysis may be misleading for data with competing failure modes. The analysis should consist in a separate Arrhenius model for each failure mode and a series-system model for the relationship between the failure times of each failure mode and the failure time of the component. Examples of products that have multiple causes of failure are given, including insulation systems, ball bearings, and industrial heaters.

A large sample data is considered by Nelson (p. 393), and is collected from a temperature-accelerated life test of motor insulation (the turn, phase, or ground insulation can fail). The experiment was conducted for observing a greater number of failures for each failure mode. When the specimen fails

by a specific failure mode, that insulation is isolated and the motorette is kept on test and runs to a second or third failure. In actual use, the first failure from any cause ends the life of the motor. This fact conducts to a pseudo-competing risk data.

Nelson analysis of motorettes data consists in a separate Arrhenius-lognormal model for each failure mode, and a competing risk (series system) model for the relationship between the failure times of different failure modes and the failure time of a specimen. The assumptions made in the Arrhenius-lognormal model are as follows:

- for temperature T , life has a lognormal distribution;
- the log standard deviation σ_k is constant for all stress levels;
- the mean log life as a function of $x = 1000/T$ is $\mu_k(x) = \alpha_k + \beta_k x$;

where α_k , β_k , and σ_k are the parameters characteristic for the failure mode k , product, and test method. For units run at stress l , the probability that failure mode k survives time t is

$$R_k(t) = \Phi(-(\log(t) - \mu_k(x))/\sigma_k) \quad (6)$$

where Φ is the normal function.

ALT and Dependent Competing Risks

Simple empirical upper and lower bounds of the system life distribution can be found, when the lives of competing failure modes are positively correlated. The lower limit corresponds to the most pessimistic case, when independent failure modes are considered. The upper limit is the lowest life distribution for a single failure mode. These bounds may give some insights with respect to the true distributional form, especially when these bounds are translated to time average failure rate bounds.

Other techniques in dealing with dependent failure modes in ALT are based on multivariate distributions. Multivariate exponential and Weibull distributions – dominant in literature – have properties not suited for all applications.

Klein and Basu [4] proposed a bivariate Weibull model. They considered two competing risks X_1 and X_2 , both Weibull distributed with shape parameter a and scale parameter $\lambda_i(T, \beta_i) + \lambda_{12}(T, \beta_{12})$, $i = 1, 2$, at each temperature (stress) level T ($\beta_1, \beta_2, \beta_{12}$ are

time transformation parameters). The joint survival function of (X_1, X_2) is given by

$$S(x_1, x_2) = \exp(-\lambda_1(T, \beta_1)x_1^a - \lambda_2(T, \beta_2)x_2^a - \lambda_{12}(T, \beta_{12})\max(x_1, x_2)^a) \quad (7)$$

To estimate the parameters of X_1 and X_2 at use conditions, classical methodology of ALT is applied. They illustrated this procedure on the motorette data. However, as shown in the section titled “Graphical Competing Risks Analysis of Motorettes Data”, simple plots of data will reject the applicability of exponential and Weibull multivariate models to this specific dataset.

More elaborate, dependent, competing risk models based on the use of copula have been presented by Hougaard [5] and Bunea and Bedford [6]. These models will be discussed in the section titled “A Copula-Dependent ALT–Competing Risks Model”.

Graphical Competing Risks Analysis of Motorettes Data

Cooke [7] presented several independent and dependent competing risks models. These models have been further investigated by Bedford and Cooke [8], Langseth and Lidqvist [9], and Bunea *et al.* [10]. Nevertheless, the choice of the most appropriate model to fit the competing risk data remains an open debate. All the above authors gave some guidelines for model selection. An important indicator for model selection is the probability of censoring after time t , $\phi(t)$, together with the conditional subsurvival functions $CS_{X_1}^*(t)$ and $CS_{X_2}^*(t)$ (see [10]). Note, that empirical version of these function can be directly obtained from a competing risk dataset.

The model selection guidelines for the most well known competing risk models – independent exponential model, random signs model, conditional independence model, mixture of exponentials model – are as follows:

- If the risks are exponential and independent, then the conditional subsurvival functions are equal and correspond to exponential distributions. Moreover, $\phi(t)$ is constant.
- Under random signs censoring, $\phi(0) > \phi(t)$ and $CS_{X_1}^*(t) > CS_{X_2}^*(t)$ for all $t > 0$.
- If the conditional independence model holds with exponential marginals, then the conditional subsurvival functions are equal and $\phi(t)$ is constant.

- If the mixture of exponentials model holds, then $\phi(t)$ is strictly increasing and $CS_{X_1}^*(t) < CS_{X_2}^*(t)$ for all $t > 0$.

In addition to the above statements, one can easily show that Klein and Basu’s multivariate exponential model requires equal conditional subsurvival function. Simple calculations give the subsurvival and conditional subsurvival functions:

$$S_{X_1}^*(t) = \frac{\lambda_1 + \lambda_{12}}{\lambda_1 + \lambda_2 + \lambda_{12}} \exp(-(\lambda_1 + \lambda_2 + \lambda_{12})t) \quad (8)$$

$$S_{X_2}^*(t) = \frac{\lambda_2 + \lambda_{12}}{\lambda_1 + \lambda_2 + \lambda_{12}} \exp(-(\lambda_1 + \lambda_2 + \lambda_{12})t) \quad (9)$$

$$CS_{X_1}^*(t) = CS_{X_2}^*(t) = \exp(-(\lambda_1 + \lambda_2 + \lambda_{12})t) \quad (10)$$

Note that the conditional subsurvival function should be equal at each stress level to have a valid ALT–competing risks model.

Figure 1 shows the empirical conditional subsurvival functions for motorettes data, with failure modes grouped in two competing risk classes: risk 1 – turn-failure modes and risk 2 – phase- and ground-failure modes. One can see a strong change in the failure mechanism as the stress level increases. At the first stress level the dominant failure mode is risk 2, with its conditional subsurvival function lying entirely above the conditional subsurvival function of risk 1. A mixture of exponentials model may be appropriate for this case. The conditional subsurvival functions are more or less equal for stress level 2 and 3, indicating that an independent exponential model may apply. Data obtained from stress level 4 indicates a dominant conditional subsurvival functions for risk 1, and thus a random signs model may be selected (risk 1 being censored).

The change in the relative prevalence of the failure modes does not allow a classical ALT–competing risk model to analyze data. Since the conditional subsurvival functions are not equal for each stress level, the assumption of independent exponential or Weibull with same shape parameter failure modes is clearly rejected by motor insulation data. On the basis of model selection criteria presented at the beginning of this section, the random signs model and the multivariate exponential (or Weibull with same shape parameter model) are also rejected. The next section will present a more complex dependent model and investigate the influence of the degree of dependence on the estimated distribution functions at use conditions.

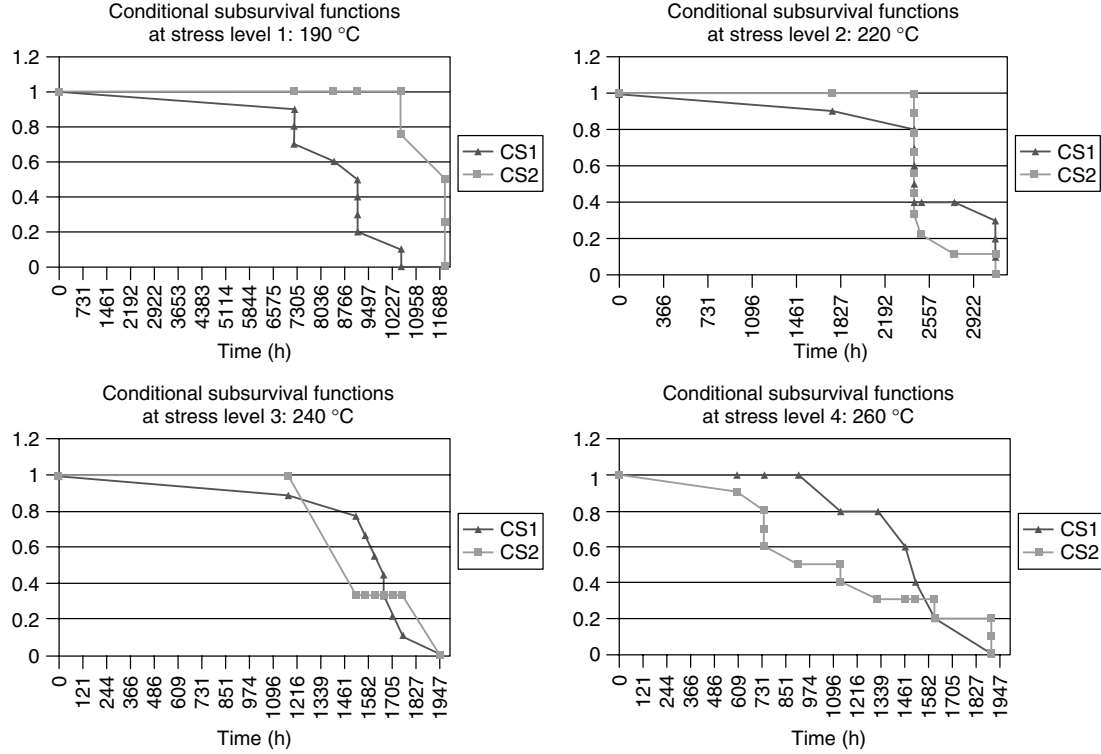


Figure 1 Empirical conditional subsurvival functions for motorettes data, at each temperature (stress) level

A Copula-Dependent ALT–Competing Risks Model

Competing risks and copula

We assume that the dependence structure between X_1 and X_2 is given by a copula. The copula of two random variables X_1 and X_2 with distribution functions $F_{X_1}(X_1)$ and $F_{X_2}(X_2)$, is the distribution C on the unit square $[0, 1]^2$ of the pair $(F_{X_1}(X_1), F_{X_2}(X_2))$. The functional form of $C: [0, 1]^2 \rightarrow \mathbf{R}$ is $C(u, v) = H(F_{X_1}^{-1}(u), F_{X_2}^{-1}(v))$, where H is the joint distribution function of (X_1, X_2) and $F_{X_1}^{-1}$ and $F_{X_2}^{-1}$ are the right-continuous inverses of F_{X_1} and F_{X_2} . Under the assumption of independence of X_1 and X_2 , the marginal distribution functions of X_1 and X_2 are uniquely determined by data. Zheng and Klein [11] showed a more general result that, if the copula of (X_1, X_2) is known, then the marginal distributions functions of X_1 and X_2 are uniquely determined by the competing risk data. More precisely, the marginal distributions functions F_{X_1} and F_{X_2} are solutions

of the following system of ordinary differential equations:

$$\begin{cases} \{1 - C_u(F_{X_1}(t), F_{X_2}(t))\} F'_{X_1}(t) = F_{X_1}^{*'}(t) \\ \{1 - C_v(F_{X_1}(t), F_{X_2}(t))\} F'_{X_2}(t) = F_{X_2}^{*'}(t) \end{cases} \quad (11)$$

with initial conditions $F_{X_1}(0) = F_{X_2}(0) = 0$, where $C_u(F_{X_1}(t), F_{X_2}(t))$ and $C_v(F_{X_1}(t), F_{X_2}(t))$ denote the first order partial derivatives $\frac{\partial}{\partial u} C(u, v)$ and $\frac{\partial}{\partial v} C(u, v)$ calculated in $(F_{X_1}(t), F_{X_2}(t))$. $F_{X_1}^{*'}(t) = S_{X_1}^*(0) - S_{X_1}^*(t)$ and $F_{X_2}^{*'}(t) = S_{X_2}^*(0) - S_{X_2}^*(t)$ are the subdistribution functions of X_1 and X_2 (see [8]).

Measures of Association

Let X and Y be two random variables. There are many **measures of association** for the pair (X, Y) , which are symmetric in X and Y . The best-known measures of association are Kendall's τ and Spearman's ρ .

For the goal of this article, we will henceforth consider only Spearman's ρ . The Spearman's ρ is

defined as the product moment correlation of $F_X(X)$ and $F_Y(Y)$:

$$\begin{aligned}\rho_r(x, y) &= \rho(F_X(X), F_Y(Y)) \\ &= \frac{\text{Cov}(F_X(X), F_Y(Y))}{\sqrt{\text{Var}(F_X(X))\text{Var}(F_Y(Y))}}\end{aligned}\quad (12)$$

The simple formula relating Spearman's ρ to copula density is

$$\rho(X, Y) = 12 \int_0^1 \int_0^1 uv dC(u, v) - 3 \quad (13)$$

Since the measure of association is to be treated as a primary parameter, it is necessary to choose a family of copulae that model all possible measures of association in a simple way. Meeuwissen and Bedford [12] have proposed to use the unique copula with the given Spearman's ρ that has minimum information with respect to the independent distribution, and also they have given a method to calculate this copula numerically. Prior knowledge or subjective information to obtain values for the measure of association is needed. Expert judgment is used to model the uncertainty over the measure of association (see [6]).

Work of Zheng and Klein suggests that the important factor for an estimate of the marginal survival function is a reasonable guess at the strength of the association between competing risks and not the functional form of the copula. For this reason we will choose a class of copula that is easy to work with from the mathematical point of view. A class with such properties is the Archimedean copula.

Archimedean Copula

Let X_1 and X_2 be continuous random variables with joint distribution H and marginal distribution F_{X_1} and F_{X_2} . The Archimedean copula is defined as

$$\varphi(H(x_1, x_2)) = \varphi(F_{X_1}(x_1)) + \varphi(F_{X_2}(x_2)) \quad (14)$$

or in terms of copula

$$\varphi(C(u, v)) = \varphi(u) + \varphi(v) \quad (15)$$

The function φ is called an *additive generator of the copula*. If $\varphi(0) = 1$, φ is a strict generator and $C(u, v) = \varphi^{-1}(\varphi(u) + \varphi(v))$ is a strict Archimedean copula. For our goal we choose a one-parameter

family of copulae, which has a strict generator. The Gumbel family is defined as follows:

$$C_\alpha(u, v) = \exp(-((- \log u)^\alpha + (- \log v)^\alpha)^{1/\alpha}) \quad (16)$$

with parameter $\alpha \in [1, \infty)$. The generator is the function $\varphi_\alpha(t) = (- \log t)^\alpha$. Using the relation between Spearman's ρ and copula (equation (13)), we can obtain α as a function of Spearman's ρ . However, no analytical solution can be found and numerical calculations are needed.

Copula ALT-Competing Risks Model

Let U_1 and U_2 be two random variables associated with the conditional subsurvival functions of X_1 and X_2 . The ALT-competing risk model proposed has two steps: the first step consists of a separate Arrhenius-lognormal model for each random variable U_1 and U_2 ; the second step models the dependence structure of X_1 and X_2 at each stress level and at use conditions by a copula, with the rank correlation coefficient given as a primary parameter.

Graphical and analytical (maximum likelihood) analysis of the Arrhenius-lognormal model (see [1]), applied to the conditional subsurvival functions at each stress level, enables us to find the conditional subsurvival functions at use conditions, and to estimate the distributional parameters. Elicitation of the rank correlation coefficient at each stress level and at use conditions using experts' opinion allows us to fully specify the copula. Further, applying Zheng and Klein's result (equation (11)) one can obtain the distribution functions of X_1 and X_2 at each stress level and at use conditions.

Data Application

We will apply the copula ALT-competing risks model to Nelson's motorettes data. Table 1 presents the estimated parameters of the lognormal distribution (for U_1 and U_2) at each stress level and at use conditions. These distributions, together with the values in point zero of the subsurvival functions of X_1 and X_2 , and with the elicited rank correlation coefficients, constitute input factors in the system of equations (11).

Figure 2 shows the distribution functions of X_1 and X_2 at each stress level. One can see a shift to the left of the plots as the stress increases, indicating a shorter time to failure as the stress is higher.

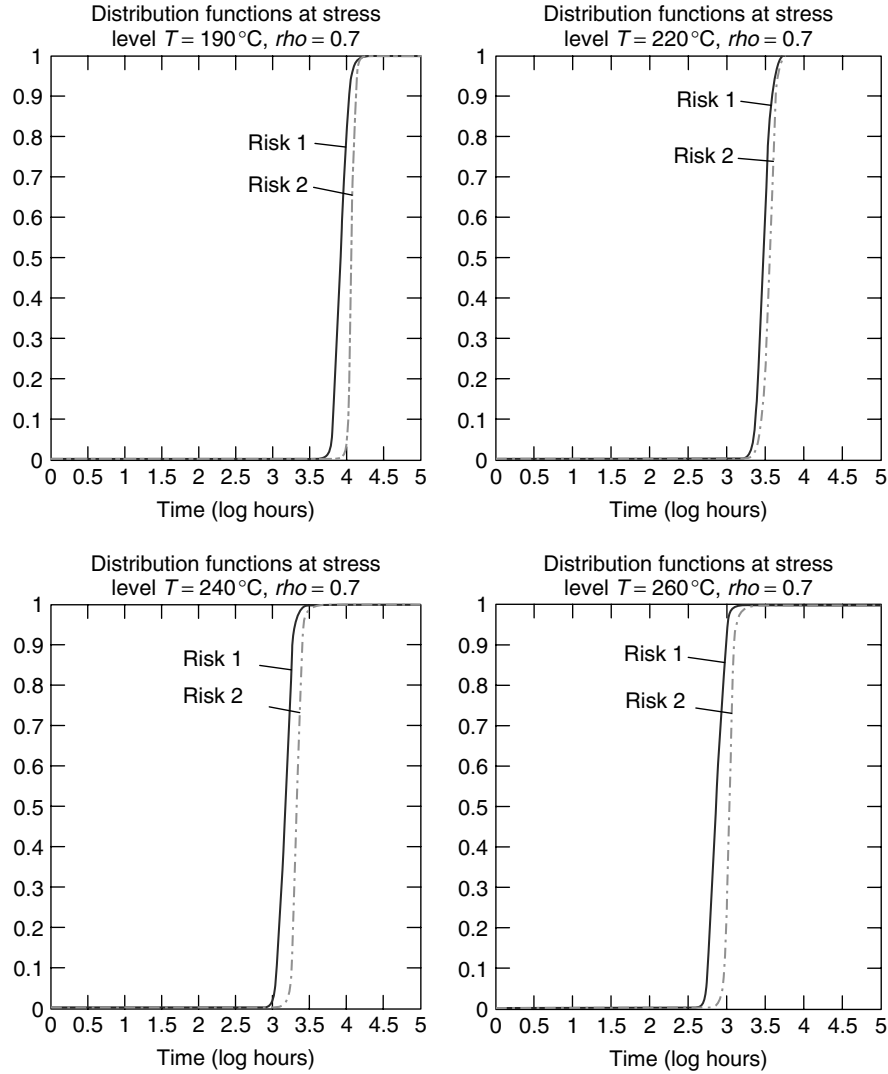


Figure 2 Distribution functions for each failure mode, at each temperature (stress) level (motorettes data)

Table 1 Estimated parameters of the lognormal distribution (for U_1 and U_2) at each stress level and at use conditions (motorettes data)

Temperature (°C)	μ_{U_1}	σ_{U_1}	μ_{U_2}	σ_{U_2}
180	1.4061	0.0201	1.4660	0.0280
190	1.3683	0.0208	1.4165	0.0206
220	1.2453	0.0236	1.2506	0.0321
240	1.1540	0.0259	1.2243	0.1026
260	1.0535	0.0287	1.1508	0.1334

Figure 3 presents the distribution functions of X_1 and X_2 at use condition, for different degree of dependence between failure modes: from independence in the upper left corner to strong dependence in the lower right corner. Even if a logarithmic scale is used, one can see the difference that the degree of dependence makes in estimating the true distribution functions at use conditions.

A more clear visualization of the above results may be obtained by analyzing the industrial heaters data presented by Nelson (p. 420). This data is

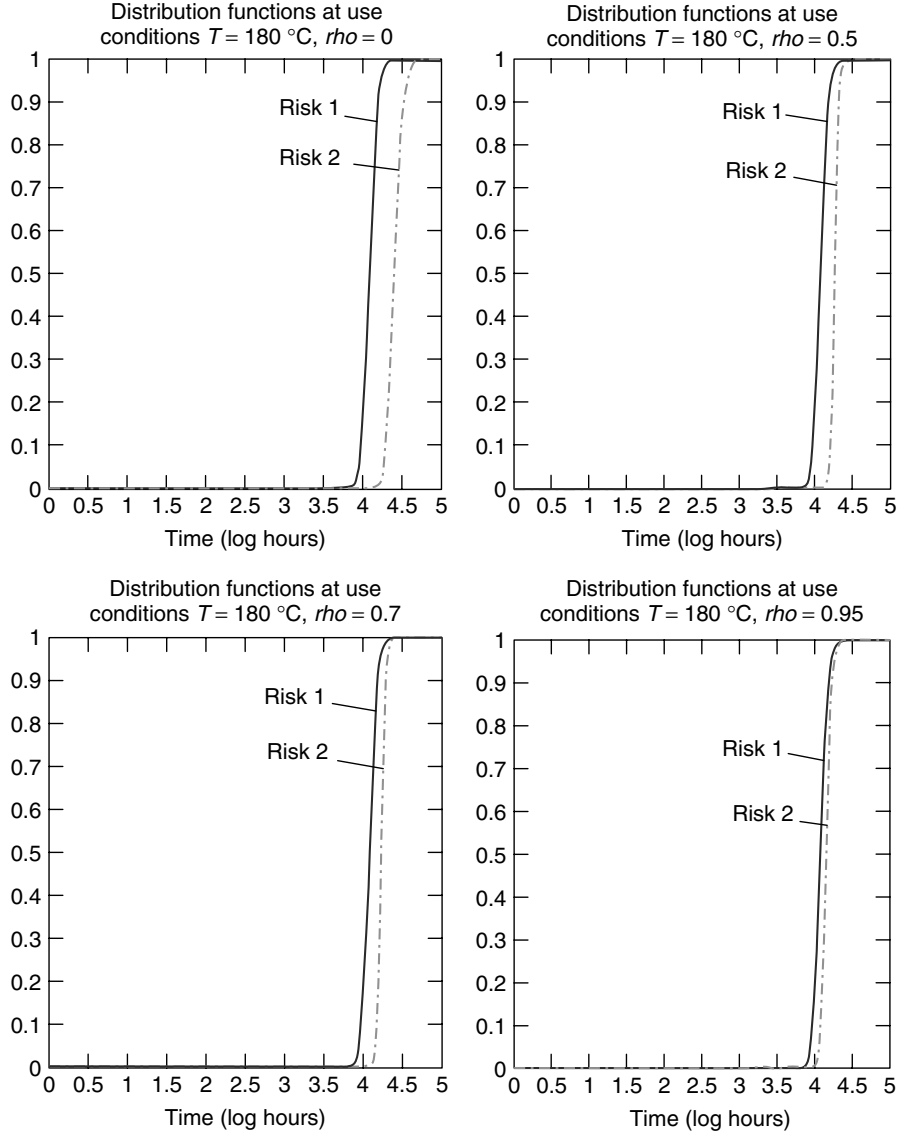


Figure 3 Distribution functions for each failure mode, at use conditions (motorettes data); different degree of dependence are considered

collected from a temperature-accelerated life test of industrial heaters, which have two failure modes – open (breaking of the heating element wire) and short (breakdown of the electrical insulation). A Bayesian Arrhenius-exponential model for random variables U_1 and U_2 is used (see [13]), and the dependence structure is modeled by an Archimedean copula, given the rank correlation coefficient.

Table 2 Estimated failure rates of U_1 and U_2 at each stress level and at use conditions (industrial heaters data)

Temperature ($^\circ\text{F}$)	λ_{U_1}	λ_{U_2}
1150	4.1075E-04	6.0985E-05
1600	6.5287E-04	6.8595E-04
1675	8.0000E-04	7.4048E-04
1750	1.0865E-03	1.3333E-03
1820	3.9393E-03	4.0000E-03

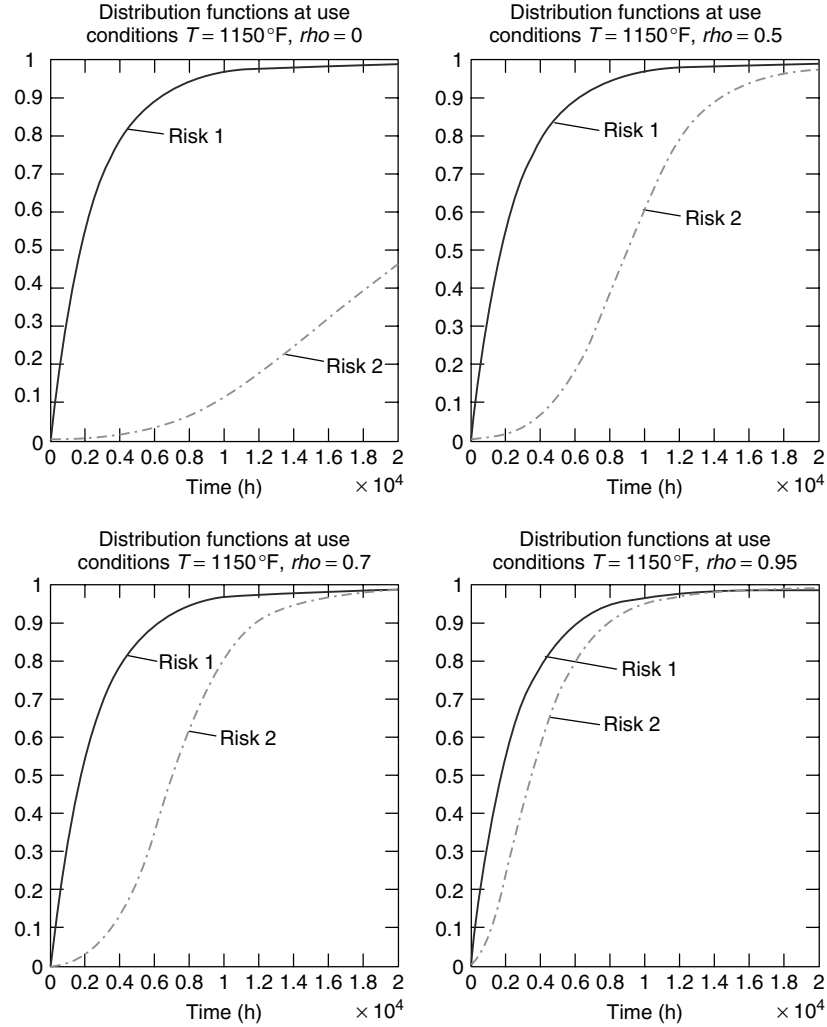


Figure 4 Distribution functions for each failure mode, at use conditions (industrial heaters data); different degree of dependence are considered

Table 2 presents the estimated failure rates of U_1 and U_2 at each stress level and at use condition. Figure 4 shows the distribution functions at use conditions of X_1 and X_2 . The sensitivity of the estimated distributions with respect to the degree of dependence is more obvious in this case.

Conclusions

Several ALT-competing risks models have been presented in this article. Classical independent and

dependent models failed to suit all types of ALT data. Simple plots of experimental data have rejected the applicability of such models, especially when the relative dominance of competing failure modes changes as the applied stress increases. A new dependent ALT-competing risks model based on copula has been proposed. The results show a high sensitivity with respect to the degree of dependence between competing failure modes. Thus, the assumption of “independent failure modes” can lead to wrong results when the risks are actually highly correlated. The disadvantage of this method consists in the large

number of measures of association that need to be estimated by experts (one for each stress level and one for use conditions).

References

- [1] Nelson, W. (1990). *Accelerated Testing*, John Wiley & Sons, New York.
- [2] McCool, J.I. (1978). Competing risk and multiple comparison analysis for bearing fatigue tests, *ASLE Transactions* **21**, 271–284.
- [3] Klein, J.P. & Basu, A.P. (1981). Weibull accelerated life tests when there are competing causes of failure, *Communications in Statistical Methods and Theory* **A10**, 2073–2100.
- [4] Klein, J.P. & Basu, A.P. (1982). Accelerated life testing under competing exponential failure distribution, *Indian Association for Productivity, Quality & Reliability Transactions* **7**, 1–20.
- [5] Hougaard, P. (2000). *Analysis of Multivariate Survival Data*, Springer, New York.
- [6] Bunea, C. & Bedford, T.J. (2002). The effect of model uncertainty on maintenance optimization, *IEEE Transaction on Reliability* **51**(4), 486–493.
- [7] Cooke, R.M. (1996). The design of reliability databases Part I and II, *Reliability Engineering and System Safety* **51**, 137–146, 209–223.
- [8] Bedford, T.J. & Cooke, R.M. (2001). *Probabilistic Risk Analysis: Foundations and Methods*, Cambridge University Press.
- [9] Langseth, H. & Lindqvist, B. (2003). A maintenance model for components exposed to several failure mechanisms and imperfect repair, *Mathematical and Statistical Methods in Reliability, Series on Quality, Reliability and Engineering Statistics*, World Scientific, pp. 415–430.
- [10] Bunea, C., Cooke, R.M. & Lindqvist, B. (2003). Competing risk perspective on reliability databases, *Mathematical and Statistical Methods in Reliability, Series on Quality, Reliability and Engineering Statistics*, World Scientific, pp. 355–370.
- [11] Zheng, M. & Klein, J.P. (1995). Estimates of marginal survival for dependent competing risks based on an assumed copula, *Biometrika* **82**, 127–138.
- [12] Meeuwissen, A.M.H. & Bedford, T.J. (1997). Minimally informative distribution with given rank correlation for use in uncertainty analysis, *Statistical Computation and Simulation* **57**, 143–174.
- [13] Bunea, C. & Mazzuchi, T.A. (2005). Bayesian accelerated life testing under competing failure modes, in *Proceedings of the RAMS Conference*, Alexandria, January 2005.

Related Articles

Accelerated Life Models; Accelerated Life Tests: Classical Methods for Design and Analysis; Accelerated Life Tests: Bayesian Design; Accelerated Life Tests: Nonparametric Approach; Accelerated Life Tests: Bayesian Models.

CORNEL BUNEA AND THOMAS A. MAZZUCHI

Prediction of Expected Fatigue Lives of Fiber Reinforced Plastic Joints

Introduction

Analysis of fatigue is typically based on the concept of a function D that represents the accumulated damage due to stresses and strains occurring at any given time. This function is presumed to increase monotonically, and failure is expected when the accumulated damage reaches some critical level. Usually the damage function is normalized so that it reaches unity at failure — $D = 1$ at the time of failure. This reasonable concept is very compatible with stochastic modeling, with $\{D\}$ becoming a stochastic process when stresses are stochastic. The problem of predicting time of failure then becomes one of studying the $\dot{D}$ rate of growth of $\{D\}$. The practical difficulty in implementing the accumulated **damage model** is that D is not observable by currently available means. D is known at only two points; it is presumed to be at least approximately zero for a specimen that has not yet been subjected to loads and it is unity when failure occurs [1].

Background

Current fatigue analysis methods are based on various approximations and assumptions about the function D . The goal in formulating such approximations is to achieve compatibility with the results of experiments. These experiments are sometimes performed with quite complicated time histories of loading, but more typically they involve simple periodic (possibly harmonic) loads in so-called *constant-amplitude tests*. It is presumed that each period of the motion contains only one peak and one valley. In this situation, the number of cycles until fatigue failure is usually found to depend primarily on the amplitude of the cycle, although this is usually characterized with the alternative nomenclature of *stress range*, which is essentially the double amplitude, being equal to a peak value minus a valley value. There is a secondary dependence on the mean stress level, which can be taken as the average of the peak and valley stresses [1].

Constant-Amplitude Fatigue

The notation S_r denotes the stress range of a cycle and N_f designate the number of cycles until failure in a constant-amplitude test. A typical experimental investigation of constant-amplitude fatigue for specimens of a given configuration and material involves performing a large number of tests including a number of values of S_r , and then plotting the (S_r, N_f) results. This is called an *S/N curve*, and it forms the basis for most information and assumptions about D . One can emphasize the dependence of fatigue life on stress range by writing $N_f(S_r)$ for the fatigue life observed for a given value of the stress range. In principle, the *S/N curve* of $N_f(S_r)$ versus S_r could be any nonincreasing curve, but experimental data commonly show that a large portion of that curve is well approximated by an equation of the form

$$N_f(S_r) = K S_r^{-m} \quad (1)$$

in which K and m are positive constants whose values depend on both the material and the geometry of the specimen. This *S/N curve* becomes a straight line on log–log paper. If S_r is given either by small or very large values, then the form of equation (1) will generally no longer be appropriate, but fatigue analysis usually consists of predicting failure due to moderately large loads, so equation (1) is often quite useful.

Because of the considerable scatter in *S/N data*, it is necessary to use a statistical procedure to account for specimen variability and choose the best $N_f(S_r)$ curve. For the common power-law form of equation (1), the value of K and m are usually found from linear regression, which minimizes the mean-squared error in fitting the data points. It should be noted that this linear regression of data does not involve the use of any probability model.

As noted, the fatigue life also has a secondary dependence on the mean stress. Although this effect is often neglected in fatigue prediction, there are simple empirical formulas that can be used to account for it in an approximate way. One such approach is called the *Goodman correction*, and it postulates that the damage done by a loading $x(t)$ with mean stress x_m and stress range S_r is the same as would be done by an “equivalent” mean-zero loading with an increased stress range of S_e , given by $S_e = S_r / (1 - x_m/x_u)$, in which x_u denotes the ultimate stress capacity of the material. The *Gerber correction* is similar, $S_e =$

2 Prediction of Expected Fatigue Lives of Fiber Reinforced Plastic Joints

$S_r/[1 - (x_m/x_u)^2]$, and has often been determined to be in better agreement with empirical data. Note that use of the Goodman or Gerber correction allows one to include both mean stress and stress range in fatigue prediction, while maintaining the simplicity of a one-dimensional S/N relationship [1].

Another power-law relationship established in metal fatigue is the Manson–Coffin relationship, usually expressed as [2]

$$\varepsilon_{ap} = \varepsilon_f(N)^c \quad (2)$$

in which ε_{ap} is the applied plastic strain amplitude, and ε_f and c are fatigue coefficients. The Manson–Coffin equation is usually employed in the context of local approaches to fatigue damage accumulation; only in such a situation can the plastic strain magnitude be quantified from local material response due to stress concentration and/or residual stresses.

Because it is difficult to define plastic strain for composite materials, the Manson–Coffin relation has not found any application. However, owing to its universality, the power-law S/N curve is often used to characterize fatigue failures, especially when little information on the fatigue behavior of a particular material exists. In composite fatigue, this relationship is usually associated with the names of Hahn and Kim [3] who preferred to express it as

$$sN^d = c \quad (3)$$

In this equation d and c are model (material) parameters, and s is the ratio of the applied stress, S , and the material's ultimate strength, S_{ult} , i.e.,

$$s = \frac{S}{S_{ult}} \quad (4)$$

This ratio is frequently applied as a measure of loading severity in fatigue studies.

Hahn and Kim also consider the semilogarithmic S/N curve model in which the applied stress level is related to $\log(N)$, expressing it as

$$s = k \log(N) + b \quad (5)$$

where k and b are material parameters and s is given by equation (4) [2].

Stochastic Fatigue

As previously noted, the fundamental problem in stochastic analysis of fatigue is to define a stochastic

process $\{D\}$ to model the accumulation of damage caused by a stochastic stress time history $\{X\}$. It seems that most predictions of fatigue life in practical problems are based on simpler models in which information about the accumulated damage is obtained from the constant-amplitude S/N curve.

Hwang and Han [4] suggested that a composite material D can be most generally expressed as

$$D = D(D_0, n, S, f, T, M, \dots) \quad (6)$$

where D_0 is the initial damage existing prior to the application of the considered loading segment, n is the number of fatigue cycles, S is the applied stress level, f is the loading frequency, T is the temperature, and M is the moisture content. The effect of the last three parameters on fatigue damage will not be considered in this section.

The combined influence of the other three parameters, D_0 , n , and S , on the fatigue damage accumulation can be considered in the assumption that when damage is caused by a multilevel stress history, not only does the damage contribution from each loading cycle need to be added up, but the influence of the loading of each cycle on the damage caused by subsequent cycles needs to be identified. This is quite a complex situation that is rarely employed in practice. The first simplification is done by neglecting the effect of D_0 . That is, by assuming that damage caused by the sequence of n cycles with magnitude S is always the same regardless of where this sequence occurs in the loading history. Such damage accumulation rules have been considered in the context of fatigue of fiber reinforced plastic (FRP) laminates and are discussed below.

The simplest damage accumulation law, which is commonly known as *Palmgren–Miner's rule*, assumes that damage contribution from each cycle of the loading history is independent of the other cycles. Therefore, the damage inflicted by n stress cycles with magnitude S is simply

$$D = \frac{n}{N} \quad (7)$$

where N denotes the cycles to failure at S from the constant-amplitude S/N curve. For all stress levels this damage rule yields

$$D = \sum_{i=1}^m \frac{n_i}{N_i} \quad (8)$$

where n_i is the number of cycles having amplitude S_i and N_i is the cycles to failure at S_i .

In many situations, the rate of fatigue damage accumulation in composites may deviate significantly from that assumed in equation (7). For example, Howe and Owen [5] reported that

$$D = B \left(\frac{n}{N} \right) - C \left(\frac{n}{N} \right)^2 \quad B, C > 0 \quad (9)$$

fits much better with constant-amplitude data that they obtained from chopped strand mat (CSM) composites. Such an equation, however, has no particular benefit over the linear damage accumulation rule, equation (7). Making damage a function of the ratio of n/N only implies that it is always the same regardless of the cycle's position in the loading sequence. Therefore, there is always a measure of damage that accumulates linearly under constant-amplitude loadings.

In connection with equation (1), Palmgren–Miner's rule equation (8) is used almost exclusively in metal fatigue. By treating a multilevel stress time history as a random process, equations (1) and (8) sometimes allow quite a simple stochastic approach to be used to estimate cumulative damage. This approach is applicable when the loading can be described as a zero-mean, Gaussian, and narrowband stochastic process. In such a case, the amplitudes of the fatigue cycles possess a Rayleigh probability distribution, which can be expressed as

$$p_S(S) = \frac{S}{S_{\text{rms}}^2} \exp \left(-\frac{S^2}{2S_{\text{rms}}^2} \right) \quad (10)$$

In the above formula, S_{rms} denotes the root-mean-square (RMS) value of the stochastic stress history. By using the Palmgren–Miner hypothesis equation (7) and the power-law S/N curve equation (1), the expected fatigue damage $E[D]$ and the expected cycles to failure $E[N]$ are expressed as

$$\begin{aligned} E[D] &= \frac{1}{E[N]} = \frac{1}{K} E[S^m] \\ &= \frac{1}{K} \int_0^\infty \frac{S^{m+1}}{S_{\text{rms}}^2} \exp \left(-\frac{S^2}{2S_{\text{rms}}^2} \right) dS \end{aligned} \quad (11)$$

The integration of the above equation yields

$$E[D] = \frac{1}{E[N]} = \frac{1}{K} 2^{\frac{m}{2}} S_{\text{rms}}^m \Gamma \left(1 + \frac{m}{2} \right) \quad (12)$$

where $\Gamma(\cdot)$ is the gamma function. Equation (12) is known as the *Rayleigh approximation*.

An expression similar to equation (12) can be derived if the Palmgren–Miner rule is applied in conjunction with the semilogarithmic S/N curve model equation (5). To this end, assume first that the semilogarithmic model is expressed as

$$\log(N) = b - aS \quad (13)$$

in which b and a are the parameters of the model obtained from a linear regression of the experimental data. This equation can be further rearranged as

$$N = 10^{b-aS} = e^{\bar{b}-\bar{a}S} \quad (14)$$

where

$$\bar{b} = b \ln(10) \quad \bar{a} = a \ln(10) \quad (15)$$

The expected damage can then be expressed as

$$\begin{aligned} E[D] &= \frac{1}{E[N]} = e^{-\bar{b}} \int_0^\infty \frac{S}{S_{\text{rms}}^2} \\ &\quad \times \exp \left(-\frac{S^2}{2S_{\text{rms}}^2} + \bar{a}S \right) dS \end{aligned} \quad (16)$$

Performing the integration in equation (16) gives

$$\begin{aligned} E[D] &= \frac{1}{E[N]} = e^{-\bar{b}} \left[1 + \bar{a}S_{\text{rms}} \exp \left(\frac{\bar{a}^2 S_{\text{rms}}^2}{2} \right) \right. \\ &\quad \times \left. \Phi(\bar{a}S_{\text{rms}}) \sqrt{2\pi} \right] \end{aligned} \quad (17)$$

in which $\Phi(\cdot)$ denotes the Gaussian cumulative function [6].

Experimental Approach

The composite test program discussed in this paper was set up to provide a more realistic insight into the static and fatigue behavior and damage accumulation of FRP joints. To study the fatigue damage accumulation in FRP specimens, several different specimen configurations were designed and fabricated for characterization under constant- and variable-amplitude fatigue loadings.

The base laminate consisted of an alternating 0/90° stacking sequence of 30 plies made of 0.025 in.

4 Prediction of Expected Fatigue Lives of Fiber Reinforced Plastic Joints

thick Woven Roving (WR) glass fabric and 510-A vinyl ester resin. This layup was fabricated by Hardcore DuPont Composites, L.L.C., into several panels measuring 2 ft by 3 ft using the Seeman composite resin infusion molding process (SCRIMP) process. The postcure cycle for the panels was 2 h at 160 °F. The panels were then cut into specimens.

The behavior of composite joints under static, constant-amplitude, and variable-amplitude fatigue was studied. Each configuration was realized in a single geometry butt-strap joint. Specimens were made of the same laminate as previously discussed with the addition of 15 ply laminates of the same layup; joint adherents were oriented along the 0/90° material direction. Figure 1 shows the typical geometries of joint specimens [6].

Static Tests

The static tests consisted of ultimate tension and compression strength tests, which were designed to characterize the static behavior of the specimens and to aid in selecting the appropriate stress levels for the fatigue tests. In all configurations, some strain gages were mounted on the surface of the joints in order

to monitor the evolution of the state at the points of interest. All bonded specimens failed in an adhesive mode. The bolted joint exhibited a combination of cleavage-tension and shear-out failure and the bolted-bonded joint, after demonstrating first a bond failure, had finally failed in a shear-out mode. In the compressive regime, the bonded joint had delaminated in both laps and one of the adherents owing to the combined effect of contraction and buckling. Similarly, the bolted joint had shortened and buckled antisymmetrically (mode 2) and all joint components had delaminated as a consequence. The bonded-bolted joint showed an early bond failure and a final failure due to a symmetric buckling and a consequent delamination in all components. The bolted-bonded joint is the most advantageous configuration, as far as the tensile strength is concerned, since it has a higher initial stiffness, a higher failure load and a very good ductility compared to the others. On the other hand, in the compressive regime, although it has a higher initial stiffness, its strength and ductility are lower than those of the bolted configuration, which may be due to the difference in the buckling mode, since a mode 2 buckling requires a higher load to occur [6].

Fatigue Tests

The fatigue tests typically consisted of two to three specimens tested at each stress level, or RMS level for constant or variable amplitude, respectively. All tests were performed at constant frequency. The variable-amplitude loading represents a segment of a sinusoidal time history composed of 10 000 end points, which with the assumption of stationary and ergodicity can be shown to be a narrowband Gaussian process. This segment is repeated until failure of the specimen occurs [6].

Tables 1 and 2 show the constant-amplitude and variable-amplitude test results for different types of composite joints.

The failure of the bonded joints is an adhesive failure of the bond, while that of the bolted joints is a bearing failure followed by a partial shear-out failure. The failure of the bonded-bolted joints is a two-stage one. First, and very early, comes the adhesive failure of the bonds and then, typically after a long time, the bearing and partial shear-out failures of the bolts occurs.

The experimental fatigue lives from the constant-amplitude tests of all joint configurations are presented in Figure 2 on a log(S)–log(N) scale and in

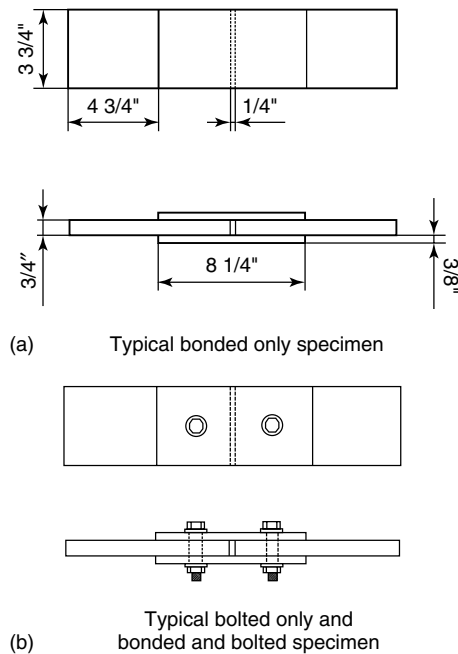


Figure 1 Geometry of different types of composite joints

Table 1 Constant-amplitude fatigue tests of FRP joints

Joint type	ID name	Maximum stress level (ksi)	Maximum load level (lbs)	Cycles to failure
Bonded only	JA-9	2	5044	2 411 800
	JA-10	2	5042	3 369 600
	JA-20	2	5035	2 748 520
	JA-7	3	7735	51 900
	JA-8	3	7591	255 400
	JA-19	3	7735	1 000 300
	JA-17	4	10 134	98 600
	JA-18	4	10 170	86 200
	JA-5	5	12 806	3170
	JA-6	5	12 846	4220
	JA-3	7.5	18 737	980
	JA-4	7.5	18 764	1060
Bolted only	JB-7	3	7607	589 400
	JB-8	3	7019	20 000 000 ^(a)
	JB-13	3	7307	23 200
	JB-14	3	7319	899 100
	JB-11	3.5	8617	12 600
	JB-12	3.5	8409	8720
	JB-9	4	9623	1810
	JB-10	4	9464	2530
	JB-5	5	11 795	890
	JB-6	5	11 695	800
	JB-3	7.5	17 774	305
	JB-4	7.5	17 793	320
Bonded and bolted	JAB-9	4	9511	6 964 000
	JAB-10	4	9667	6 578 900
	JAB-11	4.5	10 958	517 800
	JAB-12	4.5	10 765	2 733 500
	JAB-5	5	11 878	475 900
	JAB-6	5	11 775	333 500
	JAB-7	6	14 463	13 000
	JAB-8	6	14 191	16 200
	JAB-3	7.5	17 850	6560
	JAB-4	7.5	17 692	5420
	JAB-13	9	21 331	1480
	JAB-14	9	21 331	2780

^(a) Stopped

Figure 3 on a S-log(N) scale. Both logarithmic and semilogarithmic S/N curve models were fit to these data by linear regression. The curves indicate a definite deviation from the logarithmic S/N curve model for all joint configurations. Neither S/N curve model appears to represent the data well.

Variable-Amplitude Fatigue Tests Analysis

Experimental variable-amplitude fatigue lives are obtained by applying the aforementioned time

history to the specimens at different RMSs; and the analytic random fatigue lives are calculated by following the fatigue damage accumulation models (equations (12) and (17)). Comparison between the variable-amplitude fatigue results and the Rayleigh approximation, based on both logarithmic and semilogarithmic S/N curve model is provided in Table 3. It can be seen that for the lower RMS levels considered, the classical Rayleigh approximation underestimates the actual fatigue life from 37 to 67%. The same trend is observed for Gaussian approximation method and the

6 Prediction of Expected Fatigue Lives of Fiber Reinforced Plastic Joints

Table 2 Variable-amplitude fatigue tests of FRP joints

Joint type	ID number	RMS (ksi)	Maximum–Minimum stress (ksi)	Maximum load level (lbs)	Cycles to failure
Bonded only	JA-14	1.25	±5.426	13 812	3 654 000
	JA-15	1.25	±5.426	13 912	1 441 900
	JA-16	1.25	±5.426	13 839	4 419 200
	JA-11	1.75	±7.597	19 949	62 500
	JA-12	1.75	±7.597	19 103	22 500
	JA-13	1.75	±7.597	19 306	26 700
Bolted only	JB-18	1.25	±5.426	13 116	253 900
	JB-19	1.25	±5.426	13 060	492 500
	JB-20	1.25	±5.426	13 161	356 700
	JB-15	1.75	±7.597	18 671	7520
	JB-16	1.75	±7.597	18 407	9390
	JB-17	1.75	±7.597	18 357	8220
Bolted and bonded	JAB-18	1.75	±7.597	19 279	5 343 100
	JAB-19	1.75	±7.597	18 070	11 934 700
	JAB-20	1.75	±7.597	17 986	9 977 000
	JAB-15	2.25	±9.765	23 169	203 100
	JAB-16	2.25	±9.765	22 900	305 700
	JAB-17	2.25	±9.765	24 664	67 400

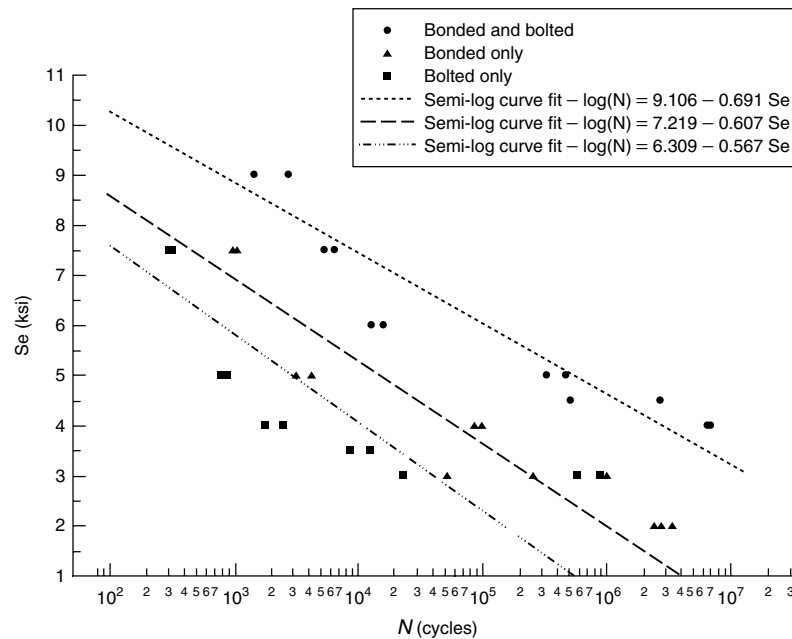


Figure 2 S/N curves of FRP joints (log–log curve fit)

expected fatigue lives are underestimated from 56 to 66%.

The fatigue lives at the higher RMS levels are significantly overestimated by using both Rayleigh and

Gaussian methods. For both bonded only and bolted only specimen types, the predictions at the higher RMS levels (1.75 ksi) are somewhat improved by the use of Rayleigh approximation method, but this

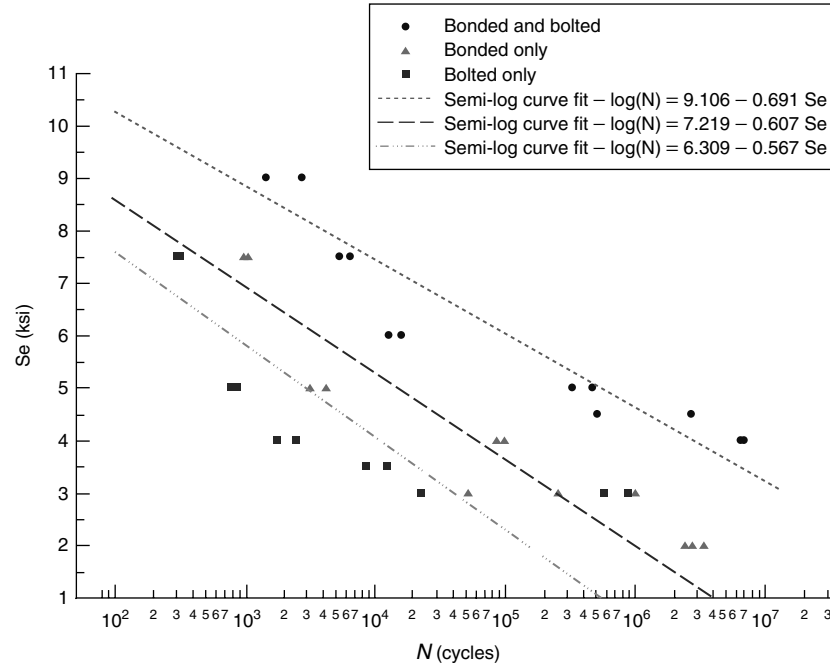


Figure 3 S/N curves of FRP joints (semi-log curve fit)

Table 3 Comparison of Rayleigh and Gaussian fatigue lives with test results

Joint type	RMS level (ksi)	Experimental (geometrical mean)	Rayleigh approximation	Gaussian method
Bonded only	1.25	2 855 600	925 300	979 200
	1.75	33 500	113 600	148 300
Bolted only	1.25	354 600	225 500	129 800
	1.75	8340	22 100	26 100
Bonded and bolted	1.75	8 600 700	3 510 200	3 789 200
	2.25	161 100	275 400	234 500

improvement is quite marginal. At the same time, Rayleigh approximation leads to further deviation from the experimental mean for the two lower stress RMS levels of bonded only and bonded and bolted specimens (1.25 and 1.75, respectively). These large deviations mainly come from the basic assumptions built in the Rayleigh approximation.

As pointed out in the previous section, the classical Rayleigh and Gaussian approximation formulas (12) and (17) is based on both equation (1) and on the Palmgren–Miner damage accumulation rule equation (8). Therefore, misrepresentation of the actual stress-life relationship caused by a poor S/N curve

function or insufficient S/N data, or disregarding the sequence effects by using the Miner rule, are main reasons for the shortcomings of the Rayleigh approximation in estimating the random fatigue lives.

Conclusion

Fatigue damage in FRP laminates occurs in the form of diffuse zones containing microcracks in the matrix, fiber–matrix debonding sites, ply delamination, and broken fibers. This damage leads to **degradation** of the elastic moduli of the laminate, and is therefore

best modeled by stiffness degradation approaches. Various S/N curve models for constant-amplitude fatigue damage characterization can be described in terms of the fatigue modulus concept of Hwang and Han. In contrast, fatigue damage accumulation techniques constitute a much less investigated area; most data available deal with two-stage loading cases only [6]. The constant- and variable-amplitude fatigue studies on the laminate joints tested show a definite trend toward deviation from the power-law S/N curve model. Variable-amplitude fatigue tests with narrowband stress histories demonstrate significant deviation of the observed fatigue lives from those predicted with the Rayleigh approximation and Gaussian cumulative function. Such a deviation can be attributed to the inadequacy of the Palmgren–Miner’s damage accumulation rule for composite materials and to poor representation or insufficiency of the constant-amplitude data as well as load frequency and sequence effects, which have not yet been quantified and incorporated into a fatigue damage accumulation model.

References

- [1] Lutes, L.D. & Sarkani, S. (2004). *Random Vibration*, Elsevier, Butterworth-Heinemann, pp. 520–523.
- [2] Momenkhani, K. & Sarkani, S. (2006). A new method for predicting the fatigue life of fiber-reinforced plastic laminates, *Journal of Composite Materials* **40**(21), 1971–1982.
- [3] Hahn, H.T. & Kim, R.Y. (1976). Fatigue behavior of composite laminates, *Journal of Composite Materials* **10**, 156–180.
- [4] Hwang, W. & Han, K.S. (1986). Cumulative damage models and multi-stress fatigue life prediction, *Journal of Composite Materials* **20**, 125–153.
- [5] Howe, R.J. & Owen, M.J. (1972). Cumulative damage in chopped strand mat polyester resin laminates, *Proceedings of the 8th International Reinforced Plastics Conference*, BPF, London, pp. 137–148.
- [6] Michaelov, G., Kharrazi, M., Momenkhani, K., Sarkani, S. & Kihl, D.P. (1997). *Fatigue Behavior and Damage Accumulation of Fiber-Reinforced Plastic Laminates and Joints*, NSWCCD-TR-65-97/36, Naval Surface Warfare Center, Carderock Division.

KOUROSH MOMENKHANI AND SHAHRAM
SARKANI

Design for Reliability

Introduction

To ensure product reliability, an organization must follow certain practices during the product development process. These practices impact reliability through the selection of materials, product design, manufacturing, assembly, shipping and handling, operation, and maintenance and repair. These practices are listed below and briefly described in this chapter.

- Define realistic product requirements determined by factors including the life-cycle application profiles, required operating and storage life, performance expectations, size, weight, and cost. The product requirements should be based on both the customer's needs and the manufacturer's capability to meet those needs.
- Define the product life-cycle conditions by specifying all relevant manufacturing, assembly, storage, handling, shipping, operating, and maintenance conditions.
- Ensure that the supply-chain participants have the capability to produce the parts (materials) or services to meet the final reliability objectives.
- Select the materials and parts that have sufficient quality, and are capable of delivering the expected performance and reliability in the application. The materials and parts must be characterized. Variabilities in material properties and manufacturing processes will impact the product's reliability.
- Identify the potential failure modes, failure sites, and failure mechanisms by which the product can be expected to fail.
- Design to the usage and process capability of the product (i.e., the quality level that can be controlled in manufacturing and assembly) considering the potential failure modes, failure sites, and failure mechanisms. The designed product must satisfy the manufacturability, quality, reliability, and budget requirements and constraints, and be available to the customers in a timely manner.
- Qualifications should be conducted to verify the reliability of the product in the expected life-cycle conditions. Qualification encompasses all activities that ensure that the nominal design and

manufacturing specifications will meet or exceed the reliability goals. The goal of this step is to provide a physics-of-failure basis for design decisions, with an assessment of all possible failure mechanisms for the anticipated product.

- All manufacturing and assembly processes must be capable of producing the product within the statistical process window required by the design. Therefore, characteristics of the process must be identified, measured, and monitored. Regardless of how well a product is designed, it cannot be reliable unless the variability in the manufacturing processes is controlled (not necessarily minimized). Each process may involve screens and tests to assess **statistical process control**.
- Manage the life-cycle usage of the product using closed-loop management procedures. This includes realistic inspection and maintenance procedures, and root-cause analysis.

Product Requirements and Constraints

A product's requirements and constraints are defined in terms of customer demands and the company's core competencies, culture, and goals. There are various reasons to justify the creation, modification, upgrade, or even discontinuation of a product. For example, a company may want to address a perceived market need or to open new markets. In some cases, a company may need to develop new products to remain competitive in a key market or to maintain market share and customer confidence. In other cases, a company may want to fill a need of specific strategic customers, or to demonstrate experience with a new technology or methodology, or to improve maintainability of an existing product. In addition, product updates are often developed to reduce the life-cycle costs of an existing product.

To make reliable products, there should be a joint effort between the supplier and the customer. According to IEEE 1332 [1], there are three reliability objectives to achieve in the stage of defining product requirements and constraints. First, the supplier, working with the customer, shall determine and understand the customer's requirements and product needs, so that a comprehensive design specification can be generated. Second, the supplier shall structure and follow a series of engineering activities so that the resulting product satisfies

2 Design for Reliability

the customer's requirements and product needs with regard to product reliability. Third, the supplier shall perform activities that assure the customer that the reliability requirements and product needs have been satisfied.

If the product is for direct sale to end users, marketing usually takes the lead in defining the product's requirements and constraints through interaction with the customer's marketplace, examination of product sales figures, and analysis of the competition. Alternatively, if the product is a subsystem that fits within a larger product, the requirements and constraints may be determined by the customer's requirements for the larger product into which the subsystem fits. The product definition process involves multiple influences and considerations. Figure 1 shows a diagram of constraints in the product definition process.

Two prevalent risks in defining the requirements and constraints are the inclusion of irrelevant requirements and the omission of relevant requirements. The inclusion of irrelevant requirements can involve unnecessary design and testing time as well as money. Irrelevant or erroneous requirements generally result from two sources: requirements created by persons who do not understand the constraints and opportunities implicit in the product definition, and inclusion of requirements for historical reasons. The omission

of critical requirements can significantly reduce the effectiveness of the product.

The initial requirements are formulated into a requirements document, where they are prioritized. The requirements document should be approved by the engineers, the management, the customers, and other involved parties. Usually, the specific people involved in the approval will vary with the organization and the product. For example, for human safety-critical products, legal representatives should attend the approval to identify legal considerations with respect to the implementation of the parts selection and management team's directives. The results of capturing product requirements and constraints allow the design team to choose parts and develop products that conform to company objectives.

Once a set of requirements has been completed, the product engineering function creates a response to the requirements in the form of a specification. The specification states the requirements that must be met; the schedule for meeting the requirements; the identification of those who will perform the work; and the identification of potential risks. Differences in the requirements document and the preliminary specification become the topic of trade-off analyses.

Once product requirements are defined and the design process begins, there should be an assessment

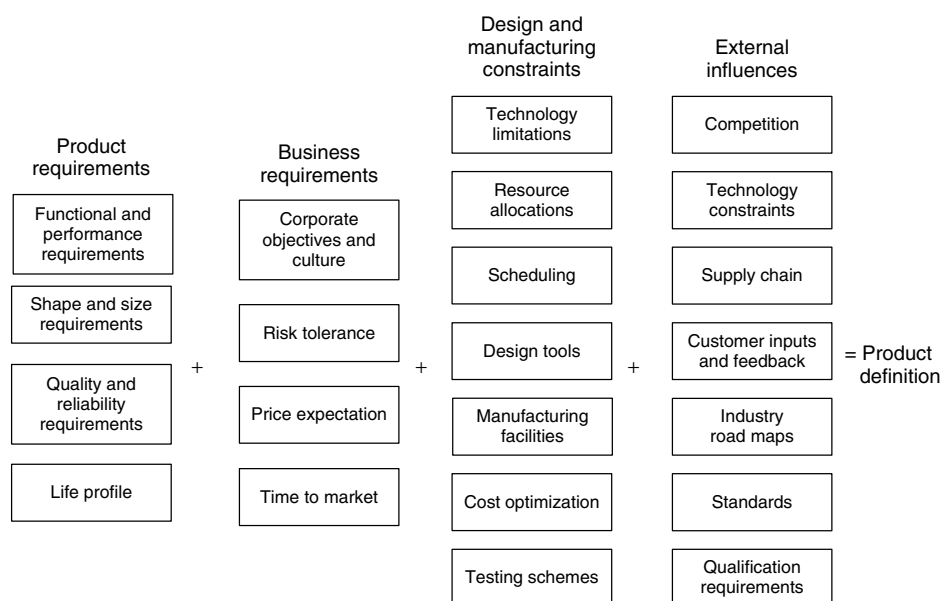


Figure 1 Constraints in the product definition process

of the product requirements against the actual product design. As the product's design becomes increasingly detailed, it becomes increasingly important to track the product's characteristics (size, weight, performance, functionality, reliability, and cost) in relation to the original product requirements. The rationale for making changes should be documented. The completeness with which the tracking of requirements is performed can significantly reduce future redesign costs of the product. Planned redesigns or design refreshes through technology monitoring and use of road maps ensure that a company is able to market new products or redesigned versions of old products in a timely, effective manner to retain their customer base and ensure continued profits.

Product Life-Cycle Conditions

The life-cycle conditions of the product play a significant role in determining the product requirements and guiding reliability assessments. It influences decisions regarding product design and development, materials and parts selection, qualification, product safety, warranty, and support.

The phases in a product's life cycle include manufacturing and assembly, testing, rework, storage, transportation and handling, operation^a (modes of operation, on-off cycles, etc.), repair and maintenance, and disposal. During each phase of its life cycle, the product experiences various external and internal environmental and usage loads. The life-cycle loads can be thermal (steady-state temperature, temperature ranges, temperature cycles, temperature gradients), mechanical (pressure levels, pressure gradients, vibrations, shock loads, acoustic levels), chemical (aggressive or inert environments, ozone, pollution and humidity levels, contamination, fuel spills), physical (radiation, electromagnetic interference, altitude), and/or electrical loading conditions (power, power surge, current, voltage, voltage spikes). These loads, either individually or in various combinations, may influence the reliability of the product. The extent and rate of product **degradation** depend upon the nature, magnitude, and duration of exposure to such loads.

Defining and characterizing the life-cycle loads is often the most uncertain element of the overall design for reliability process [2]. The challenge occurs because products can experience completely different

application conditions, depending on location and customer. The application conditions include the application duration, the number of applications, the product utilization or nonutilization profile, and maintenance or servicing conditions. For example, typically all desktop computers are designed for office environments. However, the operational profile of each unit may be completely different depending on user behavior. Some users may shut down the computer every time after it is used; others may shut down only once at the end of the day; still others may keep their computers powered on all the time. Thus, the temperature profile experienced by each product, and hence its degradation due to thermal loads, will be different. There are four methods used to estimate the life-cycle loads on a product: market studies and standards, field trial and service records, similarity analysis, and *in situ* monitoring. Each is discussed below.

Market surveys and standards^b provide a coarse estimate of the actual environmental loads expected in the field. The environmental profiles available from these sources are typically classified according to industry type, such as military, consumer, telecommunications, automotive, and commercial avionics. Often the data include the worst-case and typical-use conditions.

Field trial records provide estimates of the environmental profiles experienced by the product. The data depend on the durations and conditions of the trials, and can be extrapolated to get an estimate of actual environmental conditions. Service records provide information on the maintenance, replacement, or servicing performed. This data can give an idea of the life-cycle environmental and usage conditions that lead to servicing or failure.

Similarity analysis is a technique for estimating environmental loads when sufficient field histories for similar products are available. Before using data on existing products for proposed designs, the characteristic differences in design and application for the two products need to be reviewed. Changes and discrepancies in the conditions of the two products should be critically analyzed to ensure the applicability of available loading data for the new product. For example, electronics inside a washing machine in a commercial laundry are expected to experience a wider distribution of loads and use conditions (due to a large number of users), and higher usage rates compared with those in a domestic washing machine.

These differences should be considered during similarity analysis.

Environmental and usage conditions experienced by the product over its life cycle can be monitored *in situ*. This data is often collected using sensors, either mounted externally or integrated with the product and supported by telemetry systems. Load distributions should be developed from data obtained by monitoring products used by different customers, ideally from various geographical locations. The data should be collected over a sufficient period to provide an accurate estimate of the loads and their variation over time. *In situ* monitoring has the potential to provide the most accurate account of load history for use in health (condition) assessment and design of future products.

Reliability Capability

The selection of a supplier for electronic parts or services is often based on factors that do not explicitly address reliability, such as technical capabilities, production capacity, geographic location, support facilities, and financial and contractual factors. A selection process that takes into account the ability of suppliers to design for reliability and meet reliability objectives during manufacturing, testing, and support can improve reliability of the final product throughout its life cycle. This can provide valuable competitive advantage to an equipment manufacturer or system integrator.

Reliability capability is a measure of the practices within an organization that contribute to the reliability of the final product, and the effectiveness of these practices in meeting the reliability requirements of customers. The reliability capability of an organization is a combination of the reliability capabilities of constituent reliability activities. Reliability capability assessment is the act of quantifying the effectiveness of these activities, using a metric called the *reliability capability maturity*. From a reliability perspective, maturity indicates whether the key reliability practices within an organization are well understood, supported by documentation and training, applied to all products throughout the organization, and continually monitored and improved.

In the absence of systematic evaluation of these organizational traits, product testing and monitoring of field performance are often the primary means

for assessing reliability over time. Improvement in product reliability based on these metrics is generally limited to a single product family and is difficult to institutionalize over the long term and across product lines. Periodic assessment of organizational reliability capability offers a method for identifying practices in need of improvement and implementing required improvements on a continual basis across multiple product lines or departments.

Parts and Materials Selection

A parts selection and management methodology helps a company to make risk-informed decisions when purchasing and using parts and materials. The methodology shown in Figure 2 describes this approach to parts (materials) selection.

In the manufacturer assessment step, the part manufacturer's ability to produce parts with consistent quality is evaluated. This was already discussed in the section titled "Reliability Capability". Similarly the distributor assessment evaluates the distributor's ability to provide parts without affecting the initial quality and reliability, and to provide specific services, such as part problem and change notifications.

The goal-of-performance assessment is to evaluate the part's ability to meet the performance requirements (e.g., functional, mechanical, and electrical) of the product. There is often a minimum and a maximum limit beyond which the part will not function properly and, in general, these bounds are defined by the recommended operating conditions for the part.

Reliability assessment provides information about the ability of a part to meet the required performance specifications in its life-cycle application environment for a specified period of time. If the parametric and functional requirements of the product cannot be met within the required local environment, then the local environment may have to be modified, or a different part may have to be used. If part reliability is not ensured through the reliability assessment process, an alternate part or product redesign should be considered.

Reliability assessment is conducted through the use of reliability test data (conducted by the part manufacturer), virtual reliability assessment results, or accelerated test results. If the magnitude and duration of the life-cycle conditions are less severe than those of the reliability tests, and if the test sample

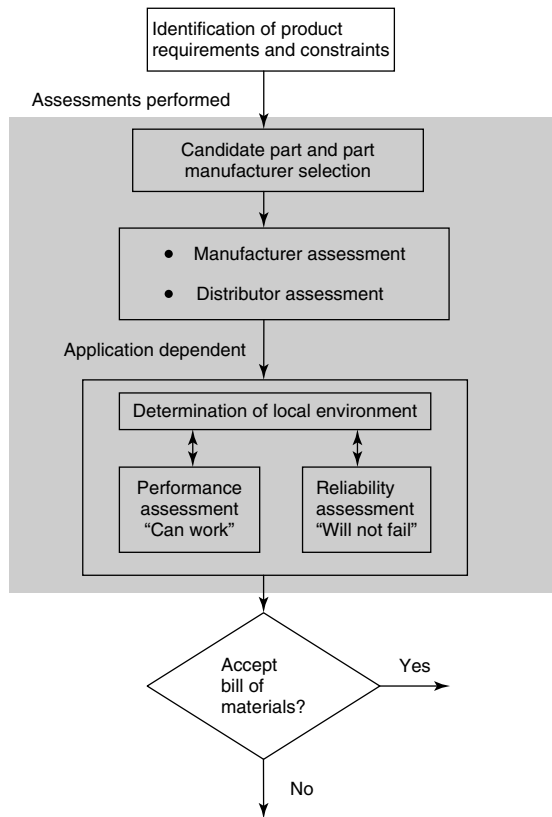


Figure 2 Parts selection and management methodology

size and results are acceptable, the part reliability is acceptable. Otherwise, a virtual reliability assessment should be considered. Virtual reliability is a simulation-based methodology used to identify the dominant failure mechanisms associated with the part under life-cycle conditions to determine the acceleration factor for a given set of accelerated test parameters, and to determine the times to failure corresponding to the identified failure mechanisms. If virtual reliability proves insufficient to validate part reliability, accelerated testing should be performed on representative part lots at conditions that represent the application condition to determine the part reliability. This decision process is illustrated in the diagram in Figure 3.

Failure Modes, Mechanisms, and Effects Analysis

A failure mode is the manner in which a failure can occur – that is, the ways in which the products fails to perform its intended design function, or performs the function but fails to meet its objectives [3, 4]. A failure mode can also be defined as any physically observable change caused by a failure mechanism. In an electrical system, failure mode can be the change of current, voltage, or electrical noise beyond some limiting values. Sometimes the failure modes are intentionally accentuated so that the user of a system is aware of the existence of a problem. The

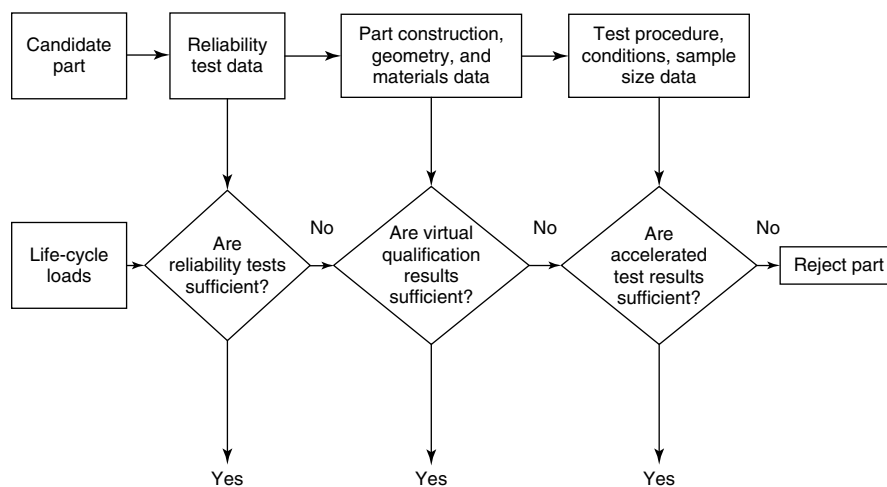


Figure 3 Decision process for part reliability assessment

addition of a compound with a distinct smell to natural gas to indicate the existence of a leak is such an example. Another example is the grinding noise when the brake pads wear out in a car. Failure modes are closely related to the functional and performance requirements of the product. Failure mechanisms are the processes by which a specific combination of physical, electrical, chemical, and mechanical stresses induces failures. Each failure mechanism is a process but not a physical condition or failure site. For example, solder failure is not a failure mechanism, but that description does identify a failure site.

The **failure mode and effects analysis** (FMEA) methodology, which was developed in the 1950s and has since continuously evolved, is a procedure to recognize and evaluate the potential failure of a product and its effects, and to identify actions that could eliminate or reduce the likelihood of the potential failure to occur [5]. This methodology can also be referred to as *failure modes effects and criticality analysis* (FMECA), where the criticality of each failure mode is also estimated and reported.

A limitation of the FMEA procedure is that it does not identify the product failure mechanisms and models in the analysis and reporting process. Investigation of the possible failure modes and mechanisms of the product aids in developing reliable designs. A design team must be aware of the possible failure mechanisms to design hardware capable of withstanding loads without failing. Failure mechanisms and their related physical models are necessary for effective tests and screens to audit nominal design and manufacturing specifications, as well as the level of defects introduced by excessive variability in manufacturing and material parameters.

The purpose of failure modes, mechanisms, and effects analysis (FMMEA) is to identify potential failure mechanisms and models for all the potential failure modes, and to prioritize failure mechanisms. FMMEA is based on understanding the relationships between product requirements and the physical characteristics of the product (and their variation in the production process), the interactions of product materials with loads (stresses at application conditions), and their influence on the product's susceptibility to failure with respect to the use conditions. This involves finding the failure mechanisms and reliability models to quantitatively evaluate the susceptibility to failure. FMMEA uses life-cycle environmental and operating conditions and the duration of the intended

application with knowledge of the active stresses and potential failure mechanisms. The FMMEA process is summarized in Figure 4.

An essential input to an FMMEA is a well-developed life-cycle environmental profile (LCEP). Information on life-cycle conditions can be used to include or eliminate failure mechanisms that are active only under certain application conditions.

The product under investigation needs to be defined. The product is then divided into convenient and logical elements. This breakdown can be either functional (i.e., according to what the system elements "do") or architectural (i.e., according to where the system elements "are"), or both (i.e., functional within a location or *vice versa*). For example, in an automobile, a functional breakdown would include the cooling system, braking system, propulsion system, and so on. An architectural breakdown would involve the engine compartment, passenger compartment, and dashboard or control panel.

After this logical breakdown is performed, all possible failure modes for each identified element should be listed. Potential failure modes may be identified using numerical stress analysis, accelerated tests to failure (e.g., Highly Accelerated Life Tests (HALT)), past experience, and engineering judgment. The failure mode identification should not imply a cause or mechanism.

A failure cause is defined as the specific process, design, and/or environmental condition that initiates a failure and whose removal will eliminate the failure. Life-cycle environmental and operational conditions are used to determine the potential causes for the potential failure modes. Knowledge of potential failure causes can help identify the failure mechanisms that create the failure modes for a given element. This list of causes should be exhaustive and input should be taken from several functional groups.

Identification of the failure mechanisms should be performed with care so that they are not confused with sites, modes, or causes. While performing FMMEA for specific hardware, specialists for those products or subsystems should be included in the physical structure.

The failure models should show the times to failure as a function of product material, geometry, environmental, and application conditions. Failure models for well-known failure mechanisms can be obtained from the literature. This may be a difficult step to perform if the failure mechanisms are not

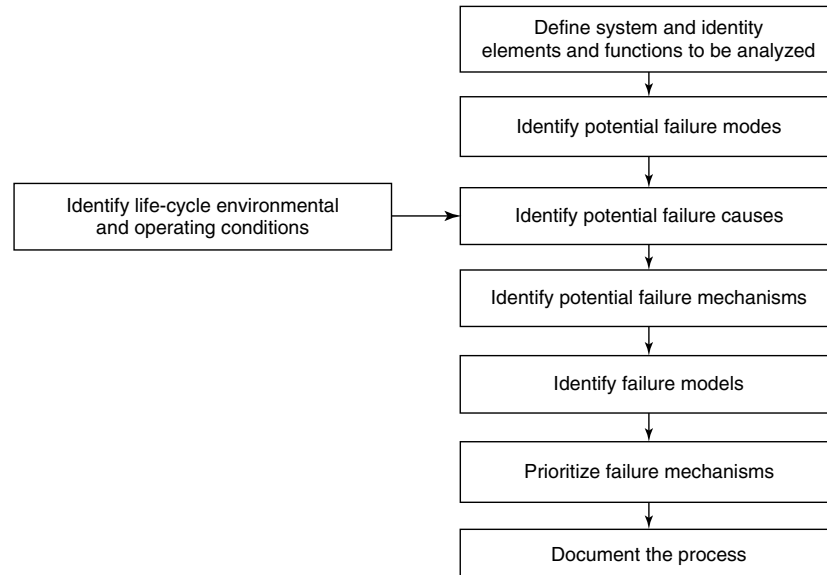


Figure 4 FMMEA methodology

clearly identified as a physical process. It may also be necessary to develop failure models or determine the values of constants associated with the models from prior failure data or from other test data. Organizations should assign resources to develop, maintain, and update failure model libraries for the failure mechanisms of concern.

Using the failure models in association with environmental conditions, product design and geometry, a process is followed to determine the highest-priority failure mechanisms. The goal then becomes to identify and eliminate or control the causes responsible for those mechanisms.

Design Techniques

Once the parts, materials, processes, and stress conditions are identified, the objective is to design a product using parts and materials that have been sufficiently characterized in terms of performance over time when subjected to the manufacturing and application profile conditions. Only through a methodical design approach, using physics-of-failure analysis, testing, and root-cause analysis can a reliable and cost-effective product be designed.

Design guidelines based on physics-of-failure models can also be used to develop tests, screens, and

derating^c factors. Tests based on physics-of-failure models can be designed to measure specific quantities, to detect the presence of unexpected flaws, and to detect manufacturing or maintenance problems. Screens can be designed to precipitate failures in the weak population while not cutting into the design life of the normal population. Derating or safety factors can be determined to lower the stresses for the dominant failure mechanisms.

The following sections describe approaches that can be employed in the design phase of electronic equipment to improve reliability.

Protective Architectures

The objective of protective architectures is to enable some form of action, after an initial failure or malfunction, to prevent additional or secondary failures. Protective techniques include the use of fuses and circuit breakers, self-sensing structures, and adjustment structures that correct for parametric shifts.

In designs where safety is an issue, it is generally desirable to incorporate some means for preventing a part, structure, or interconnection from failing or from causing further damage when it fails. Fuses and circuit breakers are examples of elements used in electronic products to sense excessive current

drain and disconnect power from the concerned part. Fuses within circuits can also safeguard parts against voltage transients or excessive power dissipation, and protect power supplies from shorted parts. Similarly, thermostats can be used to sense critical temperature limiting conditions, and to automatically disconnect the product or part of the system from the power circuit until the temperature returns to normal. Self-checking circuitry can also be incorporated to sense abnormal conditions and restore normal conditions, or to activate circuitry that will compensate for the malfunction.

In some instances, it may be desirable to permit partial operation of the product after a part failure, possibly with degraded performance, rather than completely disconnecting power supply to the product. For example, in shutting down a failed circuit whose function is to provide precise trimming adjustment within a deadband of another control product, acceptable performance may be achieved, under emergency conditions, with the deadband control product alone [6].

Sometimes, the physical removal of a part from a product can harm or cause failure in another part by removing load, drive, bias, or control. In such cases, an interlock mechanism may be implemented within the part that is to be removed to shut down or protect other parts.

By reducing the number of failures, protective techniques also affect product availability and effectiveness. In addition, protective architectures must be designed by considering the impacts of maintenance, repair, and part replacement. For example, if a fuse protecting a circuit is replaced, the following questions need to be answered: what is the impact when the product is reenergized? Which protective architectures are appropriate for post-repair operations? What maintenance guidance must be documented and followed when fail-safe protective architectures have or have not been included?

Stress Margins

A properly designed product should be capable of operating satisfactorily with parts that drift or change with time and with changes in operating conditions, such as temperature, humidity, pressure, and altitude, as long as the parameters remain within their rated tolerances.

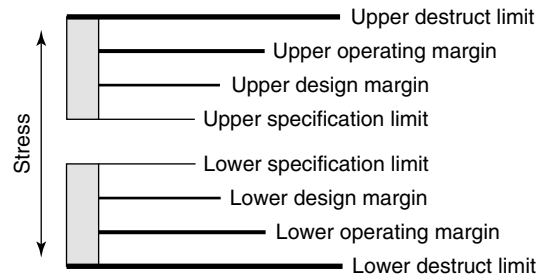


Figure 5 Stress limits and margins

Figure 5 provides a schematic representation of the hierarchy of product stress limits and margins. The specified tolerance limit is set by the manufacturer to limit the conditions of customer use. The design margin corresponds to the stress value that the product is designed to survive without failure. The operating margin is the expected stress value that may lead to a recoverable failure. The destruct margin is the expected stress value that may lead to permanent (overstress) failure.

To prevent out-of-tolerance failures, a design team must consider the combined effects on stress levels of manufacturing tolerances, including those for parts used in future repair or maintenance operations; expected environmental conditions; and drifts due to aging over the period specified in the reliability requirement. Parts and structures should be designed to operate satisfactorily at the extremes of the parameter ranges, and allowable ranges must be included in the procurement or reprocurement specifications.

Statistical analysis and worst-case analysis can be used to assess the effects of part and structural parameter variations. In statistical analysis, a functional relationship is established between the output characteristics of the structure and the parameters of one or more of its parts. In worst-case analysis, the effect that a part has on the product output is evaluated on the basis of end-of-life performance values or out-of-specification replacement parts.

To ensure that the parts used in a system remain within the predetermined margins shown in Figure 5, derating can be used. Derating is the practice of limiting thermal, electrical, and mechanical stresses to levels below the manufacturer's specified ratings to improve reliability. Derating can provide added protection from system anomalies unforeseen by the designer (e.g., transient loads, electrical surges).

For example, manufacturers of electronic hardware often specify limits for supply voltage, output current, power dissipation, junction temperature, and frequency. The equipment design team may decide to select an alternative component or make a design change that ensures that the operational condition for a particular parameter, such as temperature, is always below the rated level. The stress reduction is expected to extend the useful operating life where the failure mechanisms under consideration are of the wear-out type. This practice is also expected to provide a safer operating condition by furnishing a margin of safety when the failure mechanisms are of the overstress type.

As inherently suggested by the term *derating*, the methodology involves a two-step process: a “rated” stress value is first determined from the part manufacturer’s data book, and a reduced value is then assigned. The margin of safety that the process of derating must provide is the difference between the maximum allowable actual applied stress and the demonstrated limits of the part capabilities. The part capabilities as given by manufacturer specifications are taken as the demonstrated limits. System design teams often choose to design a product based on conservative stress assumptions due to lack of knowledge of the life-cycle conditions, at the expense of overall productivity. There are reasons to believe that the part manufacturers already provide safety margin when choosing the operating limits. When these values are derated by the users, effectively a second level of safety margin is added.

In order to be effective, derating must target the appropriate, critical stress parameters based on models of the relevant failure mechanisms. Once the failure models for the critical failure mechanisms have been identified, using for example an FMMEA, the impact of derating on the effective reliability of the part for a given load can be determined. Instead of derating the stress rating values provided by the device manufacturers, the goal should be to determine the safe operating envelope for each part and subsystem and then operate within that envelope.

Redundancy

Redundancy exists when one or more of the parts of a system can fail and the system can still function with the parts that remain operational. A design team often finds that redundancy is the quickest way to

improve product reliability if there is insufficient time to explore alternatives, or if the part is already designed. It can be the most cost-effective solution, if the cost of redundancy is economical in comparison with the cost of redesign. It may even be the only solution if the reliability requirement is beyond the state of the art.

A redundant design may also exceed the limitations on size and weight, or power limitations. When not properly implemented, redundancy can also provide a false sense of security. If the cause can affect all the redundant elements of a product at the same time, then the benefits of redundancy will be lost. Also, the failure of sensing and switching circuitry or software can result in failure even in the presence of redundancy. Besides these possibilities, the design team may find that the disadvantages of redundancy may outweigh the benefits when the implementation of redundancy proves too expensive.

Integrated Diagnostics and Prognostics

Design guidelines and techniques should involve strategies to assess the reliability of the product in its life-cycle environment. A product’s health is the extent of deviation or degradation from its expected normal (in terms of physical condition and performance) operating condition. Knowledge of a product’s health can be used for the detection and isolation of faults or failures (diagnostics) and for the prediction of an impending failure based on the current conditions (prognostics). Thus, by determining the advent of failure, based on actual life-cycle conditions, procedures can be developed to mitigate and manage potential failures and maintain the product.

Diagnostics and prognostics can be integrated into a product by (a) installing built-in fuses and canary structures that will fail faster than the actual product when subjected to life-cycle conditions [7]; (b) sensing parameters that are precursors to failure, such as defects or performance degradation [8]; (c) sensing the life-cycle environmental and operational loads that influence the system’s health, and processing the measured data to estimate the life consumed [9, 10]. The life-cycle data measured by such integrated monitors can be extremely useful in future product design and end-of-life decisions [11]. An example of integrated prognostics in electronic products is the self-monitoring analysis and reporting

technology (SMART) currently employed for hard disk drives (HDD) in certain computing equipment [12]. HDD operating parameters, including flying height of the head, error counts, variations in spin time, temperature, and data transfer rates, are monitored to provide advance warning of failures. This is achieved through an interface between the computer's start-up program (BIOS) and the HDD.

Qualification

Qualification includes all activities that ensure that the nominal design and manufacturing specifications will meet the desired reliability targets. Qualification validates the ability of the nominal design and manufacturing specifications of the product to meet the customer's expectations, and assesses the probability of survival of the product over its complete life cycle. The purpose of qualification is to define the acceptable range of variability for all critical product parameters affected by design and manufacturing, such as geometric dimensions, material properties, and operating environmental limits. Product attributes that are outside the acceptable ranges are termed *defects*, since they have the potential to compromise product reliability [13].

Qualification tests should be performed during initial product development, and immediately after any design or manufacturing changes in an existing product. A well-designed qualification procedure provides economic savings and quick turnaround during development of new products or products subject to manufacturing and process changes.

Investigating failure mechanisms and assessing the reliability of products where longevity is required may be a challenge, since a very long test period under actual operating conditions is necessary to obtain sufficient data to determine actual failure characteristics. The results from FMMEA should be used to guide this process. One approach to the problem of obtaining meaningful qualification data for high-reliability devices in shorter time periods is using methods such as virtual qualification and accelerated testing to achieve test-time compression.

Virtual Qualification

Virtual qualification has the potential to significantly accelerate the qualification process of a part for

its life-cycle environment. Virtual qualification is based on computer-aided simulation, and can assist in identifying and ranking the dominant failure mechanisms associated with the part under life-cycle loads, determining the acceleration factor for a given set of accelerated test parameters, and determining the time to failure corresponding to the identified failure mechanisms.

Each failure model comprises a stress analysis model and a damage assessment model. The output is a ranking of different failure mechanisms, based on the time to failure. The stress model captures the product architecture, while the **damage model** depends on a material's response to the applied stress. Virtual qualification can be used to optimize the product design in such a way that the minimum time to failure of any part of the product is greater than its desired life. Although the data obtained from virtual qualification cannot fully replace those obtained from physical tests, they can increase the efficiency of physical tests by indicating the potential failure modes and mechanisms that can be expected.

Ideally, a virtual qualification process will involve identification of quality suppliers, quality parts, physics-of-failure qualification, and a risk assessment and mitigation program. The process allows qualification to be readily incorporated into the design phase of product development, since it allows design, manufacturing, and testing to be conducted promptly and cost effectively. It also allows consumers to qualify off-the-shelf components for use in specific environments without extensive physical tests. Since virtual qualification reduces emphasis on examining a physical sample, it is imperative that the manufacturing technology and quality assurance capability of the manufacturer be taken into account. If the data on which the virtual qualification is performed are inaccurate or unreliable, all results are suspect. In addition, if a reduced quantity of physical tests is performed in the interest of simply verifying virtual results, the operator needs to be confident that the group of parts selected is sufficient to represent the product. The effect of manufacturing variabilities can be parametrically assessed by simulation as part of the virtual qualification process. Further, it should be remembered that the accuracy of the results using virtual qualification depends on the accuracy of the inputs to the process – that is, the system geometry and material properties, the life-cycle loads, the failure models used (e.g., constants in the failure model),

the analysis domain, and discretization approach (spatial and temporal). Hence, to obtain a reliable prediction, the variability in the inputs should be specified using distribution functions, and the validity of the failure models should be tested by conducting accelerated tests.

Accelerated Testing

Accelerated testing is based on the concept that a product will exhibit the same failure mechanism and mode in a short time under high-stress conditions as it would exhibit in a longer time under actual life-cycle stress conditions. The purpose is to decrease the total time and cost required to obtain reliability information for the product under study. It is generally possible to quantitatively extrapolate from the accelerated environment to the usage environment with some reasonable degree of assurance.

Accelerated tests can be divided into two categories: qualitative tests and quantitative tests. Qualitative tests in the form of overstressing the products to obtain failure may be the oldest form of reliability testing. Such tests typically target a single load condition, such as shock, temperature extremes, or electrical overstress. Qualitative tests generally only yield failure mode information, and do not reveal the failure mechanism or time to failure. Quantitative tests target wear-out failure mechanisms, in which failures occur as a result of cumulative load conditions. Time to failure is the major outcome of quantitative accelerated tests.

The easiest and most common form of accelerated life testing is continuous-use acceleration. For example, a washing machine is used for 10 h per week on average. If it were continuously operated, the acceleration factor would be 16.8. However, this method is not applicable to high-usage products. Under such circumstances, accelerated testing is conducted to measure the performance of the test product at loads or stresses that are more severe than would normally be encountered, so as to enhance the damage accumulation rate within a reduced time period. The goal of such testing is to accelerate time-dependent failure mechanisms and the damage accumulation rate to reduce the time to failure. On the basis of the data from those accelerated tests, life in normal-use conditions can be extrapolated.

Accelerated testing begins by identifying all the possible overstress and wear-out failure mechanisms.

The load parameter that directly causes the time-dependent failure is selected as the acceleration parameter, and is commonly called the *accelerated load*. Common accelerated loads include thermal loads, such as temperature, temperature cycling, and rates of temperature change; chemical loads, such as humidity, corrosives, acids, solvents, and salts; electrical loads, such as voltage or power; and mechanical loads, such as vibration, mechanical load cycles, strain cycles, and shocks/impulses. The accelerated environment may include a combination of these loads. Interpretation of the results for combined loads requires a quantitative understanding of their relative interactions and the contribution of each load to the overall damage.

Failure due to a particular mechanism can be induced by several acceleration parameters. For example, corrosion can be accelerated by both temperature and humidity, and creep can be accelerated by both mechanical stress and temperature. Furthermore, a single accelerated stress can induce failure by several wear-out mechanisms simultaneously. For example, temperature can accelerate wear-out damage accumulation not only by electromigration, but also by corrosion and creep. Failure mechanisms that dominate under usual operating conditions may lose their dominance as the stress is elevated. Conversely, failure mechanisms that are dormant under normal-use conditions may contribute to device failure under accelerated conditions. Thus, accelerated tests require careful planning if they are to represent the actual usage environments and operating conditions without introducing extraneous failure mechanisms or nonrepresentative physical or material behavior. The degree of stress acceleration is usually controlled by an acceleration factor, defined as the ratio of the life of the product under normal-use conditions to that under the accelerated condition. The acceleration factor should be tailored to the hardware in question, and can be estimated from an acceleration transform (that is, a functional relationship between the accelerated stress and the life-cycle stress), in terms of all the hardware parameters.

Once the failure mechanisms are identified, it is necessary to select the appropriate acceleration load; to determine the test procedures and the stress levels; to determine the test method, such as constant-stress acceleration or step-stress acceleration; to perform the tests; and to interpret the test data, which includes extrapolating the accelerated test results to normal

operating conditions. The test results provide failure information for improving the hardware through design and/or process changes.

Manufacture and Assembly

Manufacturing and assembly processes can significantly impact the quality and reliability of hardware. Improper assembly and manufacturing can introduce defects, flaws, and residual stresses that act as potential failure sites, or stress enhancers or raisers later in the life of the product. The effect of manufacturing variability on time to failure is depicted in Figure 6. A shift in the mean or increase in the standard deviation of key geometric parameters during manufacturing can result in early failure due to a decrease in strength of the product. If these defects and stresses can be identified, the design analyst can proactively account for them during the design and development phase. Generally, qualification procedures are required to ensure that manufacturing specifications do not compromise the long-term reliability of the hardware. In some cases, lot-to-lot screening is required to ensure that the variability of all manufacturing-related parameters is within specified tolerances [14]. Here, screening ensures the quality of the product by precipitating latent defects before they reach the field.

Manufacturability

The design team must understand material limits, and manufacturing process capabilities to select materials and construct architectures that promote producibility and reduce the occurrence of defects. The team must have clear definitions of the threshold for acceptable quality and of what constitutes nonconformance. Nonconformance that compromises hardware performance and reliability is considered a defect.

A defect is any outcome of a process (manufacturing or assembly) that impairs or has the potential to impair the functionality of the product at any time. The defect may arise during a single process or may be the result of a sequence of processes. The yield of a process is the fraction of products that are acceptable for use in a subsequent manufacturing sequence or product life cycle. The cumulative yield of the process is approximately determined by multiplying the individual yields of each of the individual process steps. The source of defects is not always apparent, because defects resulting from a process can go undetected until the product reaches some downstream point in the process sequence, especially if screening is not employed.

It is often possible to simplify the manufacturing and assembly processes to reduce the probability of workmanship defects. As processes become more sophisticated, however, process monitoring and control are necessary to ensure a defect-free product. The bounds that specify whether the process is within tolerance limits, often referred to as the *process window*, are defined in terms of the independent variables to be controlled within the process and the effects of the process on the product or the dependent product variables. The goal is to understand the effect of each process variable on each product parameter to formulate control limits for the process – that is, the points on the variable scale where the defect rate begins to have a potential for causing failure. In defining the process window, the upper and lower limits of each process variable beyond which it will produce defects must be determined. Manufacturing processes must be contained in the process window by defect testing, analysis of the causes of defects, and elimination of defects by process control, such as by closed-loop corrective action systems. The establishment of an effective feedback path to report process-related

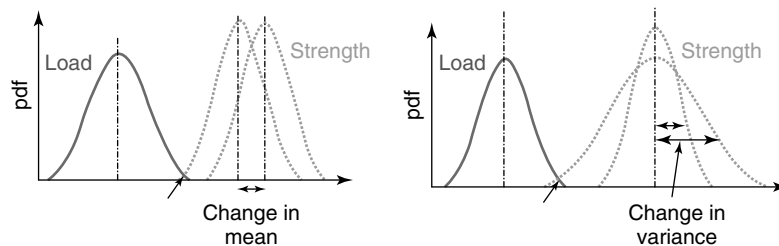


Figure 6 Influence of quality on failure probability

defect data is critical. Once this is done and the process window is determined, the process window itself becomes a feedback system for the process operator.

Several process parameters may interact to produce a different defect than would have resulted from an individual parameter acting independently. This complex case may require that the interaction of various process parameters be evaluated in a matrix of experiments. In some cases, a defect cannot be detected until late in the process sequence. Thus, a defect can cause rejection, rework, or failure of the product after considerable value has been added to it. These cost items due to defects can reduce a return on investment by adding to hidden factory costs. All critical processes require special attention for defect elimination by process control.

Process Qualification

Similar to design qualification, process qualification should be conducted at the prototype development phase. The objective is to ensure that the nominal manufacturing specifications and tolerances produce acceptable reliability in the product. The process needs requalification when process parameters, materials, manufacturing specifications, or human factors change.

Process qualification tests can be the same set of accelerated wear-out tests used in design qualification. As in design qualification, overstress tests may be used to qualify a product for anticipated field overstress loads. Overstress tests may also be exploited to ensure that manufacturing processes do not degrade the intrinsic material strength of the hardware beyond a specified limit. However, such tests should supplement, not replace, the accelerated wear-out test program, unless explicit physics-based correlations are available between overstress test results and wear-out field-failure data.

Process Verification Testing

Process verification testing is often called *screening*. Screening involves 100% auditing of all manufactured products to detect or precipitate defects. The aim of this step is to preempt potential quality problems before they reach the field. Thus screening aids in reducing warranty returns and

increases customer goodwill. In principle, screening should not be required for a well-controlled process. However, screening is often used as a safety net.

Some products exhibit a multimodal **probability density function** for failures, with peaks during the early period of their service life due to the use of faulty materials, poorly controlled manufacturing and assembly technologies, or mishandling. This type of early-life failure is often called *infant mortality*. Properly applied screening techniques can successfully detect or precipitate these failures, eliminating or reducing their occurrence in field use. Screening should be considered for use during the early stages of production, if at all and only, when products are expected to exhibit infant mortality field failures. Screening will be ineffective and costly if there is only one main peak in the failure probability density function. Further, failures arising due to unanticipated events such as acts of God (lightning, earthquakes) may be impossible to screen cost effectively.

Since screening is done on a 100% basis, it is important to develop screens that do not harm good components. The best screens, therefore, are nondestructive evaluation techniques, such as microscopic visual exams, X rays, acoustic scans, nuclear magnetic resonance, electronic paramagnetic resonance, and so on. Stress screening involves the application of stresses, possibly above the rated operational limits. If stress screens are unavoidable, overstress tests are preferred over accelerated wear-out tests, since the latter are more likely to consume some useful life of good components. If damage to good components is unavoidable during stress screening, then quantitative estimates of the screening damage, based on failure mechanism models, must be developed to allow the design team to account for this loss of usable life. The appropriate stress levels for screening must be tailored to the specific hardware. As in qualification testing, quantitative models of failure mechanisms can aid in determining screen parameters.

A stress screen need not necessarily simulate the field environment, or even utilize the same failure mechanism as the one likely to be triggered by this defect in field conditions. Instead, a screen should exploit the most convenient and effective failure mechanism to stimulate the defects

that can show up in the field as infant mortality. Obviously, this requires an awareness of the possible defects that may occur in the hardware and extensive familiarity with the associated failure mechanisms.

Unlike qualification testing, the effectiveness of screens is maximized when screening is conducted immediately after the operation believed to be responsible for introducing the defect. Qualification testing is preferably conducted on the finished product or as close to the final operation as possible; on the other hand, screening only at the final stage, when all operations have been completed, is less effective, since failure analysis, defect diagnostics, and troubleshooting are difficult and impair corrective actions. Further, if a defect is introduced early in the manufacturing process, subsequent value added through new materials and processes is wasted, which additionally burdens operating costs and reduces productivity.

Admittedly, there are also several disadvantages to such an approach. The cost of screening at every manufacturing station may be prohibitive, especially for small batch jobs. Further, components will experience repeated screening loads as they pass through several manufacturing steps, which increases the risk of accumulating wear-out damage in good components due to screening stresses. To arrive at a screening matrix that addresses as many defects and failure mechanisms as feasible with each screen test, an optimum situation must be sought through analysis of the cost effectiveness, risk, and criticality of the defects. All defects must be traced back to the root cause of the variability.

Any commitment to stress screening must include the necessary funding and staff to determine the root cause and appropriate corrective actions for all failed units. The type of stress screening chosen should be derived from the design, manufacturing, and quality teams. Although a stress screen may be necessary during the early stages of production, stress screening carries substantial penalties in capital, operating expense, and cycle time, and its benefits diminish as a product approaches maturity. If almost all of the products fail in a properly designed screen test, the design is probably faulty. If many products fail, a revision of the manufacturing process is required. If the number of failures in a screen is small, the processes are likely to be within tolerances and the observed faults may be beyond the resources of the design and production process.

Closed-Loop Monitoring and Root-Cause Analysis

Product reliability needs to be ensured using a closed-loop process that provides feedback to design and manufacturing in each stage of the product life cycle, including after the product is shipped and used in its application environment. Data obtained from periodic maintenance, inspection/testing, or health (condition) and usage monitoring methods can be used to perform timely maintenance for sustaining the product and preventing catastrophic failures. Figure 7 depicts the closed-loop process for managing the reliability of a product over the complete life cycle.

The objective of closed-loop monitoring is to analyze the failures occurring in testing and field conditions to identify the root cause of failure. The root cause is the most basic casual factor or factors that, if corrected or removed, will prevent recurrence of the situation. The purpose of determining the root cause(s) is to fix the problem at its most basic source so it does not occur again, even in other products, as opposed to merely fixing a failure symptom.

Root-cause analysis is a methodology designed to help (a) describe *what* happened during a particular occurrence; (b) determine *how* it happened; and (c) understand *why* it happened. Only when we are able to determine why an event or failure occurred, will we be able to determine the corrective measures over time. Root-cause analysis differs from troubleshooting, which is generally employed to eliminate a symptom in a given product, as opposed to finding a solution to the root cause to prevent this and other products from failing.

Correctly identifying the root cause during design and manufacturing, followed by appropriate actions to fix the design and processes, results in fewer field returns, major cost savings, and customer goodwill. Root-cause analysis of field failures can be more challenging if the conditions during and prior to failure are unknown. The lessons learned from each failure analysis need to be documented and appropriate actions need to be taken to update the design, manufacturing process, and/or maintenance actions or schedules.

After a product is developed, resources must be applied to managing its life cycle, including supply chain management, obsolescence assessment, manufacturing and assembly feedback, manufacturer

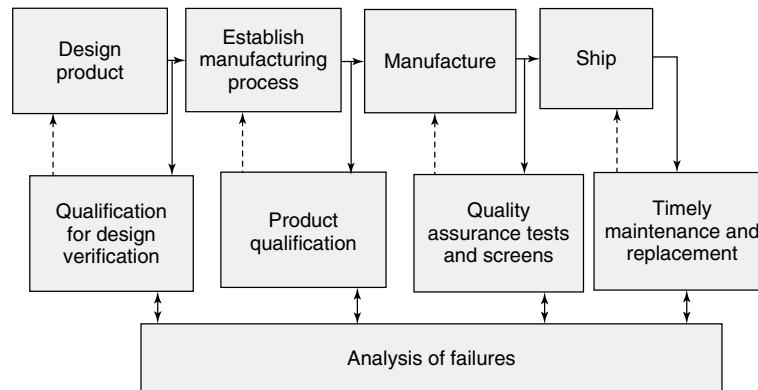


Figure 7 Reliability management using a closed-loop process

warranties management, and field failure and root-cause analysis. The risks associated with the product fall into two categories:

- managed risks – risks that the product development team chooses to proactively manage by creating a management plan and performing a prescribed monitoring regime of the field performance, manufacturer, and manufacturability; and
- unmanaged risks – risks that the product development team chooses not to proactively manage.

If risk management is considered necessary, a plan should be prepared. The plan should contain details about how the product is monitored (**data collection**), and how the results of the monitoring feed back into various product development processes. The feasibility, effort, and cost involved in management processes must be considered.

Summary

The development of a reliable product is not a matter of chance or good fortune; rather, it is a rational consequence of conscious, systematic, and rigorous efforts conducted through the entire life cycle of the product. High-product reliability can only be assured through robust product designs, capable processes that are known to be within tolerances, and qualified components and materials from vendors whose processes are also capable and within tolerances.

Quantitative understanding and modeling of all relevant failure mechanisms can guide design, manufacturing, and the planning of test specifications and tolerances.

When utilized early in the concept stage of a product's development, reliability serves as an aid to determine feasibility and risk. In the design stage of product development, reliability analysis involves the selection of materials, design of structures, design tolerances, manufacturing processes and tolerances, assembly techniques, shipping and handling methods, and maintenance and maintainability guidelines. Engineering concepts such as strength, fatigue, fracture, creep, tolerances, corrosion, and aging play a role in these design analyses. The use of physics-of-failure concepts coupled with mechanistic and probabilistic techniques are often required to understand the potential problems and trade-offs, and to take corrective actions. The use of factors of safety and worst-case studies as part of the analysis is useful in determining stress screening and **burn-in** procedures, **reliability growth**, maintenance modifications, field testing procedures, and various logistics requirements.

End Notes

^aOperational conditions are sometimes referred to as the *life-cycle application conditions*.

^bFor example, IPC SM-785 specifies the use and extreme temperature conditions for electronic products categorized under different industry sectors, such as telecommunication, commercial, and military.

^cDerating is the practice of subjecting parts to lower electrical or mechanical stresses than they can withstand, to increase the life expectancy of the part.

References

- [1] IEEE Standard 1332–1998 (1998). *IEEE Standard Reliability Program for the Development and Production of Electronic Systems and Equipment*, IEEE Reliability Society.
- [2] Vichare, N., Rodgers, P., Eveloy, V. & Pecht, M. (2004). In situ temperature measurement of a notebook computer – a case study in health and usage monitoring of electronics, *IEEE Transactions on Device and Materials Reliability* **4**(4), 658–663.
- [3] SAE Standard SAE J1739 (2002). *Potential Failure Mode and Effects Analysis in Design (Design FMEA) and Potential Failure Mode and Effects Analysis in Manufacturing and Assembly Processes (Process FMEA) and Effects Analysis for Machinery (Machinery FMEA)*.
- [4] IEEE Standard 1413.1–2002 (2003). *IEEE Guide for Selecting and Using Reliability Predictions Based on IEEE 1413*.
- [5] Bowles, J.B. (2003). Fundamentals of failure modes and effect analysis, Tutorial Notes, in *Annual Reliability and Maintainability Symposium*, Society of Automotive Engineering (SAE), Las Angeles.
- [6] Sage, A.P. & Rouse, W.B. (1999). *Handbook of Systems Engineering and Management*, John Wiley & Sons, Society of Automotive Engineering (SAE), New York.
- [7] Mishra, S. & Pecht, M. (2002). In-situ sensors for product reliability monitoring, *Proceedings of SPIE* **4755**, 10–19.
- [8] Pecht, M., Dube, M., Natishan, M. & Knowles, I. (2001). An evaluation of built-in test, *IEEE Transactions on Aerospace and Electronic Systems* **37**(1), 266–272.
- [9] Mishra, S., Pecht, M., Smith, T., McNee, I. & Harris, R. (2002). Remaining life prediction of electronic products using life consumption monitoring approach, in *Proceedings of the European Microelectronics Packaging and Interconnection Symposium*, Cracow, pp. 136–142, 16–18 June 2002.
- [10] Ramakrishnan, A. & Pecht, M. (2003). A life consumption monitoring methodology for electronic systems, *IEEE Transactions on Components and Packaging Technologies* **26**(3), 625–634.
- [11] Vichare, N., Rodgers, P., Azarian, M.H. & Pecht, M. (2004). Application of health monitoring to product take-back decisions, in *Joint International Congress and Exhibition – Electronics Goes Green 2004*+, Berlin, 6–8 September 2004.
- [12] Self-Monitoring Analysis and Reporting Technology (SMART) (accessed August 2004). PC Guide, <http://www.pcguides.com/ref/hdd/perf/qual/featuresSMART-c.html>.
- [13] Pecht, M., Dasgupta, A., Evans, J.W. & Evans, J.Y. (1994). *Quality Conformance and Qualification of Microelectronic Packages and Interconnects*, John Wiley & Sons, New York.
- [14] Upadhyayula, K. & Dasgupta, A. (1998). Guidelines for physics-of-failure based accelerated stress testing, in *Proceedings of Annual Reliability Maintainability Symposium*, New York, pp. 345–357.

MICHAEL PECHT

Software Failure Data Analysis

Some Definitions

Software failures are reported as a result of testing or regular usage by customers. They are recorded as modification requests (MRs), trouble reports (TRs), authorized program analysis reports (APARs), or some similar designation. Such reports document a departure from user requirements and performance deficiencies.

Developers can track the defect in the program that caused the failure. We call these defects *faults*. Another common term used to describe faults is *bugs*.

An error is the root cause of software faults that can potentially lead to software failures. This potential is realized when a user or a user-simulation triggers, under given conditions, an input stimulus that reveals the software fault, thus creating a failure.

These definitions of failures, faults and errors are based on the official terminology set in the IEEE Standard Glossary of Software Engineering Terminology [1] which specifies the following:

Error *Human action that results in software containing a fault. Examples include omissions or misinterpretation of user requirements in a software specification, incorrect translation, or omission of a requirement in the design specification.*

Fault (1) *An accidental condition that causes a functional unit to fail to perform its required function* (2) *A manifestation of an error in software. A fault, if encountered, may cause a failure. Synonymous with a bug.*

Failure (1) *The termination of the ability of a functional unit to perform its required function* (2) *An event in which a system component does not perform a required function within specified limits. A failure may be produced when a fault is encountered.*

Customer satisfaction is determined by experience with actual usage of the software. It is sometimes measured as *failures per 1000 CPU hours* or *failures per 1000 transactions*.

Process quality is determined by errors introduced during the process. Development process quality is often measured by *faults per 1000 source lines* or *faults per function point*.

The terms *defect*, *problem*, or *anomaly* are typically used in a generic sense. The terms *error*, *fault*, and *failure* will be reserved to have the technical definitions presented above. We will review how to track and analyze software failure data. Some basic IEEE references include [2], [3], and [4].

Problem Report Tracking

Problem report tracking provides the capability to gain insight into both product quality and the ability of the organization to maintain the software.

Figure 1 plots, by week, the total number of problem reports, those still open, and those that have been closed. At the start of a new test phase, such as software integration testing, or acceptance testing, the number of problem reports opened usually increases fairly rapidly, opening a gap between the number of problems reported and closed. If the problems are straightforward, we expect the gap to quickly decrease. If not, that could be an indication of a possibly serious condition or, alternatively, that the problems reported are of little operational consequence so that their resolution could be deferred to a later date. This suggests that to make more effective use of this chart, it is advisable to define categories of problem severity, and to overlay the number of problems open by severity category. For more on problem tracking and process improvement in software development see Kenett and Koenig [5], Humphrey [6], Kenett [7, 8], Fenton and Pfleeger [9], Kenett and Baker [10], Onoma and Yamura [11], Kan [12], and Laird and Brennan [13].

Control Charts of Newly Opened Bugs, by Week

Measures collected in real life environments often exhibit statistical variation. Such variation poses significant challenges to decision makers who attempt to filter out “signal” from “noise”. By signal we mean events with significant effects on the data such as short or long term trends and patterns or one time, sporadic spikes. We do expect the number of newly reported bugs to fluctuate.

It is in the context of such fluctuation that we want to determine if the number of opened bugs, in a certain day, exceeds the “normal” range of fluctuation. Another phenomenon we are interested

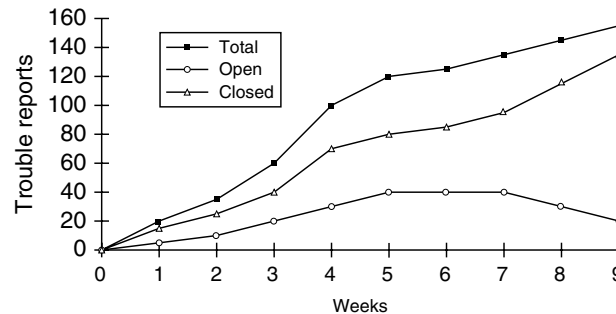


Figure 1 Problem report tracking

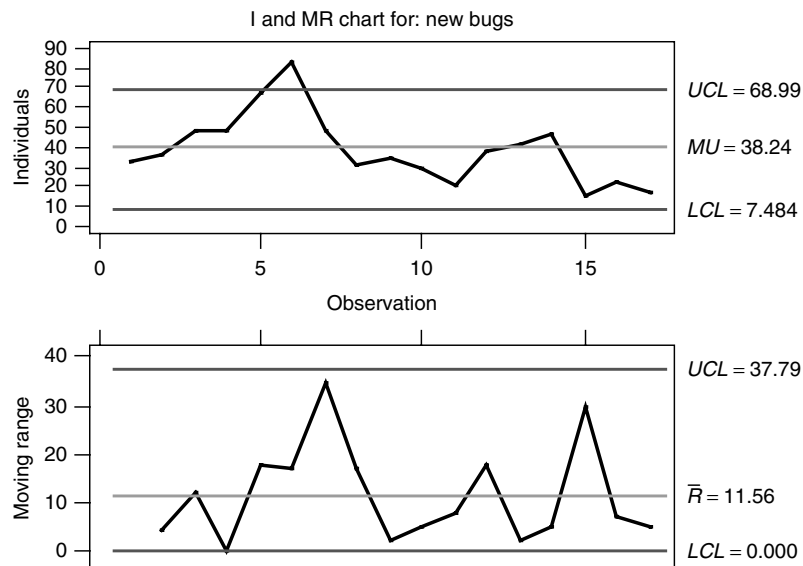


Figure 2 Control chart for number of newly opened bugs, by week

in identifying is a significantly decreasing trend in the number of opened bugs. Such a trend might be a prerequisite for a version release decision.

A primary tool for tracking data overtime, in order to identify significant trends and sporadic spikes, is the **control chart**. Control charts can be applied to individual data points such as the number of bugs opened, per week or to percentages, such as percent of line with comments. Sometimes data are collected simultaneously in several dimensions such as response times to a routine transaction at several stations in a network. Such multivariate data can be analyzed with **multivariate control charts**. The interested reader can find more information on control

chart in **Control Charts, Overview, Control Charts for the Mean, Moving Range and *R* Charts, Control Charts for Attributes**, and Kenett and Zacks [14].

Figure 2 consists of an individual (I) and a moving range (MR) control chart for the number of new bugs, per week, reported in a beta test of a new software product from a company developing rapid application development tools.

Analyzing the lower **MR chart** shows no unusual, nonrandom, patterns. The I chart however indicates a significant increase in the number of reported bugs on week 6 and a short run of three weeks below average toward the end. If this pattern is sustained

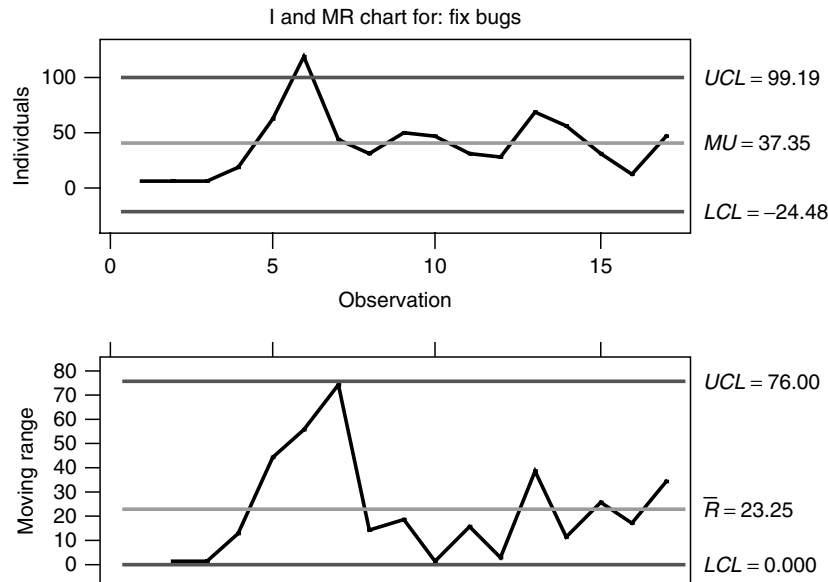


Figure 3 Control chart for number of closed bugs, by week

for three more weeks management could declare that the number of reported new bugs has significantly dropped, thus allowing for version release. Overall the average number of new bugs has been 38.2 per week. The next example tracks the number of closed or fixed bugs over the same time period.

Control Charts of Closed Bugs, by Week

In this example the number of closed bugs per week is tracked on an I and MR control chart (Figure 3). Analyzing the lower MR chart shows a period where the number of closed bugs per week remained almost constant (week 8–14), thus producing very low MRs (successive differences being almost zero). The I chart indicates a significant increase in the number of closed bugs on week 6 indicating a quick response to the significant increase in number of reported bugs identified above. Overall, the average number of bugs closed per week has been 37.3, just a bit lower than the number of reported new bugs per week.

Software Reliability Models

We proceed with an introduction to **software reliability** models that are used in the analysis of

software failure data. Software reliability estimation is an engineering management activity designed to

- Determine the reliability of delivered software
- evaluate the product at each milestone in the software life cycle
- formalize **data collection** and maintenance of historical data.

The basic steps in software reliability estimation and tracking are

1. Identify the system to be studied, including associated systems and hardware configurations that will be involved in system testing.
2. Determine reliability goals and classify failures into severity classes.
3. Develop operational profiles such as workload characterization of peak, prime, and off hours.
4. Prepare tools and procedures for testing according to the test plan.
5. Execute system tests.
6. Interpret failure data using tools such as **Pareto charts**, M-tests, and software reliability models.

The standard definition of software reliability is *the probability of execution without failure for*

some specified interval, called the mission time. This definition is compatible with that used for hardware reliability, though the failure mechanisms may differ significantly.

Software reliability is applicable both as a tool complementing development testing, in which faults are found and removed, and for certification testing, when a software product is either accepted or rejected. The now obsolete MIL-STD-1679 [15] document specified that one error per 70 000 machine instruction words, which degrades performance with a reasonable alternative work-around solution, is allowed; and one error causing inconvenience and annoyance is allowed for every 35 000 machine instruction words. However, no errors which cause a program to stop or degrade performance without alternative work-around solutions are allowed. These reliability specifications should be verified with a specified level of confidence.

Interpreting failure data is also used to determine if a product can be moved from development to beta testing with selected customers, and then from beta testing to official shipping. Both applications to development testing and certification testing rely on mathematical models for tracking and predicting software reliability. Many software reliability models have been suggested in the literature. For a review, see Wood [16], for an approach to determine when the software is ready for release, see Kenett and Pollak [17]. In this section, we cover two basic software reliability models. For more on software reliability models see **Software Reliability Modeling and Analysis**, **Software Reliability: Bayesian Analysis**, and **Sampling in Software Development**.

The Jelinski and Moranda “De-Eutrophication” Model

Jelinski and Moranda, while working for the McDonnell Douglas Astronautics Company, published one of the first practical software reliability models [18]. They developed a model for use on a number of modules of the Apollo program. The model assumptions are

1. The rate of error detection is proportional to the current error content of a program.
2. All errors are equally likely to occur and are independent of each other.

3. Each error is of the same order of severity as any other error.
4. The error rate remains constant over the interval between error occurrences.
5. The software is operated in a manner similar to the anticipated operational usage.
6. The errors are corrected instantaneously without introduction of new errors into the program.

Clearly, it is difficult to envision a situation in which a perfect error correction process is achieved. However assumption 6 can be avoided by not counting errors which were previously detected, but were not corrected. Assumption 2 can be avoided by dividing the errors into classes based upon severity. For instance, one might have a category for critical errors, a category for less serious errors, and one for minor errors. Separate software reliability models are then developed for each category. Using assumptions 1, 2, 5, and 6, the failure rate is defined as

$Z(t) = \phi[N - (i - 1)]$, where t is any time point between the discovery of the $(i - 1)$ th error and the i th error. The quantity ϕ is the proportionality constant given in assumption 1. N is the unknown total number of errors initially in the system. Hence, if $i - 1$ errors have been discovered by time t , there are $N - (i - 1)$ remaining errors so the hazard rate is proportional to this remaining number. Figure 4 is a plot of the hazard rate versus time. As can be seen, the rate is reduced by the same amount ϕ at the time of each error detection.

Let $X_i = t_i - t_{i-1}$, i.e., the time between the discovery of the i th and the $(i - 1)$ st error for $i = 1, \dots, n$ where $t_0 = 0$, using assumption 4, the X_i 's are assumed to have an exponential distribution with rate $Z(t_i)$. That is: $f(X_i) = \phi[N - (i - 1)] \exp\{-\phi[N - (i - 1)]X_i\}$.

The joint density of all the X_i , using assumption 2 is

$$\begin{aligned} L(X_1, \dots, X_n) &= \prod_{i=1}^n f(X_i) \\ &= \prod_{i=1}^n \phi[N - (i - 1)] \\ &\quad \times \exp\{-\phi[N - (i - 1)]X_i\} \quad (1) \end{aligned}$$

Taking the partial derivatives of $L(X)$ with respect to N and ϕ and setting the resulting equations equal to zero, the maximum likelihood estimators, $\hat{N}$ and

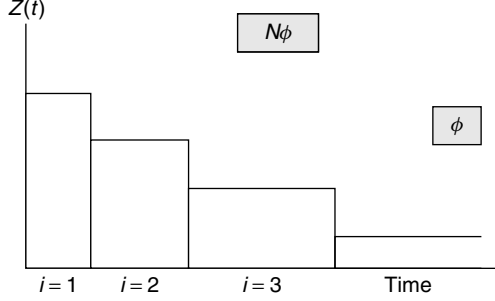


Figure 4 The De-Eutrophication process

$\hat{\phi}$, for N and ϕ are satisfy

$$\hat{\phi} = \frac{n}{\hat{N} \left(\sum_{i=1}^n X_i \right) - \sum_{i=1}^n (i-1)X_i} \quad (2)$$

and

$$\sum_{i=1}^n \frac{1}{\hat{N} - (i-1)} = \frac{n}{\hat{N} - \frac{1}{\sum_{i=1}^n X_i} \left(\sum_{i=1}^n (i-1)X_i \right)} \quad (3)$$

These equations can be solved using numerical techniques such as Newton–Raphson iterations.

The estimate of the mean time to failure (MTTF) is derived, after the j th occurrence, is

$$\begin{aligned} \text{Estimated MTTF after } j\text{th error} &= \frac{1}{\hat{Z}(t_j)} \\ &= \frac{1}{\hat{\phi}(\hat{N} - j)} \end{aligned} \quad (4)$$

The estimated time to remove the next m errors, after observing n failures, can be easily derived as:

$$\begin{aligned} \text{Estimated time to remove the next } m \text{ errors} \\ = \sum_{j=n+1}^{n+m} \frac{1}{\hat{\phi}(\hat{N} - j + 1)} \end{aligned} \quad (5)$$

The data required for the calculations of these estimates is either the time between error occurrences, X_i , or the time of error occurrences, t_i , $i = 1, \dots, n$.

The biggest problem in seeking these estimates is the difficulty in convergence of the numerical techniques employed to find the maximum likelihood

estimators. Difficulties include lack of convergence, sensitivity of the iteration scheme to the starting value, convergence to a saddle point or invalid estimate, and nonuniqueness of the estimates. Littlewood and Verall [19] show that a unique maximum at finite $\hat{N}$ and nonzero $\hat{N}$ is attained if, and only if,

$$\frac{\sum_{i=1}^n (i-1)X_i}{\sum_{i=1}^n (i-1)} > \frac{\sum_{i=1}^n X_i}{n} \quad (6)$$

If this condition is not satisfied, there is no convergence and $\hat{N}$ is infinity. Essentially this means that the model can only be applied to software that exhibits software growth, i.e. $X_i > X_{i-1}$. In any computer implementation of this model, the condition should first be verified to ensure a unique finite maximum exists.

The Jelinski–Moranda model cannot be applied to software programs which are not mature and are still under development. The program under test has to be relatively stable with a fixed but unknown total number of N errors present in the code. Various extensions of this basic model have been proposed including an adaptation to evolving a program which accounts for (i) the completion percentage, (ii) to error rates proportional to program size, (iii) to nonhomogeneous distributions of errors, (iv) to error rates proportional both to the number of errors remaining and to the time spent testing, and (v) of models accounting for varying testing efforts and programmers ability (see Kenett and Baker [10], Nelson [20], Musa *et al.* [21] and Neufelder [22]).

We present next an extension of the Jelinski–Moranda model proposed by Lipow [23]. The extension allows for more than one error occurrence during a testing interval. This is well suited to situations where data is aggregated overtime. Since failure data is usually reported on a weekly or monthly basis this extension proved very useful.

The Lipow Extension Model

The Lipow model assumes that

1. The rate of error detection is proportional to the current error content of a program.
2. All errors are equally likely to occur and are independent of each other.

3. Each error is of the same order of severity as any other error.
4. The error rate remains constant over the testing interval.
5. The software is operated in a manner similar to the anticipated operational usage.
6. During a testing interval i , f_i errors are discovered, but only n_i errors are corrected in the time frame.

Assumptions 1 to 5 are identical to the assumptions of the Jelinski–Moranda model. Assumption 6 is different. Suppose there are M periods of testing in which testing interval i is of length x_i . During this time frame, f_i errors are discovered, of which n_i are corrected. Assuming the error rate remains constant during each of the M testing periods (assumption 4), the failure rate during the i th testing period is

$$Z(t) = \phi[N - F_i], t_{i-1} \leq t \leq t_i \quad (7)$$

where ϕ is the proportionality constant, N is again the total number of errors initially present in the program, $F_{i-1} = \sum_{j=1}^{i-1} n_j$ is the total number of errors corrected up through the $(i-1)$ st testing intervals, and t_i is the time measured in either CPU or wall clock time at the end of the i th testing interval, $x_i = t_i - t_{i-1}$. The t_i 's are fixed and thus, are not fixed as in the Jelinski–Moranda model. Taking the number of failures, f_i , in the i th interval to be a Poisson random variable with mean $Z(t_i)x_i$, the likelihood is:

$$L(f_1, \dots, f_M) = \prod_{i=1}^M \frac{[\phi[N - F_{i-1}]x_i]^{f_i}}{f_i!} \times \exp\{-\phi[N - F_{i-1}]x_i\} \quad (8)$$

Taking the partial derivatives of $L(f)$ with respect to ϕ and N and setting the resulting equations to zero, we derive the following equations satisfied by the maximum-likelihood estimators $\hat{\phi}$ and $\hat{N}$ of ϕ and N :

$$\hat{\phi} = \frac{F_M/A}{\hat{N} + 1 - B/A} \quad (9)$$

and

$$\frac{F_M}{\hat{N} + 1 - B/A} = \sum_{i=1}^M \frac{f_i}{\hat{N} - F_{i-1}} \quad (10)$$

where

$F_M = \sum_{i=1}^M f_i$, the total number of errors found in the M periods of testing,

$B = \sum_{i=1}^M (F_{i-1} + 1)x_i$, and

$A = \sum_{i=1}^M x_i$, the total length of the testing period.

From these estimates, the maximum likelihood estimate of the mean time until the next failure (MTTF), given the information accumulated in the M testing periods $= 1/(\hat{\phi}(\hat{N} - F_M))$.

A Case Study of the Lipow Model

In order to demonstrate the application of Lipow's extension to the Jelinski and Moranda De-Eutrophication model, we use the data collected at Bell Laboratories during the development of the No. 5 electronic switching system, in Table 1.

The data consists of failures observed during four consecutive testing periods of equal length and intensity.

The Lipow model produces the following estimates: $\hat{N} = 762$ and $\hat{\phi} = 0.097$, i.e. an estimate of 762 for the initial number of failures and 0.097 for the proportionality constant representing the reduction in failure rate due to removing one fault from the system.

The column of predicted failures in test periods 5, 6, and 7 is computed from the Lipow model with the above estimates for ϕ and N . Comparing observed and predicted values over the first four test periods can help us evaluate how well the Lipow model fits the data. The χ^2 goodness of fit test statistics is 0.63, indicating extremely good fit (at a **significance level** better than 0.001). We can therefore provide reliable predictions for the number of failures expected at test periods 5–7, provided the test effort continues at the same intensity and that test periods are of the same length as before.

The Jelinski–Moranda model and the various extensions to this model are classified as time domain models. They rely on a physical modeling of the appearance and fixing of software failures. The different sets of assumptions are translated into differences in mathematical formulations (see Musa *et al.* [21], Neufelder [22] and Brown and Lipow [24]). A model free approach to tracking software **reliability growth**, that exploits optimality properties of Bayesian procedures known as *Shiryayev–Roberts procedures*, has been proposed by Kenett and Pollak [17, 25].

When software is in operation, failure rates can be computed by computing number of failures, say,

Table 1 Predicted failures using Lipow model

Test Period	Observed Failures	Predicted Failures
1	72	74
2	70	67
3	57	60
4	50	54
5		49
6		44
7		40

per hour of operation. The predicted failure rate, corresponding to the steady state behavior of the software, is usually a key indicator of great interest. The predicted failure rate may be regarded as high when compared to other systems. Kanoun *et al.* [26] analyze four systems (telephony, defense, interface, and management) and derive average failure rates over 10 testing periods and predicted failure rates. Their results are presented in Table 2.

As expected in any reliability growth process, the average failure rate over 10 test periods is higher than the predicted failure rates. However, the differences in the defense and management systems are smaller indicating that these systems have almost reached steady state, and no significant reliability improvements should be expected. The telephony and interface systems, however, are still evolving. A decision to promote the defense and interface systems to acceptance test status could be made, based on the data.

Pragmatic Software Reliability Estimation

A simple model for estimating software reliability has been proposed by Wall and Ferguson [27]. The model relies on the basic premise that the

failure rate of software decreases as more software is used and tested. The relationship proposed for this phenomenon is

$$C1 = C0 \left(\frac{M1}{M0} \right)^a \quad (11)$$

where,

$C1$ = future period cumulative number of errors

$C0$ = previous period cumulative number of errors – a constant determined empirically

$M1$ = units of test effort in future period (e.g. number of tests, testing weeks, CPU time)

$M0$ = test units in previous period – a scaleable constant

a = growth index – a constant determined, empirically, by plotting the cumulative number of errors on a logarithmic scale or by using the following formula:

$$a = \frac{\log \left(\frac{C1}{C0} \right)}{\log \left(\frac{M1}{M0} \right)} \quad (12)$$

The failure rate, $Z(t)$, is derived from the following formula:

$$Z(t) = \frac{dC1}{dt} = aC0 \frac{d \left(\frac{M1}{M0} \right)}{dt} \left(\frac{M1}{M0} \right)^{a-1} \quad (13)$$

This model is based on several assumptions:

1. software test domain is stable i.e. no new features are introduced.
2. software errors are independent.
3. errors detected are corrected within the same test period.

Table 2 Predicted software failure rates, by system type

System	Predicted failure rate	Average failure rate over 10 test periods
Telephony	$1.2 \times 10^{-6} \text{ h}^{-1}$	$1.0 \times 10^{-5} \text{ h}^{-1}$
Defense	$1.4 \times 10^{-5} \text{ h}^{-1}$	$1.6 \times 10^{-5} \text{ h}^{-1}$
Interface	$2.9 \times 10^{-5} \text{ h}^{-1}$	$3.7 \times 10^{-5} \text{ h}^{-1}$
Management	$8.5 \times 10^{-6} \text{ h}^{-1}$	$2.0 \times 10^{-5} \text{ h}^{-1}$

4. testing is incremental so that new tests are introduced in subsequent testing periods.

Revisiting the data used in the previous section, with an empirically derived estimate of 0.8 for a , we get the estimated number of failures for test period 5 to be 47.17 failures (see Table 3). The Lipow model predicted 49 failures.

Pareto Chart Analysis of Software Failures

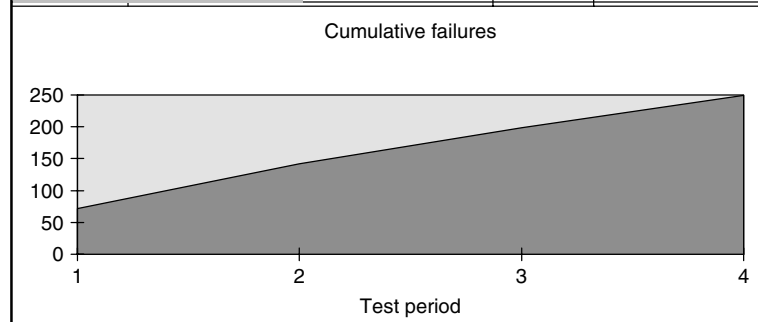
When observing events, it is a universal phenomenon that approximately 80% of events are due to 20% of the possible causes. A classical application to software is the general fact that 80% of software failures can be attributed to 20% of the code. This universal observation was first made by Joseph M. Juran who, in the early 1950s, coined the term “Pareto principle” which leads to the distinction between the “vital few” and the “useful many” (see **Pareto Chart**). The Pareto chart is a graphical display of the Pareto principle. It consists of bar graphs sorted, by category, in descending order of the relative frequency of errors. Pareto charts are used to choose the starting point for problem solving, monitor changes, or identify the basic cause of a problem. An interesting example of a Pareto chart analysis of software failures is provided by data from Knuth [28] on changes made during a period of 10 years, in the development of TEX, a software system for typesetting. Knuth’s logbook contains 516 items for

the 1978 version, which is labeled TEX78, and 346 items for the 1982 version, labeled TEX82. These entries are classified into 15 categories ($K = 15$). The 15 categories are

- A – Algorithm: a wrong procedure being implemented
- B – Blunder: a mistaken but syntactically correct instruction
- C – Cleanup: an improvement in consistency or clarity
- D – Data: a data structure debacle
- E – Efficiency: a change for improved code performance
- F – Forgotten: a forgotten or incomplete function
- G – Generalization: a change to permit future growth and extensions
- I – Interaction: a change in user interface to better respond to user needs
- L – Language: a misunderstanding of the programming language
- M – Mismatch: a module’s interface error.
- P – Portability: a change to promote portability and maintenance
- Q – Quality: an improvement to better meet user requirements
- R – Robustness: a change to include sanity checks of inputs
- S – Surprise: a change resulting from an unforeseen interaction
- T – Typo: change to correct differences between pseudocode and code.

Table 3 Predicted failures using the pragmatic model

Test period	Observed failures	Cumulative failures	a	Predicted failures
1	72	72		
2	70	142	0.8	125.36
3	57	199	0.8	196.41
4	50	249	0.8	250.50
5	47.17		0.8	297.66



The A, B, D, F, L, M, R, S, T classifications represent development errors. The C, E, G, I, P, Q classifications represent “enhancements”, consisting of unanticipated features that had to be added in late development phases. These enhancements indicate a lack of understanding of customer requirements, and, as such, can be considered failures of the requirements analysis process. Figure 5 presents a Pareto chart of the TEX78 data. The chart does not show any dominant error category so that, apparently, the Pareto principle does not apply to this data. Figure 6 is the Pareto chart of the TEX82 data. Here we notice that categories C, R, G, I, and E contribute 71% of the errors. Categories C, G, I, and E being classified as “enhancements” which are therefore the main cause behind the logbook entries.

TEX82 contains significantly *more* errors in the Cleanup (C), Efficiency (E), and Robustness (R) categories than TEX78. Significantly *less* errors are found in Blunder (B), Forgotten (F), Language (L), Mismatch (M), and Quality (Q). Knuth gives us only clues as to the possible reasons for these differences. For more on statistical comparisons of Pareto charts see Kenett and Zacks [14] and Kenett and Baker [10]. We conclude by discussing the analysis of software failures uncovered in inspection and design reviews.

Software Trouble Assessment

The IEEE 730.1 standard on Software Quality Planning [29] stipulates that, as a minimum, eight types of reviews are to be conducted:

1. Software requirements review (SRR)
2. Preliminary design review (PDR)
3. Critical design review (CDR)
4. Software verification review
5. Functional audit
6. Physical audit
7. In-process audit
8. Managerial reviews.

These reviews generate data that can be further analyzed through root cause analysis.

The Software Trouble Assessment Matrix (STAM) is a tool that software developers can use to evaluate the design and effectiveness of software inspection and testing processes so that they can be improved [10, 30]. STAM is used to organize relationships between three dimensions:

- Where were failures detected in the software life cycle?
- Where were those failures actually created?
- Where could the failures have been detected?

Three measures are computed from a STAM analysis:

- **Negligence ratio:** this ratio indicates the amount of failures that escaped through the inspection process filters. In other words, it measures inspection efficiency.
- **Evaluation ratio:** this ratio measures the delay of the inspection process in identifying failures relative to the phase in which they occurred. In other words, it measures inspection effectiveness.
- **Prevention ratio:** this ratio is an index of how early failures are detected in the development life cycle relative to the total number of reported errors. This is a combined measure of the development and inspection processes. It assesses the software developer’s ability to generate and identify failures as early as possible in the development life cycle.

Process improvements are the result of attempts to learn from current and past failures. Prerequisites to such improvement efforts are that the processes have been identified and that process ownership has been established. Typical development processes consist of requirements analysis, top-level design, detailed design, coding, and testing. A causal analysis of software faults classifies faults as being attributable to errors in any one of these activities. For such an analysis to be successful, it is essential that there be agreement on the boundaries of the processes represented by these activities. In particular, the entry and exit criteria for each process have to be clarified and documented. Such definitions permit effective data collection.

As an example, consider a certain software version with a total of 110 reported failures at the completion of acceptance testing (see Table 4). The failures were reported throughout the software life cycle, with the following distribution:

Of the seven failures detected during top-level design, it was determined that two failures could have been detected in the review session that occurred after requirements analysis. These failures were missed by the reviewers. Of the 13 failures detected during acceptance testing, it was determined that one failure

could have been detected at the requirements analysis review session, one could have been detected after detailed design, six could have been detected during system testing, and five could have been detected only during acceptance testing. The implication is that eight failures escaped the inspection process filters. A similar analysis indicated that, of the five failures that could only have been detected during acceptance testing, one failure was due to a

requirement error, three were the result of errors in preliminary design, and one was due to a detailed design error.

A typical failure analysis begins with assessing where failures could have been detected and concludes with classifying failures into the life-cycle phases in which they were created. This procedure requires a root cause analysis of each recorded failure. As previously mentioned, the success of a failure

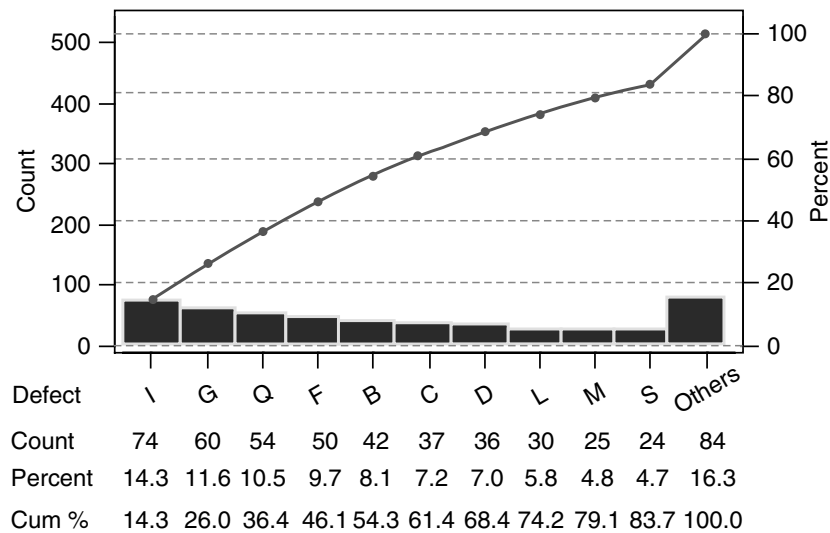


Figure 5 Pareto chart of TEX78 logbook data

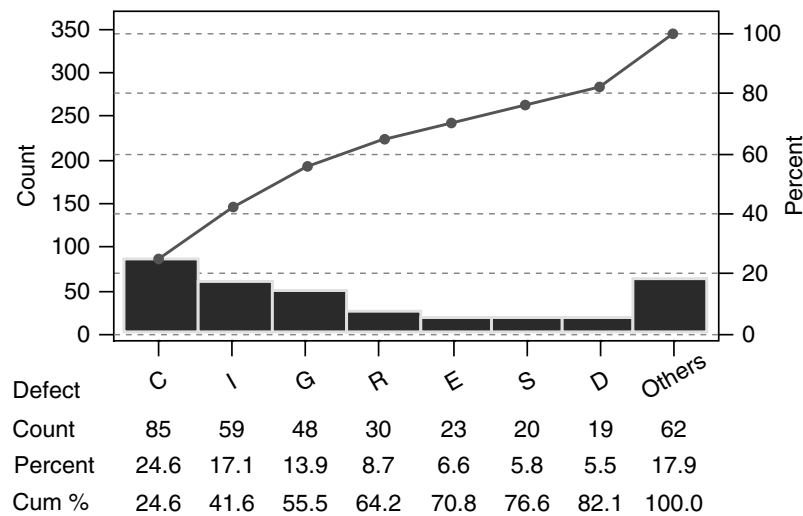


Figure 6 Pareto chart of TEX82 logbook data

Table 4 Software failures by detection phase

Life-cycle phase	Number of failures
Requirements analysis	3
Top-level design	7
Detailed design	2
Programming	25
Unit tests	31
System tests	29
Acceptance test	13

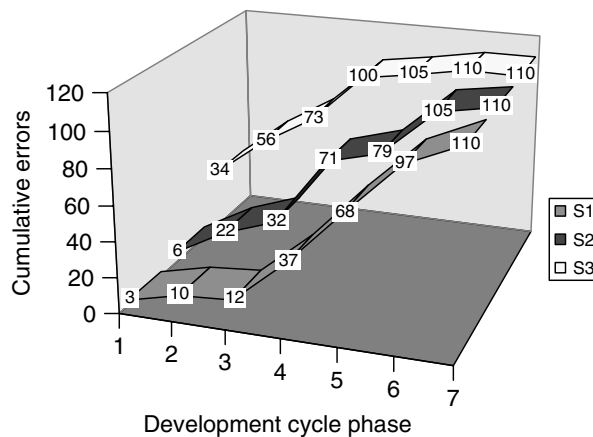
causal analysis is highly dependent on clear entry and exit criteria for the various development phases.

Once the STAM analysis is completed, cumulative failure profiles are drawn depicting where errors were detected, where they could have been detected, and where they were created. The areas under these three cumulative frequency curves are defined as $S1$, $S2$, and $S3$, respectively (see Figure 7). These areas are determined by computing the cumulative totals over the seven software development phases.

A curve is drawn by connecting the results of these additions along the seven development phases. The area under the curve, labeled $S1$, is approximated by adding these numbers: $S1 = 3 + 10 + 12 + 37 + 68 + 97 + 110 = 337$. By referring to Figure 7, we see that the same procedure is followed to get the $S2$ and $S3$ computations ($S2 = 427$ and $S3 = 588$). The negligence ratio is computed using the formula: $100 \times (S2 - S1)/S1$. As previously mentioned, it measures the amount of errors that escaped through the inspection process filters, indicating inspection

efficiency. High inspection efficiency corresponds to low negligence ratios. For the data in Figure 7, the negligence ratio is $100 \times (427 - 337)/337 = 26.7\%$, which indicates an average gap of 26.7% of a life-cycle phase between actual error detection time and perfect detection under the current inspection process.

The evaluation ratio is derived using the formula: $100 \times (S3 - S2)/S2$. As previously mentioned, it measures the delay of the inspection process in identifying errors from the phase where they were created, indicating inspection effectiveness. High evaluation ratios correspond to low inspection effectiveness. For the data in Figure 7, the evaluation ratio is $100 \times (588 - 427)/427 = 37.7\%$, which signals the need to redesign the inspection process so that it can detect errors closer to their creation. The current inspection filters, under perfect conditions, detect errors with a delay of 37.7% of a life-cycle phase. The prevention ratio is computed using the formula: $100 \times S1/(7 \times total)$. As previously mentioned, it indicates how early errors were detected relative to the total number of reported errors (*total*) in the seven development phases. If all errors were created and detected during requirements analysis, the prevention ratio would be 100%. Since early errors are less costly to correct, a high prevention ratio implies a less costly development process. Conversely, a low prevention ratio indicates that errors were detected late in the process and therefore might negatively affect delivery schedules and customer satisfaction. In Figure 7,

**Figure 7** Graph of the curves determining the $S1$, $S2$, and $S3$ areas

the prevention ratio is 43.7%, which indicates that errors were detected late in the life cycle. Using the negligence, evaluation, and prevention ratios, software developers can better understand and improve their inspection and development processes. They also can use STAM to benchmark different projects within their companies and against those of different companies.

References

- [1] ANSI/IEEE 729. (1983). IEEE standard glossary of software engineering terminology, IEEE Standards Office, P.O. Box 1331, Piscataway.
- [2] ANSI/IEEE 982.1. (1988). Standard dictionary of measures to produce reliable software, IEEE Standards Office, Piscataway, revised draft 2005.
- [3] IEEE 1044.1. (1995). Guide to IEEE standard for classification for software anomalies, IEEE Standards Department, Piscataway.
- [4] IEEE Standard 1044. (1993). Standard classification for software anomalies, IEEE Standards Office, Piscataway.
- [5] Kenett, R.S. & Koenig, S. (1988). A process management approach to software quality assurance, *Quality Progress* 66–70.
- [6] Humphrey, W. (1989). *Managing the Software Process*, Addison-Wesley, New York.
- [7] Kenett, R.S. (1989). Managing a continuous improvement of the software development process, in *Proceedings of the 8th IMPRO Conference*, Atlanta.
- [8] Kenett, R.S. (1992). Understanding the software process, in *Total Quality in Software Development*, G. Schulmeyer & J. McManus, eds, Van Nostrand Reinhold.
- [9] Fenton, N. & Pfleeger, S. (1997). *Software Metrics: A Rigorous and Practical Approach*, PWS Publishing, Boston.
- [10] Kenett, R.S. & Baker, E. (1999). *Software Process Quality: Management and Control*, Marcel Dekker, New York.
- [11] Onoma, A. & Yamura, T. (1995). Practical steps toward quality development, *IEEE Software* 68–76.
- [12] Kan, S. (2001). *Metrics and Models in Software Quality Engineering*, Addison-Wesley, Boston.
- [13] Laird, L.M. & Brennan, M.C. (2006). *Software Measurement and Estimation: A Practical Approach*, John Wiley & Sons, Hoboken.
- [14] Kenett, R.S. & Zacks, S. (1998). *Modern Industrial Statistics*, Duxbury Press, San Francisco.
- [15] MIL-STD-1679 (Navy). (1978). Military standard weapon system software development, Department of Defense, Washington, DC.
- [16] Wood, A. (1996). Predicting software reliability, *IEEE Software* 69–77.
- [17] Kenett, R.S. & Pollak, M. (1986). A semi-parametric approach to testing for reliability growth, with application to software systems, *IEEE Transactions on Reliability* **R-35**(3), 304–311.
- [18] Moranda, P. & Jelinski, Z. (1972). Final report on software reliability study, MDC Report No. 63921, McDonnell Douglas Astronautics Company.
- [19] Littlewood, B. & Verall, B. (1973). A Bayesian reliability growth model for computer software, *The Journal of the Royal Statistical Society. Series C* **22**(3), 332–346.
- [20] Nelson, E. (1978). Estimating software reliability from test data, *Microelectronics and Reliability* **17**, 67–73.
- [21] Musa, J., Iannino, A. & Okumoto, K. (1987). *Software Reliability Measurement, Prediction, Application*, McGraw-Hill, New York.
- [22] Neufelder, A. (1993). *Ensuring Software Reliability*, Marcel Dekker.
- [23] Lipow, M. (1978). Models for software reliability, in *Proceedings of the Winter Meetings of the Aerospace Division of the American Society for Mechanical Engineers*, 78-WA/Aero-18, pp. 1–11.
- [24] Brown, J.R. & Lipow, M. (1975). Testing for software reliability, *Proceedings of the International Conference on Reliable Software*, IEEE.
- [25] Kenett, R.S. & Pollak, M. (1996). Data-analytic aspects of the Shiryayev-Roberts control chart: surveillance of a non-homogeneous Poisson process, *Journal of Applied Statistics* **23**(1), 125–127.
- [26] Kanoun, K., Kaaniche, M. & Laprie, J.-C. (1997). Qualitative and quantitative reliability assessment, *IEEE Software* 77–87.
- [27] Wall, J.K. & Ferguson, P.A. (1977). Pragmatic software reliability prediction, *Proceedings of the Annual Reliability and Maintainability Symposium* 485–488.
- [28] Knuth, D. (1988). The errors of TEX, Report No. STAN-CS-88-1223, Department of Computer Science, Stanford University, Stanford.
- [29] NASI/IEEE 730.1, (1989). Software quality planning, IEEE Standards Office, Piscataway.
- [30] Kenett, R.S. (1994). Assessing software development and inspection errors, *Quality Progress*, **27**(10), 109–112, October 1994 with correction in Vol. 28, number 2, February 1995.

Related Articles

Software Reliability Modeling and Analysis; Software Reliability; Bayesian Analysis.

RON S. KENETT

Life Testing from a Bayesian Perspective

The Role of Life Testing

Life testing is the name given to a variety of testing methods carried out to obtain failure data (*see Reliability Sampling*). Roughly speaking, life testing consists in putting on test a random sample of n units drawn from a (hypothesized) population of identical units under specified operating and environmental conditions. Clearly, the testing method changes according to a large variety of factors, among them the fact that the unit can be repaired or is discarded at failure.

In the case of unrepairable units, the life testing can be stopped when the failure of some or all the n units put on test occurs. If the failure time of each unit is observed, then the test is said to be a complete life test. When it is not practical, or economically convenient, to test all the units, for example due to a very high reliability level of the unit, an incomplete life testing is performed. Clearly, a large variety of incomplete tests can be performed. For example, the test can be stopped after a fixed period has elapsed, so that a random number m ($m < n$) of units is observed to fail. In this case, the test is said to be a time-truncated test. Again, if the test is terminated when a prefixed number of units failed, the test is said to be a failure-truncated test. In both the cases, some units may also have been withdrawn from the life test before they fail, so only their survival times are observed.

In case of repairable units, which are repaired at failure, the test is always truncated because the unit can achieve a theoretically infinite life, and more that one failure of each unit is generally observed.

When testing time is limited, the units are often subjected to accelerated life tests (*see Accelerated Life Tests: Classical Methods for Design and Analysis; Accelerated Life Tests: Bayesian Design*). Accelerated life testing is the process by which the units are forced to fail more quickly than they would under normal use conditions, so test time can be greatly reduced.

Irrespective of the testing method, the number of observed failures determines the accuracy of the resulting reliability estimates.

In classical estimation procedures, such as maximum-likelihood and least-squares methods, life tests constitute the exclusive source of information, whereas in a Bayesian perspective the purpose of life testing is to “confirm or deny expected performance as predicted from past experience” [1, p. 169]. In other words, life testing provides the Bayesian analyst new information (the so-called sampling information), which allows the current belief held prior to observing the new data to be revised (*learning from experience*). The revised belief, expressed in terms of subjective probability, constitutes the inferential result. Clearly, in a decision-making approach, revising belief leads the expert to confirm or change the choice he would have made among alternative courses of action on the basis of prior belief only.

The way in which the observed data modify the current belief is formalized by Bayes’ theorem [2] (*see Bayesian Reliability Analysis*). Bayes’ theorem for parameters and observations states that the posterior knowledge of Φ , formalized by the posterior distribution $\pi(\Phi|\mathbf{x})$, is obtained from the observed data $\mathbf{x}$ and the prior information of Φ , formalized by the prior distribution $g(\Phi)$, according to

$$\pi(\Phi|\mathbf{x}) = \frac{f(\mathbf{x}|\Phi)g(\Phi)}{f(\mathbf{x})} \quad (1)$$

where Φ is the vector of parameters which index the failure time model, $f(\mathbf{x}|\Phi)$ is the joint conditional distribution of $\mathbf{x}$ given Φ , and $f(\mathbf{x})$ is the marginal distribution of the data given by

$$f(\mathbf{x}) = \begin{cases} \int f(\mathbf{x}|\Phi) g(\Phi) d\Phi, & \Phi \text{ continuous} \\ \sum f(\mathbf{x}|\Phi) g(\Phi), & \Phi \text{ discrete} \end{cases} \quad (2)$$

If no prior information is available and hence the noninformative prior distribution (*see Prior Distribution Elicitation*) is used [3, pp. 25–59], life testing becomes the only source of information in a Bayesian framework also.

The main difficulties in performing a Bayesian analysis regard just the identification, selection, and justification of the prior distribution (*see Prior Distribution Elicitation*). The problem of prior identification concerns whether to use a continuous or discrete distribution, a noninformative or informative prior, or a mathematically convenient prior

distribution, such as a conjugate prior. Once the prior has been identified, the problem of prior selection concerns the fitting of the identified prior distribution on the basis of the available subjective data. This problem mainly concerns the choice of the procedure to be used in estimating the parameters of the prior distribution (often called the *hyperparameters*) and the way the subjective data have to be solicited and elicited. Finally, adequate justification for the prior selection should be given by the Bayesian analyst, by including a clear description of the data sources used to fit the prior distribution, the method used to fit the prior to the data, and a preposterior analysis of the fitted prior. Otherwise, the prior distribution can be suspected to have been chosen with bias to provide self-serving results and hence the resulting inferential results can be criticized [1, p. 186]. Wide overviews on the identification, selection, and justification of the prior distribution can be found, for example, in [1, Chapter 6] and [4, Chapter 3].

In calculating the posterior distribution, the concept of sufficiency can be helpful. Intuitively, a sufficient statistic is a function of the data $\mathbf{x}$ that summarizes all the available sampling information regarding Φ . Then, if $T = T(\mathbf{x})$ is a sufficient statistic for Φ with probability distribution $f(T|\Phi)$, Bayes theorem can be rewritten as

$$\pi(\Phi|\mathbf{x}) = \pi(\Phi|T) = \frac{f(T|\Phi)g(\Phi)}{f(T)} \quad (3)$$

The reason for deriving the posterior distribution $\pi(\Phi|\mathbf{x})$ from equation (3) is that both $f(T|\Phi)$ and the marginal distribution $f(T)$ are generally much easier to handle than $f(\mathbf{x}|\Phi)$ and $f(\mathbf{x})$ [4, p. 93].

Thus, the way the observed data modify the prior knowledge of Φ is through the joint conditional distribution $f(\mathbf{x}|\Phi)$ (or, equivalently, through the $f(T|\Phi)$). Bayes theorem works in a way that increases the belief in parameters values that are highly compatible with the observed data, whereas the belief in parameters values which are unlikely to have resulted in the data actually observed decreases. Then, if different experts possess different prior beliefs on Φ , the sampling information makes their posterior beliefs closer to each other than the prior beliefs are. Clearly, the more the sampling information dominates over the prior beliefs, the closer the posterior distributions are.

As newer data are observed, today's posterior information becomes tomorrow's prior information,

so that each time new data come in the current belief can be revised. In this context, the use of conjugate prior distributions (*see Prior Distribution Elicitation*) is very attractive since the posterior distributions are members of the same family of distributions but with parameters updated by the observed data. This updating of parameters provides an easy way of revising the current belief in the light of new observed data and of seeing the effect of prior and sample information on posterior distribution.

However, when the prior information dominates over the sampling information, slight changes in the prior can cause significant changes in the posterior results [1, p. 85]. Then, especially when the observed data are scarce, a sensitivity analysis should be performed to investigate how possible changes in the prior parameters (or, more generally, in the prior information) affect the posterior results. Such an analysis can identify the most important prior information and focus the effort of the analyst to better refine the form of the prior distribution or the values of the prior parameters. Again, in a decision-making context, the sensitivity analysis allows the analyst to assess whether reasonable changes in the prior distribution and/or in the loss function can lead to a different decision (see, for example, references [4, 5]). A comprehensive overview of the main topics in **Bayesian robustness** is provided by reference [6].

Simple illustrations of Bayesian learning process are given in the following examples.

Example 1

An upgraded type of reciprocating pump has been designed and the reliability team must verify, before starting mass production, whether the pumps satisfy a given reliability target at the lifetime $\tau = 0.7 \times 10^6$ cycles. The reliability target is assumed to be satisfied if the 90% lower probability limit on $R_\tau \equiv R(\tau)$, say $R_{0.10}$, is greater than 0.60. The failure times are assumed to be exponentially distributed

$$f(x|\theta) = \theta \exp(-\theta x), \quad x \geq 0 \quad (4)$$

where the parameter $\theta > 0$ is the (constant) failure rate. On the basis of previous experience, expert A possesses prior information on θ expressed in terms of a prior mean $E_A\{\theta\} = 0.58 \times 10^{-6}$ per cycle and a prior standard deviation $\sigma_A\{\theta\} = E_A\{\theta\}/2 = 0.29 \times$

10^{-6} per cycle. Expert A uses a gamma distribution to formalize the prior belief on θ :

$$g_A(\theta) = \frac{b^a \theta^{a-1}}{\Gamma(a)} \exp(-b\theta), \quad a, b > 0 \quad (5)$$

whose parameters a and b (often called *hyperparameters*) are $a = (E_A\{\theta\}/\sigma_A\{\theta\})^2 = 4.0$ and $b = E_A\{\theta\}/\sigma_A^2\{\theta\} = 6.9 \times 10^6$ cycles. By making the change of variables $R_\tau = \exp(-\theta\tau)$, the prior information on θ can be easily converted into prior information on the pump reliability at the time τ :

$$g_A(R_\tau) = \frac{(b/\tau)^a}{\Gamma(a)} [-\ln(R_\tau)]^{a-1} R_\tau^{b/\tau-1} \quad (6)$$

which is a log-gamma distribution with parameters a and b/τ [7]. The prior expectation and standard deviation are

$$E_A\{R_\tau\} = \left(\frac{b}{b+\tau}\right)^a = 0.679 \quad (7)$$

and

$$\sigma_A\{R_\tau\} = \sqrt{\left(\frac{b}{b+2\tau}\right)^a - \left(\frac{b}{b+\tau}\right)^{2a}} = 0.126 \quad (8)$$

respectively. From equation (5), the prior 90% upper credibility limit on θ is given by $\theta_{0.90,A} = \chi_{0.90}^2(2a)/(2b) = 0.969 \times 10^{-6}$ per cycle, where $\chi_{0.90}^2(2a)$ denotes the 90th percentile of the χ^2 distribution with $\nu = 2a$ **degrees of freedom**. The 90% lower credibility limit on R_τ results in $R_{0.10,A} = \exp(-\theta_{0.90,A} \tau) = 0.508$. Thus, on the basis of the prior information, the pumps do not satisfy the reliability requirements and its design should be improved.

However, before making a final decision regarding the reliability attainment, a life testing is performed. Fifteen pumps are then tested under use conditions (see Example 4.10 in reference [1, p. 125]). The test is initially stopped at the time of the fifth failure (failure-truncated sampling) and the observed cycles to failure are $x_1 = 237\,217$, $x_2 = 346\,704$, $x_3 = 430\,057$, $x_4 = 447\,737$, and $x_5 = 507\,746$. The joint conditional distribution of the observed data $\mathbf{x}_1$, given θ , is

$$f(\mathbf{x}_1|\theta) = \frac{n!}{(n-m_1)!} \prod_{i=1}^{m_1} f(x_i) \cdot [R(x_{m_1})]^{n-m_1} \\ \propto \theta^{m_1} \exp(-T_1 \theta) \quad (9)$$

where $m_1 = 5$ is the number of observed failures and $T_1 = \sum_{i=1}^5 x_i + (n - m_1)x_{m_1} = 7\,046\,921$ cycles is the **total time on test**.

Using Bayes theorem (equation 1), the prior information on θ is updated by the sampling information $\mathbf{x}_1$, and the posterior information on θ , given $\mathbf{x}_1$, results in

$$\pi_A(\theta|\mathbf{x}_1) = \frac{(b + T_1)^{a+m_1} \theta^{a+m_1-1}}{\Gamma(a + m_1)} \exp[-(b + T_1)\theta] \quad (10)$$

which is still a gamma distribution with parameters $a + m_1 = 9.0$ and $b + T_1 = 13.94 \times 10^6$ cycles. Then, the gamma density (equation 4) is the conjugate prior for the joint conditional distribution in equation (9), and the prior parameters a and b can be thought of as number of failures and time on test, respectively, in a pseudosample. Hence, the posterior parameters $a + m_1$ and $b + T_1$ can be thought of as the number of combined failures and the combined time on test, respectively. The posterior expectation of θ , i.e. the Bayesian estimator under the squared-error loss function, is

$$E_A\{\theta|\mathbf{x}_1\} = \frac{a + m_1}{b + T_1} = 0.645 \times 10^{-6} \text{ cycles}^{-1} \quad (11)$$

and the posterior 90% upper limit is $\theta_{0.90,A}|\mathbf{x}_1 = 0.932 \times 10^{-6}$ per cycle.

The posterior distribution of R_τ is still a log-gamma distribution:

$$\pi_A(R_\tau|\mathbf{x}_1) = \frac{[(b + T_1)/\tau]^{a+m_1}}{\Gamma(a + m_1)} [-\ln(R_\tau)]^{a+m_1-1} \\ \times R_\tau^{(b+T_1)/\tau-1} \quad (12)$$

with parameters $a + m_1 = 9.0$ and $(b + T_1)/\tau = 19.92$, so that the posterior expectation is $E_A\{R_\tau|\mathbf{x}_1\} = [(b + T_1)/(b + T_1 + \tau)]^{a+m_1} = 0.643$. The posterior 90% lower limit on R_τ results in $R_{0.10,A}|\mathbf{x}_1 = 0.521$.

The sampling information has slightly modified the prior belief on the pump reliability, with a slight improvement of the credibility limits, so that the reliability team decides to continue the life testing until three more pumps fail.

The observed cycles to failure during the second phase of life testing are $x_6 = 575\,877$, $x_7 = 1\,399\,019$, and $x_8 = 1\,502\,329$ [1, p. 125]. The resulting joint conditional distribution of the new data $\mathbf{x}_2$,

4 Life Testing from a Bayesian Perspective

given θ and data $\mathbf{x}_1$ observed during the first testing phase, is

$$f(\mathbf{x}_2|\theta, \mathbf{x}_1) \propto \frac{\prod_{i=m_1+1}^{m_1+m_2} f(x_i) \times [R(x_{m_1+m_2})]^{n-m_1-m_2}}{[R(x_{m_1})]^{n-m_1}} = \theta^{m_2} \exp(-T_2 \theta) \quad (13)$$

where $m_2 = 3$ is the number of failures observed during the second testing phase and $T_2 = \sum_{i=6}^8 (x_i - x_{m_1}) + (n - m_1 - m_2)(x_{m_1+m_2} - x_{m_1}) = 8916068$ cycles is the total time on test.

The posterior distribution $\pi_A(\theta|\mathbf{x}_1)$ of equation (10) becomes the prior distribution of θ for the new data $\mathbf{x}_2$. Then, combining equation (10) with the joint conditional distribution $f(\mathbf{x}_2|\theta, \mathbf{x}_1)$ in equation (13), the posterior distribution of θ , given the whole sampling information $\mathbf{x}_1$ and $\mathbf{x}_2$, say $\pi_A(\theta|\mathbf{x}_1, \mathbf{x}_2)$, is obtained, which is still a gamma distribution with parameters $a + m_1 + m_2 = 12.0$ and $b + T_1 + T_2 = 22.86 \times 10^6$ cycles. Likewise, the posterior distribution $\pi_A(R_\tau|\mathbf{x}_1, \mathbf{x}_2)$ is still a log-gamma distribution with parameters $a + m_1 + m_2 = 12.0$ and $(b + T_1 + T_2)/\tau = 32.66$. Note that the posterior distributions $\pi_A(\theta|\mathbf{x}_1, \mathbf{x}_2)$ and $\pi_A(R_\tau|\mathbf{x}_1, \mathbf{x}_2)$ depend on the sampling data $\mathbf{x}_1$ and $\mathbf{x}_2$ only through the whole number of observed failures, say $m_1 + m_2 = 8$, and the entire total time on test, say $T_1 + T_2 = 15.96 \times 10^6$ cycles.

The posterior expectations of θ and R_τ are $E_A\{\theta|\mathbf{x}_1, \mathbf{x}_2\} = 0.525 \times 10^{-6}$ per cycle and $E_A\{R_\tau|\mathbf{x}_1, \mathbf{x}_2\} = 0.696$, respectively, and the posterior credibility limits are $\theta_{0.90,A}|\mathbf{x}_1, \mathbf{x}_2 = 0.726 \times 10^{-6}$ per cycle and $R_{0.10,A}|\mathbf{x}_1, \mathbf{x}_2 = 0.602$.

Given the prior information on the failure rate and the observed data $\mathbf{x}_1$ and $\mathbf{x}_2$, the upgraded type of pumps satisfies the reliability requirements and mass production can start. Thus, the sampling information has led expert A to change the decision that would have been made before observing data.

Figures 1 and 2 show the updating effect of sampling information on expert belief as life testing goes on. In particular, the standard deviations of θ and R_τ decrease noticeably, changing from $\sigma_A\{\theta\} = 0.290 \times 10^{-6}$ per cycle and $\sigma_A\{R_\tau\} = 0.126$ to $\sigma_A(\theta|\mathbf{x}_1, \mathbf{x}_2) = 0.151 \times 10^{-6}$ per cycle and $\sigma_A(R_\tau|\mathbf{x}_1, \mathbf{x}_2) = 0.072$, respectively.

Now, suppose that another expert, say expert B, is more vague on the value of θ than expert A was.

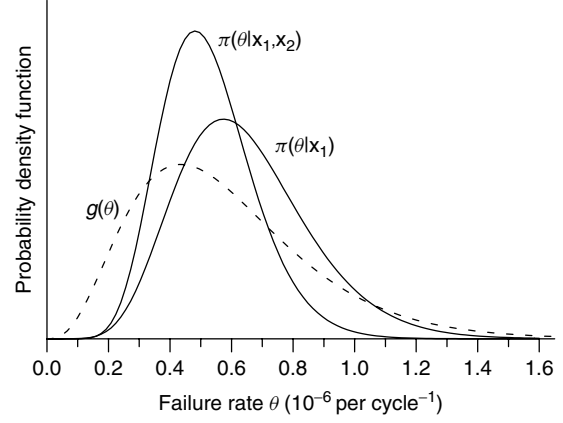


Figure 1 Prior and posterior distributions of the failure rate θ in Example 1

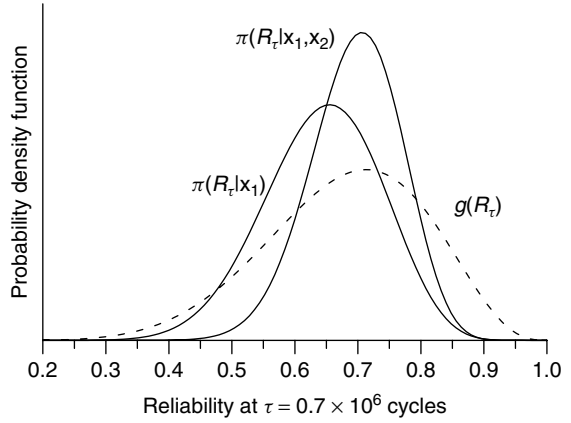


Figure 2 Prior and posterior distributions of the reliability at $\tau = 0.7 \cdot 10^6$ cycles in Example 1

Before performing life testing, expert B elicits the prior belief on θ through a gamma density $g_B(\theta)$ having the same prior expectation elicited by expert A, say $E_B\{\theta\} = E_A\{\theta\} = 0.58 \times 10^{-6}$ per cycle, but larger standard deviation, say $\sigma_B\{\theta\} = 2 \cdot \sigma_A\{\theta\} = 0.58 \times 10^{-6}$ per cycle, so that the parameters of $g_B(\theta)$ are $a' = 1$ and $b' = 1.724 \times 10^6$ cycles. The prior distribution of R_τ , say $g_B(R_\tau)$, is still a log-gamma distribution with $E_B\{R_\tau\} = 0.711$ and $\sigma_B\{R_\tau\} = 0.214$. The sampling information $\mathbf{x}_1$ and $\mathbf{x}_2$ updates the belief on θ and, hence, on R_τ as shown in Figure 3, where the prior and posterior beliefs on R_τ of expert B are compared to the beliefs of expert A. It is evident that the sampling information made the posterior beliefs

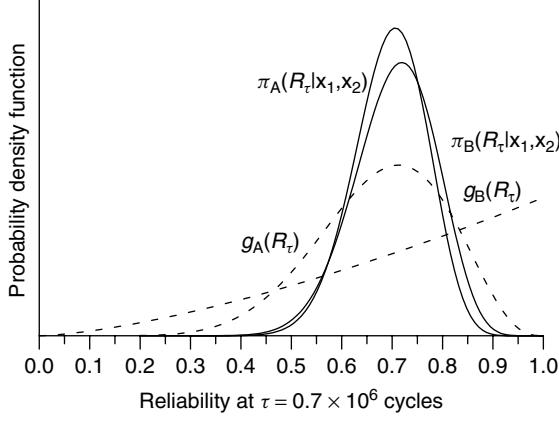


Figure 3 Prior and posterior beliefs of experts A and B on the pump reliability R_t in Example 1

of experts A and B much closer to each other than the prior beliefs were.

Example 2

The reliability of the hydraulic system of a load–haul–dump (LHD) machine has to be estimated to verify if the probability that a new system, undergoing **minimal repair** at failure, experiences more than five failures within the first 1000 h is less than 0.40.

Owing to the minimal repair assumption, the failure process will be described through a **Poisson process** (see **Intensity Functions for Nonhomogeneous Poisson Processes**), so that the probability that a new system will experience k failures in the time interval $(0, T)$ is given by

$$\Pr\{N(T) = k | M_T\} = \frac{M_T^k}{k!} \exp(-M_T) \quad (14)$$

where $M_T \equiv M(T)$ denotes the expected number of failures in $(0, T)$. Hence, the reliability requirement

is met if

$$\Pr\{N(T) > 5 | M_T\} = 1 - \sum_{k=0}^5 \frac{M_T^k}{k!} \exp(-M_T) < 0.40 \quad (15)$$

for $T = 1000$ h.

Previous experience allows the expert to formulate a prior belief just on the value of M_T , in terms of a prior mean $E\{M_T\} = 6$ and a prior standard deviation $\sigma\{M_T\} = 3$. Again, he describes his prior belief on M_T through a gamma density

$$g(M_T) = \frac{b^a}{\Gamma(a)} M_T^{a-1} \exp(-bM_T), \quad a, b > 0 \quad (16)$$

whose parameters are given by $a = (E\{M_T\}/\sigma\{M_T\})^2 = 4$ and $b = E\{M_T\}/\sigma^2\{M_T\} = 2/3$. Then, on the basis of the prior belief, the prior probability $\Pr\{N(T) = k\}$ is

$$\begin{aligned} \Pr\{N(T) = k\} &= \int_0^\infty \Pr\{N(T) = k | M_T\} g(M_T) dM_T \\ &= \frac{\Gamma(a+k)}{\Gamma(a)k!} \frac{b^a}{(b+1)^{a+k}} \end{aligned} \quad (17)$$

which is a Pólya distribution (commonly known as *negative binomial distribution* [8, pp. 122–142]) with expectation $E\{k\} = a/b = 6$. The prior cumulative probability results is $\Pr\{N(T) > 5\} = 0.483$.

A copy of the hydraulic system is put on test, under use conditions, to verify the reliability attainment also on the basis of life testing information. The life testing is stopped when the 26th failure is observed. The times (in hours) to failure (drawn from [9]) are given in Table 1.

Because the truncation time $t_{26} = 3230$ h differs from $T = 1000$ h, the expert has to select the functional form for the **intensity function** $\lambda(t)$ of the Poisson process. On the basis of previous experiences, a nonhomogeneous Poisson process (NHPP), which has the power-law intensity function

$$\lambda(t) = \frac{\beta}{\alpha} \left(\frac{t}{\alpha}\right)^{\beta-1}, \quad \alpha, \beta > 0 \quad (18)$$

Table 1 Times to failures of the hydraulic system of an LHD machine (given in reference [10])

401	437	455	614	955	1126	1150	1500	1572
1875	1909	1954	2278	2280	2350	2407	2510	2521
2526	2529	2673	2753	2806	2890	3108	3230	–

6 Life Testing from a Bayesian Perspective

is chosen, which is often referred to as *power-law process* (PLP) (see **Intensity Functions for Nonhomogeneous Poisson Processes**, [11]). The intensity function (equation 18) increases (decreases) with t when the power-law parameter $\beta > 1$ ($\beta < 1$), and hence the PLP can describe both deterioration and reliability improvement. The expected number of failures $M(t) = \int_0^t \lambda(x) dx = (t/\alpha)^\beta$ is linear with t in a log-log scale.

Then, the joint conditional distribution of $\mathbf{x} = (t_1, t_2, \dots, t_n)$, given the PLP parameters α and β , is

$$f(\mathbf{x}|\alpha, \beta) = \left(\frac{\beta^n}{\alpha^{n\beta}}\right) U^{\beta-1} \exp\left[-\left(\frac{t_n}{\alpha}\right)^\beta\right] \quad (19)$$

where $U = \prod_{i=1}^n t_i$. By changing variables $M_T = (T/\alpha)^\beta$, the joint distribution (19) can be easily reparameterized in terms of M_T and β :

$$f(\mathbf{x}|M_T, \beta) = \beta^n U^{-1} M_T^n \left(\frac{U}{T^n}\right)^\beta \times \exp\left[-\left(\frac{t_n}{T}\right)^\beta M_T\right] \quad (20)$$

and hence a prior distribution for the parameter β has to be elicited (see **Repairable Systems: Bayesian Analysis**). The only belief the expert possesses on β arises from the fact that the hydraulic system is subject to deterioration phenomena with operating time, and hence its failure intensity should increase with t . Thus, within the PLP assumption, β should be greater than 1 and, assuming M_T and β prior independent, the expert formalizes this belief by using a uniform distribution for β over the interval $(1, \beta_U)$ [12]: $g(\beta) = 1/(\beta_U - 1)$ (with $1 \leq \beta \leq \beta_U$). Then, the joint posterior distribution of M_T and β , given the life testing data $\mathbf{x}$, is

$$\pi(M_T, \beta|\mathbf{x}) = \frac{1}{f(\mathbf{x})} \beta^n \left(\frac{U}{T^n}\right)^\beta M_T^{a+n-1} \exp[-BM_T] \quad (21)$$

where $B = b + (t_n/T)^\beta$ and $f(\mathbf{x}) = \Gamma(a+n) \int_1^{\beta_U} \beta^n (U/T^n)^\beta B^{-(a+n)} d\beta$. The marginal posterior distribution of M_T results in

$$\pi(M_T|\mathbf{x}) = \frac{1}{f(\mathbf{x})} M_T^{a+n-1} \int_1^{\beta_U} \beta^n \left(\frac{U}{T^n}\right)^\beta \times \exp[-BM_T] d\beta \quad (22)$$

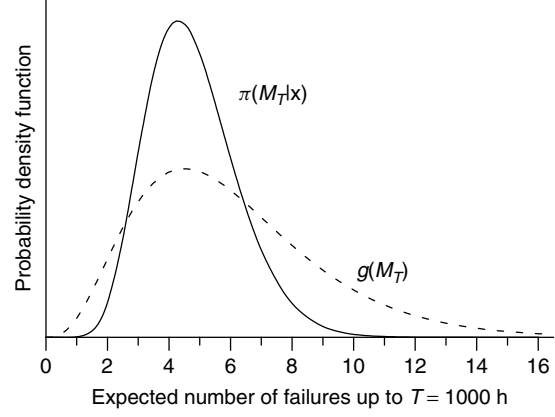


Figure 4 Prior and posterior distributions of the expected number of failures M_T up to $T = 1000$ h in Example 2

and represents the updated belief on M_T , as depicted in Figure 4 that compares the prior and posterior distribution of M_T .

Using equation (22), the posterior probability $\Pr\{N(T) > 5|\mathbf{x}\}$ follows:

$$\begin{aligned} \Pr\{N(T) > 5|\mathbf{x}\} &= \int_0^\infty \Pr\{N(T) > 5|M_T\} \\ &\quad \times \pi(M_T|\mathbf{x}) dM_T \\ &= 1 - \frac{1}{f(\mathbf{x})} \sum_{k=0}^5 \frac{\Gamma(a+n+k)}{k!} \\ &\quad \times \int_1^{\beta_U} \frac{\beta^n (U/T^n)^\beta}{(1+B)^{a+n+k}} d\beta = 0.347 \end{aligned} \quad (23)$$

The reliability requirement is then verified. Figure 5 gives the prior and posterior distribution of $N(T)$ and shows the updating effect of sampling information on prior belief.

To assess whether reasonable changes in the prior information can lead to a different result, a sensitivity analysis is performed with respect to the prior quantities $E\{M_T\}$, $\sigma\{M_T\}$, and β_U . As shown in Table 2, the posterior probability $\Pr\{N(T) > 5|\mathbf{x}\}$ remains less than 0.40 in all the cases but one, where the prior mean $E\{M_T\} = 7.0$. However, even in this worse case, the system reliability is very close to the reliability requirement, and then the reliability requirement is assumed to be met.

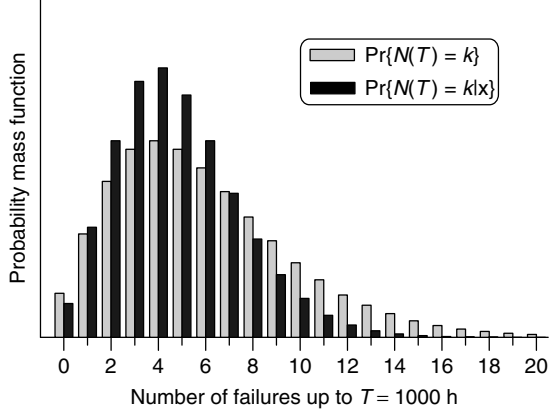


Figure 5 Prior and posterior probabilities of observing k failures until $T = 1000$ h in Example 2

Table 2 Sensitivity analysis *versus* the prior information in Example 2

$E\{M_T\}$	$\sigma\{M_T\}$	β_U	$\Pr\{N(T) > 5 \mathbf{x}\}$
6.0	3.0	5.0	0.347
5.0	3.0	5.0	0.300
7.0	3.0	5.0	0.405
6.0	2.0	5.0	0.399
6.0	4.0	5.0	0.321
6.0	3.0	3.0	0.347
6.0	3.0	7.0	0.347

It is interesting to note that simpler results could be obtained if the test were stopped just at $T = 1000$ h (time-truncated sampling). In this case the joint conditional distribution $f(\mathbf{x}|M_T, \beta)$ can be written as

$$f(\mathbf{x}|M_T, \beta) = \beta^n U^{-1} \left(\frac{U}{T^n} \right)^\beta \times M_T^n \exp[-M_T] \\ = f(\mathbf{x}|\beta) \times f(\mathbf{x}|M_T) \quad (24)$$

and, irrespective of the prior distribution chosen for β , provided β be prior independent of M_T , the posterior distribution of M_T is still a gamma density with parameters $a + n$ and $b + 1$. Thus, the marginal posterior probability $\Pr\{N(T) = k|\mathbf{x}\}$ becomes

$$\Pr\{N(T) = k|\mathbf{x}\} = \frac{\Gamma(a + n + k)}{\Gamma(a + n)k!} \frac{(b + 1)^{a+n}}{(b + 2)^{a+n+k}} \quad (25)$$

and belongs to the Pólya family as the prior $\Pr\{N(T) = k\}$ in equation (17), with expectation $E\{k|\mathbf{x}\} = (a + n)/(b + 1)$.

The Design of Life Testing

Previous examples illustrated the role and the effect of life testing in updating expert belief. Clearly, such an effect depended on the sample size and the manner in which the data have been collected, i.e. on the *design of life testing*. In Example 1, life testing design involved the number of units put on test (namely 15 units) and the truncation time. In general, the design of life testing is a very complex activity, which involves several quantities. For example, in accelerated life testing, given the total number of units to test, the experimenter has to choose the number of levels of the stresses, the values of the stresses, the number of units to test at each stress level, and the amount of censoring at each level before stopping the test (*see Accelerated Life Tests: Bayesian Design*). Because the design of life testing must be chosen before the data (and hence the posterior distribution) can be obtained, life testing design is frequently called *preposterior analysis* in Bayesian framework (*see Bayesian Design of Reliability Experiments*).

The choice of life testing design is commonly made by optimizing the design with respect to a given overall loss function, say $\Lambda(\Phi, \mathbf{d}, \mathbf{x}) = L(\Phi, \mathbf{d}) + C(\mathbf{x})$, which is the sum of a *decision loss* $L(\Phi, \mathbf{d})$, representing the loss associated with the decision $\mathbf{d}$ when Φ is the true state of nature, and a *cost of observation* (or *sampling cost*) $C(\mathbf{x})$, representing the cost of observing the data $\mathbf{x}$. Then, the Bayesian approach to life testing design optimization requires the specification of the decision loss function $L(\Phi, \mathbf{d})$ and the specification of the cost function $C(\mathbf{x})$, whereas the cost of conducting and analyzing the experiment is usually ignored. Note that decision loss and cost functions are in opposition to each other since “to lower the decision loss it will generally be necessary to run a larger experiment, whereby the experimental cost will be increased” [4, p. 281].

Hence, given the data $\mathbf{x}$ observed under a specified design D , the expected overall loss is

$$E_{\Phi|\mathbf{x}, D}\{\Lambda(\Phi, \mathbf{d}, \mathbf{x})\} = \int_{\Phi} \Lambda(\Phi, \mathbf{d}, \mathbf{x}) \pi(\Phi|\mathbf{x}, D) d\Phi \\ = \int_{\Phi} L(\Phi, \mathbf{d}) \pi(\Phi|\mathbf{x}, D) d\Phi \\ + C(\mathbf{x}) \quad (26)$$

where the expectation is made with respect to the posterior distribution $\pi(\Phi|\mathbf{x}, D) \propto f(\mathbf{x}|\Phi, D)g(\Phi)$. Thus, the optimal decision $\mathbf{d}_{\text{opt}}$, given the data $\mathbf{x}$, is the value of $\mathbf{d}$ that minimizes the expected decision loss $E_{\Phi|\mathbf{x}, D}\{L(\Phi, \mathbf{d})\} = \int_{\Phi} L(\Phi, \mathbf{d})\pi(\Phi|\mathbf{x}, D) d\Phi$ over the space of decisions. Then, the expected overall loss of acting optimally, given the data $\mathbf{x}$ under the design D , is $E_{\Phi|\mathbf{x}, D}\{\Lambda(\Phi, \mathbf{d}_{\text{opt}}, \mathbf{x})\} = E_{\Phi|\mathbf{x}, D}\{L(\Phi, \mathbf{d}_{\text{opt}})\} + C(\mathbf{x})$, often referred to as the *minimum posterior risk*. The preposterior expected overall loss, given the design D , also known as the *minimum Bayes risk*, is

$$\begin{aligned} E_{\mathbf{x}|D}\{E_{\Phi|\mathbf{x}, D}\{\Lambda(\Phi, \mathbf{d}_{\text{opt}}, \mathbf{x})\}\} \\ = \int_{\mathbf{x}} E_{\Phi|\mathbf{x}, D}\{\Lambda(\Phi, \mathbf{d}_{\text{opt}}, \mathbf{x})\} f(\mathbf{x}|D) d\mathbf{x} \end{aligned} \quad (27)$$

where $f(\mathbf{x}|D) = \int_{\Phi} f(\mathbf{x}|\Phi, D)g(\Phi) d\Phi$, and provides the overall loss of acting optimally averaged over all possible data $\mathbf{x}$ under the design D . From Bayes theorem (equation 1), the preposterior expected overall loss (equation 27) can also be rewritten as

$$\begin{aligned} E_{\mathbf{x}|D}\{E_{\Phi|\mathbf{x}, D}\{\Lambda(\Phi, \mathbf{d}_{\text{opt}}, \mathbf{x})\}\} \\ = \int_{\mathbf{x}} \int_{\Phi} \Lambda(\Phi, \mathbf{d}_{\text{opt}}, \mathbf{x}) f(\mathbf{x}|\Phi, D)g(\Phi) d\Phi d\mathbf{x} \end{aligned} \quad (28)$$

The **optimal design** is the one that provides the minimum preposterior expected overall loss [4].

Example 3

Again consider Example 1, where life testing design consists of choosing the number m of failures ($m \leq n = 15$) to observe before stopping life testing that minimizes the preposterior expected overall loss. The squared-error decision loss function $L(\theta, d) = c_1 \cdot (\theta - d)^2$ and the linear sampling cost $C(\mathbf{x}) = c_2 \cdot T$ are chosen, where the total time on test $T = \sum_{i=1}^m x_i + (n - m)x_m$ is a sufficient statistic for θ . Note that the (symmetric) squared-error loss is used here for its simplicity although an asymmetric loss function, such as the Linex loss [13], say $L(\theta, d, q) = c_1 \cdot [\exp[q(d - \theta)] - q(d - \theta) - 1]$, or the general entropy loss [14], say $L(\theta, d, p) = c_1 \cdot [(d/\theta)^p - p \ln(d/\theta) - 1]$, is usually more appropriate since an underestimation of the failure rate is

generally much more serious than an overestimation. The optimal decision d_{opt} under the squared-error loss, given the data $\mathbf{x}$ under the design $D = m$, is the posterior expectation $E\{\theta|\mathbf{x}, m\}$ of θ , and the expected decision loss $E_{\theta|\mathbf{x}, D}\{c_1 \cdot (\theta - d_{\text{opt}})^2\}$ of acting optimally is proportional to the posterior variance $V(\theta|\mathbf{x}, m)$. Then, under the gamma prior distribution of equation (5), the minimum expected overall loss $E_{\theta|\mathbf{x}, m}\{c_1 \cdot (\theta - d_{\text{opt}})^2 + c_2 \cdot T\}$ results in

$$E_{\theta|\mathbf{x}, m}\{c_1 \cdot (\theta - d_{\text{opt}})^2 + c_2 \cdot T\} = c_1 \cdot \frac{a + m}{(b + T)^2} + c_2 \cdot T \quad (29)$$

Under failure-censored sampling, the sampling distribution of T , given θ and m , is gamma [10, p. 103]:

$$f(T|\theta, m) = \frac{\theta^m}{\Gamma(m)} T^{m-1} \exp(-\theta T) \quad (30)$$

so

$$\begin{aligned} f(T|m) &= \int_{\theta} f(T|\theta, m)g(\theta) d\theta \\ &= \frac{\Gamma(a + m)}{\Gamma(a)\Gamma(m)} \frac{b^a T^{m-1}}{(b + T)^{a+m}} \end{aligned} \quad (31)$$

The preposterior expected overall loss, given the design m , is

$$\begin{aligned} E_{\mathbf{x}|m}\{E_{\theta|\mathbf{x}, m}\{c_1 \cdot (\theta - d_{\text{opt}})^2 + c_2 \cdot T\}\} \\ = c_1 \cdot \frac{\Gamma(a + m + 1)b^a}{\Gamma(a)\Gamma(m)} \int_0^{\infty} \frac{T^{m-1}}{(b + T)^{a+m+2}} dT \\ + c_2 \cdot \frac{\Gamma(a + m)b^a}{\Gamma(a)\Gamma(m)} \int_0^{\infty} \frac{T^m}{(b + T)^{a+m}} dT \end{aligned} \quad (32)$$

Table 3 gives the preposterior expected overall loss for $m = 2, \dots, n$, when the prior parameters are $a = 4.0$ and $b = 6.9 \times 10^6$ cycles, the loss factor is $c_1 = 10^6 \$ \cdot (10^6 \text{ cycles})^2$, and the cost factor is $c_2 = 1000 \$ / (10^6 \text{ cycles})$ or $c_2 = 2000 \$ / (10^6 \text{ cycles})$. The optimal life testing design is $m = 9$ when $c_2 = 1000 \$ / (10^6 \text{ cycles})$, and is $m = 5$ when $c_2 = 2000 \$ / (10^6 \text{ cycles})$, and the minimum preposterior expected overall loss is equal to \$50 700 and \$65 000, respectively. As could be anticipated, the optimal m value decreases when the sampling cost increases. When the sampling cost is null ($c_2 = 0$), the optimal choice is $m \equiv n = 15$ and the minimum preposterior expected overall loss is \$21 000.

Table 3 Preposterior expected overall loss (in $10^3\$$) for selected m values and cost factors c_2

c_2	Design m													
	2	3	4	5	6	7	8	9	10	11	12	13	14	15
1000 $\$/ (10^6 \text{ cycles})$	63.6	59.4	55.9	53.5	52.0	51.1	50.8	50.7	51.0	51.6	52.3	53.2	54.3	55.5
2000 $\$/ (10^6 \text{ cycles})$	68.5	66.3	65.1	65.0	65.8	67.2	69.1	71.4	74.0	76.9	79.9	83.1	86.5	90.0

References

- [1] Martz, H.F. & Waller, R.A. (1991). *Bayesian Reliability Analysis*, Krieger Publishing, Malabar.
- [2] Bayes, T. (1958). Essay towards solving a problem in the doctrine of chances, *Biometrika* **45**, 298–315.
- [3] Box, G.E.P. & Tiao, G.C. (1973). *Bayesian Inference in Statistical Analysis*, Addison-Wesley, Reading.
- [4] Berger, J.O. (1980). *Statistical Decision Theory. Foundations, Concepts, and Methods*, Springer-Verlag, New York.
- [5] Rios Insua, D., Ruggeri, F. & Martin, J. (2000). Bayesian sensitivity analysis, in *Sensitivity Analysis*, A. Saltelli, K. Chan & M. Scott, eds, JohnWiley & Sons, New York, pp. 225–244.
- [6] Rios, D. & Ruggeri, F. (eds) (2000). *Robust Bayesian Analysis*, Springer-Verlag, New York.
- [7] Consul, P.C. & Jain, G.C. (1971). On the log-gamma distribution and its properties, *Statistische Hefte* **12**, 100–106.
- [8] Johnson, N.L. & Kotz, S. (1969). *Distributions in Statistics. Discrete Distributions*, Houghton-Mifflin, Boston.
- [9] Kumar, U. & Klefsjö, B. (1992). Reliability analysis of hydraulic systems of LHD machines using the power law process model, *Reliability Engineering and System Safety* **35**, 217–224.
- [10] Lawless, J.F. (1982). *Statistical Models & Methods for Lifetime Data*, John Wiley & Sons, New York.
- [11] Crow, L.H. (1974). Reliability analysis for complex repairable systems, in *Reliability and Biometry*, F. Proschan & R.J. Serfling, eds, SIAM, Philadelphia, pp. 379–410.
- [12] Guida, M., Calabria, R. & Pulcini, G. (1989). Bayes inference for a non-homogeneous Poisson process with power intensity law, *IEEE Transactions on Reliability* **38**, 603–609.
- [13] Varian, H.R. A bayesian approach to real estate assessment, in *Studies in Bayesian Econometrics and Statistics*, S.E. Fienberg & A. Zellner, eds, North Holland, Amsterdam, pp. 195–208.
- [14] Calabria, R. & Pulcini, G. (1996). Point estimation under asymmetric loss functions for left-truncated exponential samples, *Communications in Statistics: Theory and Methods* **25**, 585–600.

Related Article

Accelerated Life Tests: Bayesian Design; Accelerated Life Tests: Classical Methods for Design and Analysis; Bayesian Design of Reliability Experiments; Bayesian Reliability Analysis; Intensity Functions for Nonhomogeneous Poisson Processes; Prior Distribution Elicitation; Reliability Sampling; Repairable Systems: Bayesian Analysis

GIANPAOLO PULCINI

Masked Failure Data

Introduction

In many life-testing situations one is only in a position to establish a subset of causes responsible for a given failure, but not the actual cause. This phenomenon is termed *masking*. For example, a typical failure of a personal computer results in error codes that enable one to identify the card (printed circuit board) responsible for the failure. However, in the vast majority of failures the actual offending component on the card remains unidentified. Situations of this type occur because of a number of reasons, the leading one being the cost of failure analysis (FA). Typically, the top priority in the wake of a failure is bringing the system back to operational state as quickly as possible. This is usually achieved through a modular system design and modular diagnostic tests. This approach to system design is the leading cause for the prevalence of masked in manufacturing, field, and warranty areas.

The key statistical problems related to masked data are those of modeling, **estimation**, and diagnostics. Most of the literature has been focused on **competing risk models**, which assume that there is only one cause of failure. However, models in which more than one component could be a culprit have also been considered [1]. As the concept of masked failure appears as a natural extension of the concept of completely identified failure, the whole spectrum of statistical models for lifetime data can be extended to accommodate these types of failures. The estimation and inference problems are typically handled *via* a conventional likelihood approach. One of the key estimation problems in this context relates to survival curves of the individual causes of failure. This problem is of special importance in the area of quality control associated with masked data. For example, one can find that there exists enough evidence to alert the manufacturer of a specific component that the reliability of that component is problematic, even though in most cases the component was not directly implicated in causing the system failure.

Another class of problems is related to estimation of parameters that govern the masking phenomenon itself. For example, in many cases one would want to estimate the probability that a failure of a particular component will lead to a given type of masking.

In situations where one is interested in setting up an efficient repair operation, the inverse problem is also of importance. For example, if we know that replacing a particular set of five components will take care of the failure, could we assert that replacing only components 1 and 3 will eliminate the failure with probability 0.99 or higher? Furthermore, what is the probability that each one of the components 1, 2, ..., 5 was the cause of the failure? The latter probabilities are not only important in the problem of diagnostics, but they also play a key role in quality reporting where it is frequently necessary to “assign the blame” or “allocate the failure” to several causes. Such reporting, in turn, enables one to use, in conjunction with masked data, a common set of reporting and monitoring tools, such as **Pareto** analysis and **control charts** for attribute data.

Model Specification and Parameters

Consider a system that contains k components in series. For simplicity, let us associate these components with causes of failure: we will say that a particular system failed because of component i (or because of cause i). Then one can typically prespecify the *masking groups*: for example, in the case $k = 5$ we could have just three masking groups, $g_1 = (1, 2)$, $g_2 = (1, 3)$ and $g_3 = (3, 4, 5)$. In other words, given that the masking group g_2 is feasible, we could have a situation where, in the wake of failure, we are definitely able to establish that the failure is due to either component 1 or 3, but we are not able to establish which of these components is culpable. In what follows, we will denote by G the set of the possible masking groups and use the symbol g to index these groups. If the cause i is part of the masking group g , we will express it as $i \in g$.

We must note that the problem of analysis of masked data has been studied extensively not only in the engineering and reliability literature, but also in literature related to biostatistics and medical research. Therefore, in what follows we will mention a number of publications from journals related to these areas.

A typical dataset associated with masked failures contains the following information [2]:

N is the number of systems tested.

n_0 is the number of systems that survived.

2 Masked Failure Data

$n_i, i = 1, \dots, k$ is the number of failed systems for which the failure is known to be caused by component i .

n_g is the number of failed systems for which it is known that the cause belongs to a masking group g (given for every $g \in G$).

$t_j^{(i)}, j = 1, \dots, n_i$ are the lifetimes of the n_i failures identified as caused by component i .

$t_j^{(g)}, j = 1, \dots, n_g$ are the lifetimes of the n_g failures identified as caused by one of the components in the masking group g .

$t_j^{(c)}, j = 1, \dots, n_0$ are the censored lifetimes.

The above data must satisfy the conditions $n_0 + \sum_{i=1}^k n_i + \sum_{g \in G} n_g = N$.

A typical model used in conjunction with the above data is defined in terms of the following functions:

$P_i(t), i = 1, \dots, k$, is called the i th *identification probability*. This is the probability that a system failure at time t that is due to cause i is identified as such (i.e., no masking occurs).

$P_{g|i}(t)$ is called the *masking probability*. This is the probability that a system failure at time t that is actually due to cause i will result in a masking group g .

Under the above setting, the above functions must satisfy the condition

$$P_i(t) + \sum_{g \supset i} P_{g|i}(t) = 1, \quad i = 1, \dots, k \quad (1)$$

The lifetime distributions of the individual components (causes of failure) are also of key importance. We will denote the density function and survival function corresponding to the i th component by $f_i(t)$ and $S_i(t)$, respectively ($i = 1, \dots, k$). A number of publications (for example, [3–5]) consider a *nonparametric* approach, in which no further assumptions are made about the nature of the distributions of individual causes of failure. However, in many problems originating in the area of engineering and reliability, there is a basis for assuming a more refined structure. In this case we could represent the density and survival functions in a parameterized form, $f_i(t; \beta_i)$ and $S_i(t; \beta_i)$, where β_i is the vector of parameters corresponding to cause i , and then use a *parametric* approach [6–12].

The primary quantities of interest are the survival curves corresponding to the individual causes, $S_i(t)$, the identification probabilities and the masking probabilities. Some derived parameters are also of interest in many applications. For example, one could be interested in estimating a survival curve of a system of different architecture constructed (wholly or partially) using the k components our system is based on. Problems of this type occur frequently in the process of designing subsequent versions of the system. Another derived quantity is the *diagnostic probability* $\pi_{i|g}(t)$ defined as a probability that a failure at time t that is associated with the masking group g is actually caused by the i th component:

$$\pi_{i|g}(t) = \frac{f_i(t)P_{g|i}(t)}{\sum_{j \subset g} f_j(t)P_{g|j}(t)} \quad (2)$$

Modeling and Analysis of Masked Data

The most popular and productive methods of analysis have been based on the likelihood. The complexity of the likelihood function depends greatly on further assumptions about the nature of failures, and we will only show the form that corresponds to *parametric analysis* under the assumptions that the competing risks are independent. Under this assumption the survival probability $S(t)$ of the system is given by

$$S(t; \beta_1, \dots, \beta_k) = \prod_{i=1}^k S_i(t; \beta_i) \quad (3)$$

and the likelihood function can be written in the form:

$$\begin{aligned} L = & \prod_{j=1}^{n_0} S(t_j^{(c)}; \beta_1, \dots, \beta_k) \\ & \times \prod_{i=1}^k \prod_{j=1}^{n_i} \left\{ f_i(t_j^{(i)}; \beta_i) P_i(t_j^{(i)}) \prod_{m \neq i} S_m(t_j^{(i)}; \beta_m) \right\} \\ & \times \prod_{g \in G} \prod_{j=1}^{n_g} \left[\sum_{r \subset g} f_r(t_j^{(g)}; \beta_r) P_{g|r}(t_j^{(g)}) \prod_{m \neq r} S_m(t_j^{(g)}; \beta_m) \right] \end{aligned} \quad (4)$$

In situations where the identification and masking probabilities are completely specified, the likelihood (equation 4) can be easily analyzed using

standard techniques. In other cases, however, one needs to deal with the parameters that determine these probabilities – and the likelihood analysis becomes rather difficult, because of both numeric issues and issues related to the identifiability and strength of the resulting inference. Therefore, a number of simplifying assumptions have been made in the literature. A number of papers discuss the situation where the identification and masking probabilities are not dependent on time, i.e.,

$$P_i(t) \equiv P_i; \quad P_{g|i}(t) \equiv P_{g|i} \quad (5)$$

(for example, see [12, 13]; however, some more general situations have also been discussed [5, 14]).

Estimability issues

Though the assumptions (equation 5) greatly simplify the likelihood analysis, it could still remain quite involved, even in the parametric case with relatively simple component distributions. One of the reasons for this is that in many applications one needs to take into account *covariates* corresponding to items on test that could reflect, for example, vintage characteristics or environmental conditions – and these usually inject additional parameters into the model [15, 16]. Furthermore, in the case of *proportional hazards*, i.e. when

$$S_i(t) = [S(t)]^{\phi_i}, \quad i = 1, \dots, k \quad (6)$$

with $\phi_i > 0$ and $\sum_{i=1}^k \phi_i = 1$, the problem is not estimable, even in the case where the survival function of the system $S(t)$ is completely identified (see Flehinger *et al.* [2]). For example, in the case $k = 2$ the likelihood function is proportional to $(\phi_1 P_1)^{n_1} \times (\phi_2 P_2)^{n_2} \times [1 - (\phi_1 P_1 + \phi_2 P_2)]^{n_{(1,2)}}$; therefore, only the products $\phi_1 P_1$ and $\phi_2 P_2 = (1 - \phi_1) P_2$ are estimable. Given that in most practical situations the assumption of proportional hazards (possibly involving only a subgroup of competing risks, which would be sufficient to undermine estimability) cannot be *a priori* excluded, the issue of estimability plays a prominent role in the analysis of masked data. In light of this issue, a number of further simplifications have been considered. One of the most popular ones is based on the *symmetry assumption* [17, 18]:

$$P_{g|i} = P_{g|j} \text{ for every } i, j \in g \quad (7)$$

Plausibility of the symmetry assumption typically depends on the reasons for masking. It can be justified, for example, in situations where any failed component in the group causes so much destruction that identification of the culprit is not possible without considerable expense (such situations sometimes occur in crashes of hard drives). This assumption could also be relevant in cases when masking occurs at random; in this case the masking phenomenon is quite similar to random censoring. However, in cases where masking patterns are determined by the FA strategy [19] one could find situations where the symmetry assumption could be *a priori* declared irrelevant. For example, when

1. the diagnostic procedure in the wake of a failure proceeds to examine all potential culprits in a given order (represented by a list),
2. the search for the cause of failure is time-censored, and
3. the inspection time of a given component is considerably shorter in the case when this component is faulty than in the case when it is healthy,

one could expect that for any given masking group (whose members are ordered in accordance with the list), the masking probability for the first component is smaller than that for the other components.

When the symmetry assumption is not justified, the estimability issues frequently necessitate use of more elaborate approaches to the problem of inference. In some cases, one could obtain fairly good estimates of the identification and masking probabilities based on external information. For example, in the situation described above where masking is determined by the postfailure diagnostics policy, one could sometimes obtain reasonable estimates based on the characteristics of this policy and information about the duration (or costs) of inspection under various scenarios. The evaluation can be typically carried out *via* a simulation study of various types of failure.

Two-Stage Models

In many cases one could take advantage of the fact that some of the masked cases can actually be resolved. This leads to *two-stage* analysis [2]. The data for the two-stage model is similar to one given in the section titled “Model Specification and Parameters”, with addition of the information on

how the selected cases were resolved. Use of two-stage models eliminates the estimability issue: one can carry out a likelihood analysis of the expanded dataset in a straightforward way, provided the policy that determines which masked cases are subject to resolution does not depend on the model parameters that we seek to estimate. Research results show that one can obtain a strong boost in statistical power by resolving only a small fraction of cases. Furthermore, one can control the strength of inference *via* a policy that determines which cases are to be resolved.

The two-stage data can be found in a large number of cases involving system manufacturing. The fact is that engineering departments normally do not allow situations generating exclusively masked data to persist for too long: inability to establish in a timely matter that a given component is causing reliability problems could mean exposure to significant financial losses related to warranty and maintenance costs, as large volumes of defective components continue to propagate through the supply chain and the field. However, the databases related to field failure/warranty records and to FA labs are frequently controlled by different organizations, which typically results in problems with compatibility of formats and policies. However, we found that in many cases creation of unified databases suitable for two-stage statistical analysis of masked data could lead to substantial reduction in quality and maintenance costs.

It is also important to note that in some cases, the data for two-stage analysis is potentially available at almost no additional investment and can be obtained *via* better defect documentation procedures. For example, some of the failures that are recorded as nonmasked in the data described in the section titled “Introduction” sometimes represent cases that were originally masked; however, the information about this masking was simply not recorded.

Dependent Causes of Failure

In some cases the basic assumption that the competing hazards are independent cannot be justified. In general, this assumption cannot be overturned on the basis of the observed failure data alone, since the marginal survival probabilities $S_i(t)$ are no longer identifiable when the causes of failure are dependent, even in the case where there is no masking. To see why this is the case, imagine that we have two types

of a two-component system, A and B. For a system of type A the components are independent, with survival curves $S_1(t) = S_2(t) = \sqrt{S(t)}$. For a system of type B the survival curves of both components are equal to $S(t)$, but the lifetimes are dependent in such a way that they are both almost identical, and the probability that component 1 has a shorter life (always by a minuscule amount) than component 2 is 0.5. Under such conditions, distribution of data from lifetime tests related to system A is basically indistinguishable from the similar data for system B.

The assumption of independence can, however, be overturned on the basis of the *a priori* engineering knowledge. For example, it may be known that the lifetime of a system is affected by a “frailty” parameter that affects several components, inducing dependence of the corresponding lifetimes. In the case of depending failures, analysis of masked data can be performed on the basis of *cause-specific hazards* defined by:

$$\lambda_i(t) = \lim_{\Delta t \rightarrow \infty} \left[\frac{\text{Prob}(t \leq T \leq t + \Delta t, C = i | T \geq t)}{\Delta t} \right] \quad i = 1, \dots, k \quad (8)$$

where (T, C) denotes the system lifetime and cause of failure. In this case, instead of representing $S(t)$ in the form of a product (equation 3), we represent $1 - S(t)$ as a sum of *cumulative incidence probabilities* $I_i(t)$ defined *via*

$$I_i(t) = \text{Prob}(T \leq t, C = i) = \int_0^t \lambda_i(u) S(u) du \quad i = 1, \dots, k \quad (9)$$

Now the likelihood function in the case of dependent causes of failure can be written in the form that is very similar to equation (4), but is based on the cause-specific density functions $g_i(t)$ defined by

$$g_i(t) = \frac{d}{dt} I_i(t) = \lambda_i(t) S(t) \quad (10)$$

see [2]. This approach enables one to estimate the parameters related to the cause-specific hazards, identification, masking probabilities, and other quantities of interest. Craiu and Reiser [20] developed an approach based on this technique for estimation related to a model with depending competing risks and piecewise-constant, cause-specific hazards.

Bayesian Approach and Computational Issues

Another method for incorporating external information about the model parameters (or even internal information, presented in the form of so-called non-informative priors) is to use Bayesian techniques. Literature on this approach to the problem of masking is quite extensive [11, 21–24]. In this context, the authors typically use the **Markov chain** (MCMC) techniques for estimation and inference.

In the frequentist setting, the inference can be carried out by maximizing the likelihood (for example, of type shown in equation (4), or its two-stage versions) *via* conventional nonlinear optimization techniques. Confidence statements and test statistics can then be obtained by using the *profile-likelihood* method. However, since the problem is naturally fitting the framework of estimation with missing data, the expectation-maximization (EM) algorithm is also most useful, e.g. see [5, 25, 26].

References

- [1] Reiser, B., Flehinger, B.J. & Conn, A.R. (1996). Estimating component defect probability from masked system success/failure data, *IEEE Transactions in Reliability* **45**, 238–243.
- [2] Flehinger, B.J., Reiser, B. & Yashchin, E. (2001). Statistical analysis of masked data, in *Handbook of Statistics, Advances in Reliability*, N. Balakrishnan & C.R. Rao, eds, Elsevier, Amsterdam, Vol. 20, pp. 499–522.
- [3] Dinse, G.E. (1986). Nonparametric prevalence and mortality estimators for animal experiments with incomplete cause-of-death data, *Journal of American Statistical Association* **81**, 328–336.
- [4] Flehinger, B.J., Reiser, B. & Yashchin, E. (1998). Survival with competing risks and masked causes of failures, *Biometrika* **85**, 151–164.
- [5] Craiu, R.V. & Duchesne, T. (2004). Inference based on EM algorithm for the competing risks model with masked causes of failure, *Biometrika* **91**, 543–558.
- [6] Miyakawa, M. (1984). Analyses of incomplete data in competing risk model, *IEEE Transactions in Reliability* **33**, 293–296.
- [7] Usher, J.S. & Hodgson, T.J. (1988). Maximum likelihood estimation of component reliability using masked system life test data, *IEEE Transactions in Reliability* **37**, 550–555.
- [8] Usher, J.S. & Guess, F.M. (1989). An iterative approach for estimating component reliability from masked system life data, *Quality and Reliability Engineering International* **5**, 257–261.
- [9] Guess, F.M., Usher, J.S. & Hodgson, T.J. (1991). Estimating system and component reliabilities under partial information on the cause of failure, *Journal of Statistical Planning and Inference* **29**, 75–85.
- [10] Usher, J.S. (1996). Weibull component reliability – prediction in the presence of masked data, *IEEE Transactions in Reliability* **45**, 229–232.
- [11] Basu, S., Basu, A. & Mukhopadhyay, C. (1999). Bayesian analysis for masked system failure data using non-identical Weibull models, *Journal of Statistical Planning and Inference* **78**, 255–275.
- [12] Flehinger, B.J., Reiser, B. & Yashchin, E. (2002). Parametric modeling for survival with competing risks and masked failure causes, *Lifetime Data Analysis* **8**, 177–203.
- [13] Lo, S.-H. (1991). Estimating a survival function with incomplete cause-of-death data, *Journal of Multivariate Analysis* **39**, 217–235.
- [14] Lin, D.K.J. & Guess, F.M. (1994). System life data analysis with dependent partial knowledge on the exact cause of system failure, *Microelectronics Reliability* **34**, 535–544.
- [15] Andersen, J.W., Goetghebeur, E. & Ryan, L.M. (1996). Missing cause of death information in the analysis of survival data, *Statistics in Medicine* **15**, 2191–2201.
- [16] Nagai, Y. (2004). Maximum likelihood analysis of masked data in competing risks models with an environmental stress, *IEICE Transactions of Fundamentals of Electronic Communications* **E87A**, 3389–3396.
- [17] Guttman, I., Lin, D.K., Reiser, B. & Usher, J.S. (1995). Dependent masking and system life data analysis: Bayesian inference for two-component systems, *Lifetime Data Analysis* **1**, 87–100.
- [18] Sen, A., Basu, S. & Banerjee, M. (2001). Analysis of masked failure data under competing risks, in *Handbook of Statistics, Advances in Reliability*, N. Balakrishnan & C.R. Rao, eds, Elsevier, Amsterdam, Vol. 20, pp. 523–522.
- [19] Gupta, S.S. & Gastaldi, T. (1996). Life testing for multi-component systems with incomplete information on the cause of failure: a study on some inspection strategies, *Computational Statistics and Data Analysis* **22**, 373–393.
- [20] Craiu, R.V. & Reiser, B. (2006). Inference for the dependent competing risks model with masked causes of failure, *Lifetime Data Analysis* **12**, 21–33.
- [21] Reiser, B., Guttman, I., Lin, D.K.J., Usher, J.S. & Guess, F.M. (1995). Bayesian inference for masked system life time data, *Applied Statistics* **44**, 79–90.
- [22] Lin, D.K.J., Usher, J.S. & Guess, F.M. (1996). Bayes estimation of component reliability from masked system life data, *IEEE Transactions in Reliability* **45**, 233–237.
- [23] Mukhopadhyay, C. & Basu, A.P. (1997). Bayesian analysis of incomplete time and cause of failure data, *Journal of Statistical Planning and Inference* **59**, 79–100.
- [24] Basu, S., Sen, A. & Banerjee, M. (2003). Bayesian analysis of competing risks with partially masked cause-of-failure, *Journal of the Royal Statistical Society. Series C: Applied Statistics* **52**, 77–93.

- [25] Kodell, R.K. & Chen, J.J. (1987). Handling cause of death in equivocal cases using the EM algorithm, *Communications in Statistics A: Theory and Methods* **16**, 2565–2603.
- [26] Zhao, M. & Xie, M. (1994). EM algorithms for estimating software reliability based on masked data, *Microelectronics Reliability* **34**, 1027–1038.

Related Articles

Masked Failure Data: Competing Risks Masked Failure Data: Bayesian Modeling.

EMMANUEL YASHCHIN

Cox Proportional Hazard Model

Introduction

D.R. Cox's article in the *Journal of the Royal Statistical Society* "Regression Models and Life-Tables" (Series B (Methodological), Vol. 34, No. 2 (1972), pp. 187–220) introducing the proportional hazard model is one of the most famous papers in statistics. According to Google Scholar, as of February 2007 it has received over 8500 citations. The volume of literature on the Cox proportional hazard model is wide and deep [1]. The theoretical foundations were secured years later with the development of the counting-process approach (see [2–4]).

What explains the unparalleled success of the Cox model? This contribution attempts to give insight into the reasons for its success and also its limitations. The presentation is quasiformal, ideas are introduced and motivated from sampling considerations rather than from theorems.

The guiding idea is model adequacy. Suppose you want to explain human life span in terms of "covariates" that influence the life span. The first idea might be to build a regression model; you obtain data on how long people live, how much they smoke, their body mass index, their income, years of schooling, sex, race, exposure to pollution, and whatever else you can think of and measure. You regress life span, or maybe log life span on these variables, to be able to argue "For someone in your situation, if you reduce smoking by A you can expect to live B years longer." But first the question of model adequacy must be addressed, how much variance in the life-span data do you explain with your model? If you are told someone's values on all these explanatory variables, how much better can you predict his/her life span than without this knowledge? If the answer is "not very much" then your predictions should not carry much persuasive force.

Or should they? Cox's key idea was to regress, not the life spans themselves, but the *failure* or *hazard rates* onto explanatory variables. However, we do not observe failure rates, just failures. How do we minimize mean square difference between observed and predicted failure rates? We do not; there is no residual in this sense. The beauty of Cox

regression is that it provides another way of finding optimal estimates of the covariates' coefficients. Very roughly, we find values of the **covariate** coefficients such that the people with shorter life spans tend to have larger failure rates. The method is called *maximization of partial likelihood*. It works when we can factorize the failure rate, or its time integral, the cumulative hazard function, into a time-independent part depending on covariates and a time-dependent part independent of covariates. That explains the cohabitation of the terms *proportional hazard model* and *partial likelihood*.

Is it all too good to be true? The goal of this contribution is to enable the reader to judge this for him/herself.^a

Proportional Hazard Model

To simplify the presentation, we consider the case of time-invariant covariates X, Y, Z without censoring and without ties. We consider data to be generated by the following hazard rate

$$h(X, Y, Z) = \Lambda_0(t)e^{XA+YB+ZC} \quad (1)$$

where the cumulative baseline hazard function $\Lambda_0(t) = \int_0^t \lambda_0(u) du$ is the time integral of the baseline hazard rate λ_0 . The covariates (X, Y, Z) are considered as random variables. The coefficients (A, B, C) and the baseline hazard Λ_0 will be estimated from life data. If this hazard function holds, then for an individual with covariate values (x, y, z) the survivor function is

$$e^{-h(x,y,z)} \quad (2)$$

Suppose we observe times of death $t_1, \dots, t_n$ such that $t_i < t_j$ for $i < j$. Let the covariates for the individual dying at time t_i be denoted (x_i, y_i, z_i) . The coefficients A, B, C are estimated by maximizing the *partial likelihood*

$$\prod_{i=1}^N \frac{e^{x_i A + y_i B + z_i C}}{\sum_{j \geq i} e^{x_j A + y_j B + z_j C}} \quad (3)$$

Note that the times of death t_i do not appear in equation (3). The intuitive explanation is as follows. *Given* that the first death in the population occurs at time t_1 , the probability that it happens

2 Cox Proportional Hazard Model

to individual 1 is $(e^{x_1 A + y_1 B + z_1 C} / \sum_{j \geq 1}^n e^{x_j A + y_j B + z_j C})$. After individual 1 is removed from the population, the same reasoning applies to the surviving population; *given* that the second time of death is t_2 , the probability that it happens to individual 2 is $(e^{x_2 A + y_2 B + z_2 C} / \sum_{j \geq 2}^n e^{x_j A + y_j B + z_j C})$, and so on. Kalbfleisch and Prentice [5] note that for constant covariates, equation (3) is the likelihood for the *ordering* of times of death. The baseline hazard can be estimated from the data as described in [5, p. 114].

Model Adequacy

With (x, y, z) fixed and T random, *and* with constant baseline hazard rate scaled to one, the survivor function (2), considered as a function of the random variable T is uniformly distributed on $[0, 1]$, that is

$$T \sim -\frac{\ln(U)}{h} \quad (4)$$

where U is uniform on $[0, 1]$. As this holds for each individual in the population $i = 1 \dots N$. If we order the values

$$e^{-t_i e^{x_i A + y_i B + z_i C}}; \quad i = 1 \dots N \quad (5)$$

and plot them against their number, the points should lie along the diagonal if the proportional hazard model is true with coefficients A, B, C and constant baseline hazard rate. The values in equation (5) are the exponentials of the Cox–Snell **residuals**; equal up to a constant to the martingale residual, used in the counting-process approach. The Cox–Snell residuals are exponentially distributed if the model is correct.

This would provide an easy heuristic check of model adequacy if the baseline hazard were indeed known to be constant and scaled to one. However, if the baseline hazard is also estimated from the data, then this simple test does not apply.

Testing model adequacy for the Cox model is not straightforward: This is a sampling of statements found in the literature regarding model evaluation: “... it is not apparent what kinds of departures one would expect to see in the residuals if the model is incorrect, or even to what extent agreement with the anticipated line should be expected” [5, p. 128]. “For most purposes, you can ignore the Cox–Snell and martingale residuals. While Cox–Snell residuals were useful for assessing the fit of the parametric

models in Chapter 4, they are not very informative for the Cox models estimated by partial likelihood” [6, p. 173]. “Unfortunately, this distribution theory [of the Cox–Snell residuals as exponentially distributed] has not proven to be as useful for model evaluation as the theory derived from the counting process approach”. [7 p. 202], “... there is not a single, simple, easy to calculate, useful, easy to interpret measure [of model performance] for a proportional hazards model” [7, p. 229]. “... the martingale residuals cannot play all the roles that linear model residuals do; in particular the overall distribution of the residuals does not aid in the global assessment of fit” [8, p. 81]. In many important studies, model adequacy is not examined, and only individual coefficients for the covariate of interest are reported, with Wald confidence bounds (e.g. [9, 10]). The coefficients are used to compute relative risk, and form the basis of (dis)utility calculations for different risk mitigation measures.

It may well arise that data generated with a constant baseline hazard appear to acquire a time-dependent baseline hazard as a result of missing covariates. Letting $\hat{\diamond}$ denote values estimated from the data, we may well find that the values

$$e^{-\hat{\Lambda}_0(t_i) e^{x_i \hat{A} + y_i \hat{B} + z_i \hat{C}}}; \quad i = 1 \dots N \quad (6)$$

plot as uniform, while the estimates do not resemble the values which generated the data. In particular, this may arise in the case of missing covariates. We identify some covariates but many others may not be represented in our model. For example, in considering the influence of airborne fine particulate matter on nonaccidental mortality [9, 10], covariates like smoking, sex, age, socioeconomic status, air quality, and weather are studied. However, time to death is obviously influenced by myriad other factors like occupation, genetic disposition, stress, disease prevalence, medical care, diet, alcohol consumption, home environment (e.g. radon), travel patterns, and so on.

The following type of simple numerical experiment, which the reader may verify for him/herself will illustrate the problems with model adequacy.^b

1. Choose coefficients (A, B, C) , choose a constant baseline hazard rate scaled to one, and choose a distribution for (X, Y, Z) .

2. Sample independently 100 values of (X, Y, Z) and 100 values from the uniform distribution on $[0, 1]$; compute failure times using equation (4).
3. Estimate the coefficients by maximizing equation (3), and estimate the baseline hazard.

This procedure does not require that the distributions of the covariates be centered at their means; indeed, centering is not a standard procedure in applications. However, the uniform distribution on $[-1, 1]$ used here is centered.

Let the model (1) be termed h_{XYZ} . To study the effects of model incompleteness estimate the coefficient A with a model h_{XY} using only covariates X and Y , and with a model h_X using only covariate X . For each of the models h_{XYZ} , h_{XY} and h_X , we repeat the above procedure 100 times with the same values for (A, B, C) , with (X, Y, Z) sampled independently from the (centered) uniform distribution on $[-1, 1]$; however, we change equation (3) so as to estimate A in the models h_{XY} and h_X . Figure 1 plots the ordered estimates of coefficient A .

Evidently, the models h_{XY} and h_X tend to underestimate the coefficient A . A theoretical explanation of this underestimation is given in [11, 12]. Hougaard [13] shows that missing covariates sometimes result in nonproportional hazard functions, which could be detected with sufficient replications. He concludes that “the regression coefficients cannot be evaluated without consideration of the variability of neglected covariates. This is a serious drawback of relative

risk models” [13, p. 92], leading him to promote accelerated failure or frailty models. The tendency to underestimate becomes more pronounced in Figure 2, where the missing covariate Z has coefficient $C = 5$. In spite of this, the ordered values of equation (6) plot along the diagonal, as shown in Figure 3. If we knew that the data was created with $\lambda_0(t) \equiv 1$, then we may impose this constraint on the survivor functions. From Figure 4 we see that uniformity is lost for the incomplete models h_{YXY} , h_X ; but not for the complete model h_{XYZ} . This would provide an excellent diagnostic for completeness if we had *a priori* knowledge of the baseline hazard rate; unfortunately in practice we do not have this knowledge. We can, however, find another diagnostic (see below).

Figure 5 shows the Wald 95% confidence bounds for A in model h_X , in each of the 100 repetitions of the experiment whose estimates are shown in Figure 2. These bounds are derived assuming asymptotic normality of the Wald statistic

$$\frac{\hat{A} - A}{\sigma_A} \quad (7)$$

where $\hat{A}$ is the estimate of A and σ_A is derived from the observed information matrix. If the likelihood function is correct, then the Wald statistic is asymptotically standard normal. In as much as these 95% confidence bands contain the true value $A = 1$ in only 7% of the cases, the wisdom of stating such

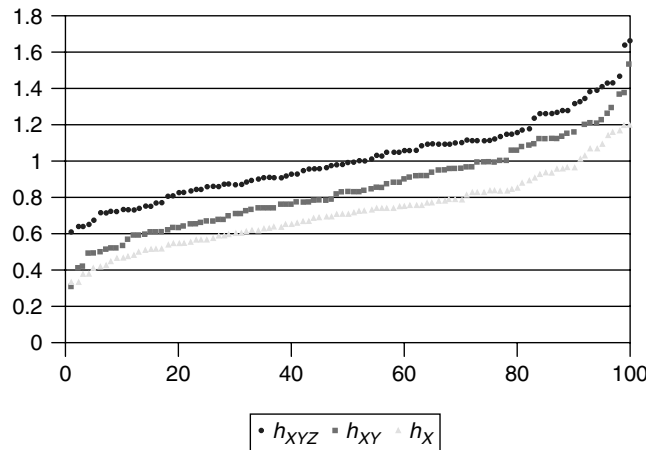


Figure 1 100 ordered estimates of A for h_{XYZ} , h_{XY} , h_X , $X, Y, Z \sim U[-1, 1]$, $(A, B, C) = (1, 1, 1)$; each estimate based on 100 samples

4 Cox Proportional Hazard Model

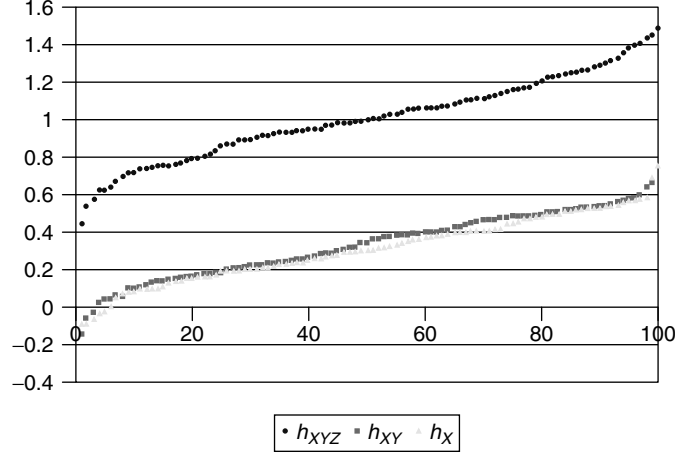


Figure 2 100 ordered estimates of A for h_{XYZ}, h_{XY}, h_X , $X, Y, Z \sim U[-1, 1]$, $(A, B, C) = (1, 1, 5)$ each estimate based on 100 samples

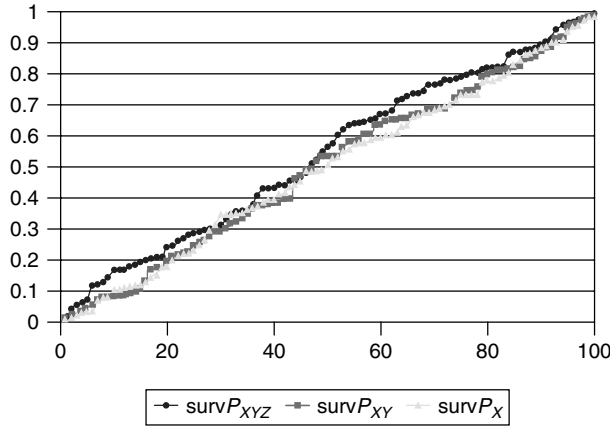


Figure 3 Ordered values of equation (6) for h_{XYZ}, h_{XY}, h_X , $X, Y, Z \sim U[-1, 1]$, $(A, B, C) = (1, 1, 5)$

confidence bounds when model adequacy cannot be demonstrated may be questioned.

The models h_{XY} and h_X are clearly incorrect and misestimate the covariate A . Relative risk coefficients based on these models would be biased. Without *a priori* knowledge of the baseline hazard function, their incorrectness cannot be diagnosed using Cox–Snell or martingale residuals, echoing the statements cited at the beginning of this section. The problem is that the **lack of fit** caused by missing covariates is compensated in the estimated baseline hazard function.

This observation suggests that we might detect

lack of fit in the covariates by comparing the estimated baseline hazard function with the population cumulative hazard function, that is, minus the natural logarithm of the population survivor function. From equation (6) it is evident that adding a constant to any covariate is equivalent to multiplying the baseline hazard by a constant. We therefore standardize the covariates by centering their distributions on the means (the distributions here already centered). This standardization can be motivated as follows. If the covariates X , Y , and Z have no effect, then their means in the set R_i of individuals at risk at the time t_i of death of the i th individual, are just their means

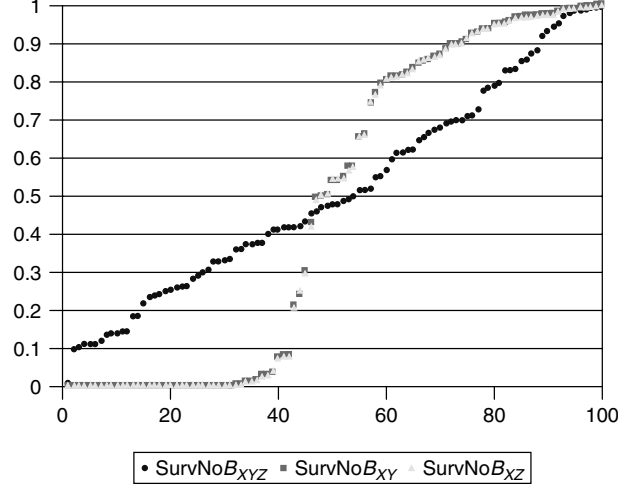


Figure 4 Ordered values of equation (6) for $h_{XYZ}, h_{XY}, h_X, X, Y, Z \sim U[-1, 1], (A, B, C) = (1, 1, 5)$ with $\hat{\Lambda}_0(t) \equiv t$

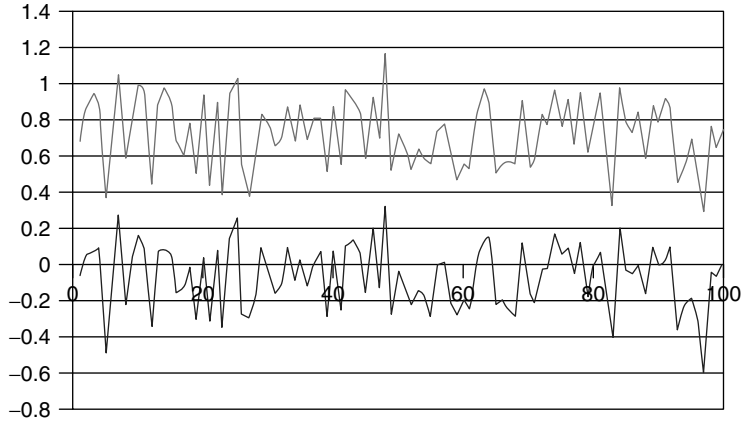


Figure 5 Wald 95% confidence bounds for A with model, h_X of Figure 2; each estimate based on 100 samples

$\bar{x}, \bar{y}, \bar{z}$. We write

$$\sum_{j \in R_i} e^{x_j A + y_j B + z_j C} \sim \sum_{j \in R_i} (1 + x_j A)(1 + y_j B)(1 + z_j C) \quad (8)$$

If X, Y , and Z are independent with mean zero then the mean of the right-hand side is just $\#R_i$. The

Breslow estimate [14] of the baseline hazard rate is

$$d\Lambda_0(t_i) = \left[\sum_{j \in R_i} e^{x_j A + y_j B + z_j C} \right]^{-1} \quad (9)$$

and the population cumulative hazard rate at t_i is $[\#R_i]^{-1}$. Hence, under these conditions, centering the covariates at their means, then replacing covariates with their mean value (zero), would equate the cumulative baseline hazard function and the population cumulative hazard function.

6 Cox Proportional Hazard Model

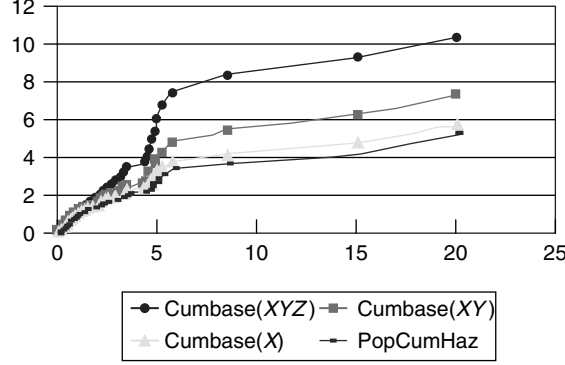


Figure 6 Cumulative population and baseline hazard functions for $h_{XYZ}, h_{XY}, h_X, X, Y, Z \sim U[-1, 1]$, $(A, B, C) = (1, 1, 1)$

Figures 6 and 7 show these comparisons for the two cases from Figures 1 and 2. Note the difference in survival times (horizontal axis); this is caused by the heavier loading of covariate Z in Figure 7. The Nelson–Aalen estimator is used for the population cumulative hazard function.

We see in Figure 7 that the cumulative baseline hazard functions for h_{XY} and h_X have moved closer to the population cumulative hazard, reflecting the heavier loading on the missing covariate Z .

If a Cox model had *none* of the actual covariates, this would be equivalent to having zero coefficients on all covariates; and in this case the baseline hazard would coincide with the population cumulative hazard. A simple heuristic test of model adequacy would test the null hypothesis that the cumulative baseline

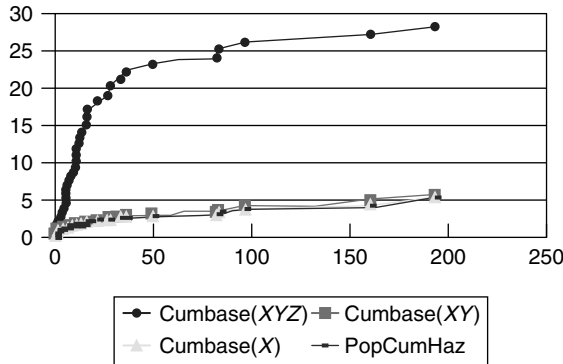


Figure 7 Cumulative population and baseline hazard functions for $h_{XYZ}, h_{XY}, h_X, X, Y, Z \sim U[-1, 1]$, $(A, B, C) = (1, 1, 5)$

hazard function is equal to the population cumulative hazard function. If the null hypothesis cannot be rejected, then using the Cox model would not be indicated. In Figures 8 and 9 the asymptotic 2σ bands on the asymptotic variance of the Nelson–Aalen estimator of the population cumulative hazard function [5, p. 25] have been added to Figures 6 and 7. We see that with this simple test we would fail to reject the null hypothesis for model h_X after 100 observations in both cases. The greater loading of the missing covariate Z in Figure 9 causes the model h_{XY} to fail to reject the null hypothesis as well.

The more familiar partial-likelihood-ratio test calculates the test statistic G as twice the difference between the log partial likelihood of the model containing the covariates and the log partial likelihood for the model not containing the covariates. G is asymptotically χ^2 distributed under the null hypothesis. The above test may have some advantage in that it does not appeal to partial likelihood. However, it is unable to detect the lack of fit in the model h_{XY} when $C = 1$.

We note that for all the results mentioned above, the covariates are independent. In practice independence is not usually checked, and not always plausible. Figure 10 shows 100 estimates of the coefficient A for the models h_{XYZ}, h_{XY}, h_X where the covariates are uniformly distributed on $[0, 1]$ with correlations $\rho(X, Z) = 0.98, \rho(Y, Z) = 0.41$. Whereas missing covariates produce underestimation in the case of independence, we see that dependence in Figure 10 produces overestimation. Note also that the spread of estimates for the complete model h_{XYZ} is very wide.

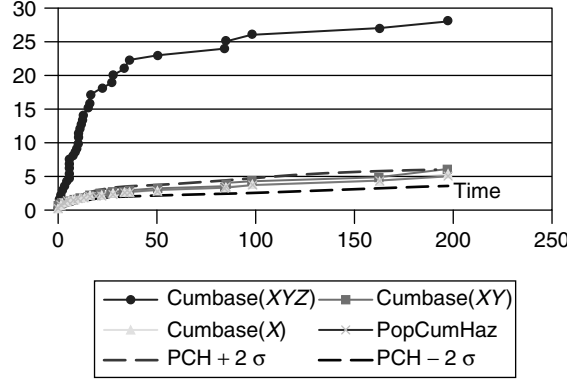


Figure 8 Cumulative population and baseline hazard functions for $h_{XYZ}, h_{YXY}, h_X, X, Y, Z \sim U[-1, 1]$, $(A, B, C) = (1, 1, 5)$ with 2σ confidence bands (dashed lines)

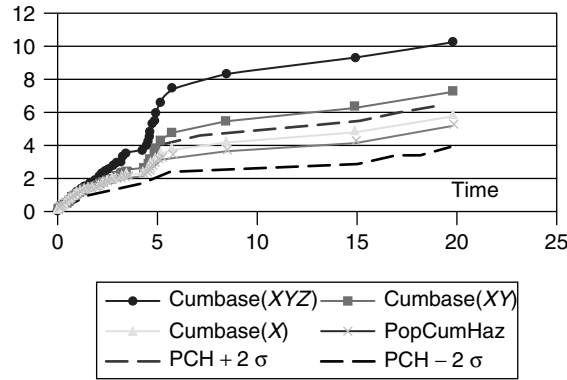


Figure 9 Cumulative population and baseline hazard functions for $h_{XYZ}, h_{YXY}, h_X, X, Y, Z \sim U[-1, 1]$, $(A, B, C) = (1, 1, 1)$ with 2σ confidence bands (dashed lines)

Censoring and Competing Risk

The discussion of model adequacy with the proportional hazard model is sometimes clouded by the role of censoring. The following statement is representative: “A perfectly adequate model may have what, at face value, seems like a terribly low R^2 due to a high percent of censored data” [7, p. 229]. The reference to R^2 must be taken as metaphorical. The proportional hazard model proposes a linear regression of the log hazard function. The hazard function is not observed, and hence a measure of the difference between observed and predicted values, like R^2 is not meaningful. The point is that the ability of a proportional hazard model to “explain the data” might be obscured by censoring.

Right censoring of course is a form of **competing risk**. In this section we review some recent results from the theory of competing risk, and indicate how they may yield diagnostic tools in proportional hazard modeling. In the competing-risks approach we model the data as a sequence of independent and identically distributed (i.i.d.) pairs (T_i, δ_i) , $i = 1, 2, \dots$. Each T is the minimum of two or more variables, corresponding to the competing risks. We will assume that there are two competing risks, described by two random variables D and C such that $T = \min(D, C)$. D will be time of death, which is of primary interest, while C is a censoring time corresponding to termination of observation by other causes. In addition to the time T one observes the indicator variable $\delta = I(D < C)$,

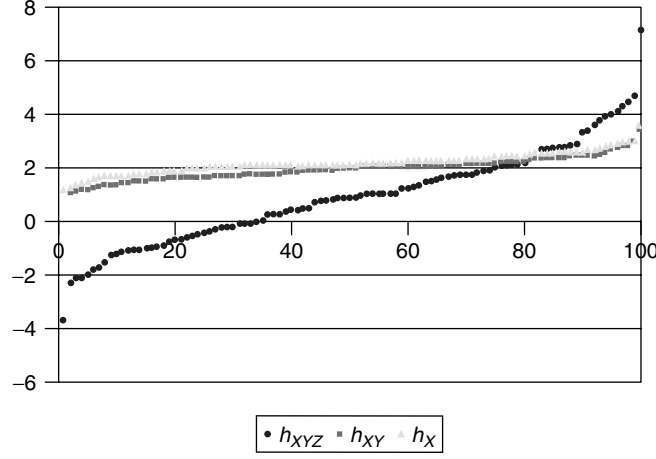


Figure 10 Hundred ordered estimates of A for $h_{XYZ}, h_{XY}, h_X, X, Y, Z \sim U[0, 1]$, $(A, B, C) = (1, 1, 1)$ with dependent covariates

which describes the cause of the termination of observation. For simplicity we assume that $P(D = C) = 0$.

It is well known (Tsiatis [15]) that from observation of (T, δ) we can identify only the subsurvivor functions of D and C :

$$S_D^*(t) = P(D > t, D < C) = P(T > t, \delta = 1) \quad (10)$$

$$S_C^*(t) = P(C > t, C < D) = P(T > t, \delta = 0) \quad (11)$$

but not in general the true survivor functions of D and C , $S_D(t)$ and $S_C(t)$. Note that $S_D^*(t)$ depends on C , though this fact is suppressed in the notation. Note also that $S_D^*(0) = P(D < C) = P(\delta = 1)$ and $S_C^*(0) = P(C < D) = P(\delta = 0)$, so that $S_D^*(0) + S_C^*(0) = 1$.

The conditional subsurvivor functions are defined as the survivor functions conditioned on the occurrence of the corresponding type of event. Assuming continuity of $S_D^*(t)$ and $S_C^*(t)$ at zero, these functions are given by

$$\begin{aligned} CS_D^*(t) &= P(D > t | D < C) \\ &= P(T > t | \delta = 1) = S_D^*(t) / S_D^*(0) \end{aligned} \quad (12)$$

$$\begin{aligned} CS_C^*(t) &= P(C > t | C < D) \\ &= P(T > t | \delta = 0) = S_C^*(t) / S_C^*(0) \end{aligned} \quad (13)$$

Closely related to the notion of subsurvivor functions is the probability of censoring beyond time t ,

$$\begin{aligned} \Phi(t) &= P(C < D | T > t) \\ &= P(\delta = 0 | T > t) = \frac{S_C^*(t)}{S_D^*(t) + S_C^*(t)} \end{aligned} \quad (14)$$

This function has some diagnostic value, aiding us to choose among competing-risk models to fit the data. Note that $\Phi(0) = P(\delta = 0) = S_C^*(0)$.

As mentioned above, without any additional assumptions on the joint distribution of D and C , it is impossible to identify the marginal survivor functions $S_D(t)$ and $S_C(t)$. However, by making extra assumptions, one may restrict to a class of models in which the survivor functions are identifiable. A classical result on competing risks [15, 16] states that, assuming independence of D and C , we can determine uniquely the survivor functions of D and C from the joint distribution of (T, δ) , where at most one of the survivor functions has an atom at infinity. In this case the survivor functions of D and C are said to be identifiable from the censored data (T, δ) . Hence, an independent model is always consistent with data.

If the censoring is assumed to be independent then the survivor function for T , the minimum of D and C , can be written as

$$S_T(t) = S_D(t)S_C(t) \quad (15)$$

If we assume that D obeys a proportional hazard model, and that the censoring is independent, then we may estimate the coefficients by maximizing the partial-likelihood function adapted to account for censoring:

$$\prod_{i \in D_N} \frac{e^{x_i A + y_i B + z_i C}}{\sum_{j \geq i}^n e^{x_j A + y_j B + z_j C}} \quad (16)$$

where D_N is the subset of observed times $t_1, \dots, t_N$ at which death is observed to occur, and j runs over all times corresponding to death or censoring.

If we now substitute the survivor function with estimated coefficients in equation (15), and use the familiar Kaplan–Meier estimator for S_C , then we may apply the ideas of the previous section to assess model adequacy [17].

Conclusions

The user of the Cox proportional hazard model may be assured that *if* the model is correct, then the coefficient estimates will converge to the true values, and the confidence bands will accurately reflect sampling fluctuations around these true values. If the model is *not* correct, however, then the coefficients' estimates may not be correct, and the confidence bands may give a misleading picture. If the covariates are independent, and if one or more covariates is omitted, then the coefficients of the other covariates will be underestimated in absolute value. If there is dependence between missing and included covariates, then nothing can be said about the direction of the error.

How can we check whether the model is correct? This remains the difficult point. We can test whether covariates are significantly different from zero. However, when covariates pass this test, it does not mean that they will be estimated correctly.

End Notes

^a The author gratefully acknowledges many helpful comments from Bo Lindqvist.

^b The following simulations were performed with EXCEL and checked with S+.

References

- [1] Oakes, D. (2001). Biometrika centenary: survival analysis, *Biometrika* **88**, 99–142.
- [2] Andersen, P.K. & Gill, R.D. (1992). Cox's regression model for counting processes: a large sample study, *The Annals of Statistics* **10**(4), 1100–1120.
- [3] Andersen, P.K., Borgan, O., Gill, R.D. & Keiding, N. (1992). *Statistical Models Based on Counting Processes*, Springer-Verlag, New York.
- [4] Fleming, T.R. & Harrington, D.P. (1991). *Counting Processes and Survival Analysis*, John Wiley & Sons, New York.
- [5] Kalbfleisch, J.D. & Prentice, R.L. (2002). *The Statistical Analysis of Failure Time Data*, 2nd Edition, John Wiley & Sons, New York.
- [6] Allison, P.D. (2003). *Survival Analysis using SAS a Practical Guide*, SAS Publishing, Cary.
- [7] Hosmer, D.W. & Lemeshow, S. (1999). *Applied Survival Analysis*, John Wiley & Sons, New York.
- [8] Therneau, T.M. & Grambsch, P.M. (2000). *Modeling Survival Data*, Springer, New York.
- [9] Dockery III, D., Pope, C.A.P., Xu, X., Spengler, J.D., Ware, J.H., Fay, M.E., Ferris, B.G. & Speizer, F.E. (1993). An association between air pollution and mortality in six U.S. cities. *New England Journal of Medicine* **329**, 1753–1759.
- [10] Pope, C.A., Thun, M.J., Namboodiri, M.M., Dockery, D., Evans, J.S., Speizer, F.E. & Heath, C.W. (1995). Particulate air pollution as a predictor of mortality in a prospective study of U.S. adults, *American Journal of Respiratory and Critical Care Medicine* **151**, 669–674.
- [11] Bretagnolle, J. & Huber-Carol, C. (1988). Effects of omitting covariates in Cox's Model for survival data, *Scandinavian Journal of Statistical* **15**, 125–138.
- [12] Keiding, N., Andersen, P.K. & Klein, J.P. (1997). The role of frailty time models in describing heterogeneity due to omitted covariates, *Statistics in Medicine* **16**, 215–224.
- [13] Hougaard, P. (2000). *Analysis of Multivariate Survival Data*, Springer-Verlag, New York.
- [14] Breslow, N. (1972). Discussion on professor Cox's paper, *Journal of the Royal Statistical Society. Series B* **34**, 216–217.
- [15] Tsiatis, A. (1975). A nonidentifiability aspect of the problem of competing risks, *Proceedings of the National Academy of Sciences of the United States of America* **72**, 20–22.
- [16] van der Weide, J.A.M. & Bedford, T. (1998). Competing risks and eternal life, in *Safety and Reliability (Proceedings of ESREL'98)*, S. Lydersen, G.K. Hansen & H.A. Sandtorv, eds, Balkema, Rotterdam, Vol. 2, pp. 1359–1364.
- [17] Cooke, R.M. & Morales, O. (2006). Competing risk and the Cox proportional hazard model, *Journal of Statistical Planning and Inference* **136**(5), 1621–1637.

Further Reading

- Bunea, C., Cooke, R.M. & Lindqvist, B. (2002). Maintenance study for components under competing risks, *Safety and Reliability*, 1st volume, (European Safety and Reliability Conference, ESREL 2002, Lyon, 18–21 March 2002), pp. 212–217.
- Bunea, C., Cooke, R.M. & Lindqvist, B. (2003). Competing risk perspective over reliability data bases, in *Mathematical and Statistical Methods in Reliability*, B.H. Lindqvist & K.A. Doksum, eds, World Scientific Publishing, Singapore, pp. 355–370.
- Cooke, R.M. (1996). The design of reliability databases Part I and II, *Reliability Engineering and System Safety* **51**, 137–146; 209–223.
- Cooke, R.M. & Bedford, T.J. (2002). Reliability databases in perspective, *IEEE Transactions on Reliability* **51**(3), 294–310.
- Cox, D.R. (1972). Regression models and life-tables, *Royal Statistical Society. Series B* **34**(2), 187–220.

ROGER M. COOKE

Process Reengineering

Background and Definition

Business process reengineering (BPR) is the *fundamental* rethinking and *radical* redesign of *business processes* [1, 2] – that is, the flow of work and information within and between organizations. Its purpose is to achieve *drastic* improvement over existing performance in quality, cost, service, and cycle time. It is also known as business process *redesign*, business process *transformation*, or business process *innovation*.

Rather than pursuing continuous, but incremental improvement like some of its predecessor methodologies such as **Kaizen**, and total quality management (TQM), BPR seeks radical improvement in a single effort.

BPR reached the peak of its popularity in the mid-to-late 1990s after Michael Hammer and James Champy published their best-selling book *Reengineering the Corporation* in 1993.

Use of the Methodology

BPR was initially intended to significantly improve, in current terms, the business performance and competitiveness of the enterprise. The extensive reengineering of administrative, information systems, procurement, and manufacturing processes in the early 1990s is the well-known result of this approach. More recently, management has discovered that business performance can further be improved by the systematic reengineering of its *core* processes: the aggregation of work that creates customer value.

Accordingly, the selection of business processes for reengineering predominantly targets those core processes that directly affect the customers of the firm: manufacturing, distribution, customer service, invoicing, and order entry.

Process reengineering is not intended to replace or substitute for *automation*, *software reengineering*, or traditional *reorganization*. Its original proponents (Hammer [2] and Davenport [3]) and other lead practitioners also emphasized that reengineering should not lead to *downsizing* – its basic concept is to “do more with the same”.

Although it is difficult to offer general guidelines on the use of BPR, certain conditions are conducive to its application:

- There are serious quality and customer satisfaction concerns.
- Continuous improvement methods are not producing satisfactory results.
- Competition is consistently outperforming the enterprise.
- The organization exhibits considerable internal organizational conflict.
- The company is experiencing excessive product and administrative cost overruns.

Most of the companies with one or more of these conditions have clearly differentiated functions: vertical organizational “silos” – structures built around the concept of the division of labor and hierarchical management. These functional silos comprise only narrow pieces of a multitude of processes that in total make up the workings of the enterprise. Managers and executives are in charge of increasingly large segments of the *function*, but no one is in charge of the *processes*. This is known as *process fragmentation*. Because of process fragmentation, handoff problems between functions proliferate throughout the organization.

As a result, processes become highly error prone as each handoff requires coordination and communication and entails wait times and mismatches in schedules. And while these processes have always existed – they are, in essence, the way work is done by the enterprise – only a few people are aware of them, and they are not managed in any consistent or rational way.

Traditional BPR is about the reintegration of work and the elimination of process fragmentation and the related handoff problem.

Basic Concepts and Methods

In many cases, radical redesigning and reorganization of an enterprise is necessary to lower costs and improve the quality of products and services. Information technology (IT) is the key enabler of such radical change. The nine principles of successful reengineering to achieve quantum performance improvement include:

2 Process Reengineering

1. Develop the business vision and process goals.
2. Identify all the processes in an organization and prioritize them in order of redesign need (based on the importance of the process, its level of dysfunctionality, and the likelihood of success).
3. Understand and measure the performance of existing processes.
4. Organize around results and customers, not tasks or departments.
5. Place decision points where the work is performed by building appropriate controls into the process.
6. Treat geographically dispersed resources as if they were concentrated.
7. Capture information only once and do this at its source.
8. Identify IT resources and capabilities; integrate IT into the material workflow for maximum leverage.
9. Design and build a prototype of the new process as the base for successive design iterations and tests.

IT plays a critical role in implementing the BPR concept. It is viewed as an important enabler of new ways of communicating and collaborating within an organization and across organizational boundaries.

Generally, the application of IT to process reengineering requires *inductive* thinking in that its power of disrupting existing rules and paradigms is first recognized, and then process situations are sought out where this provides a radically new solution. As part of the BPR approach, *workflow management* systems are considered as significant contributors to process improvement.

Where an enterprise resource planning (ERP) solution is an integral part of the BPR effort, the project planning and management function of the ERP system may be utilized. Major ERP systems offer these capabilities as part of their system packages.

BPR treats business processes as proceduralized workflows. It links functional expertise with process management teams to reintegrate the existing, fractionated workflow into contiguous, cross-functional work processes that are driven by *customer needs*.

Infrastructure

Reengineering is done by trained and qualified people. Selecting and organizing those who actually

implement reengineering is critical to success. Certain people have unique roles that focus either on the individual process, or on aligning multiple processes at the enterprise level; others play dual roles as they carry out the tasks of bridging and integration at both levels.

Although there is no one-size-fits-all list of roles, many years of reengineering experience have helped identify the following generic reengineering infrastructure:

- *Enterprise transformation executive*—the senior executive who has the authority and resources to decide and motivate the overall reengineering endeavor.
- *Enterprise transformation council*—a group of influential policy makers who develop and champion the overall process reengineering strategy of the enterprise, commit the necessary resources, and oversee its progress. They are responsible for achieving synergy across the separate, individual reengineering projects of the enterprise and for creating optimal leverage on overall business performance.
- *Process management executive* – a senior manager who is the owner of the enterprise methodology of business process management, and is responsible for developing and maintaining methods, tools, and techniques – including the latest IT and computer software – related to BPR.
- *Business process owner* – an executive or middle manager with the responsibility and accountability for a specific, single business process, its reengineering, and subsequent operational implementation. There are usually a number of business process owners within an enterprise.
- *Business process management team* – a group of selected individuals (usually managers or senior professionals) working with the process owner and committed to the reengineering and subsequent operation of a specific, single business process.
- *Business process management team leader* – a member of the business process management team selected by the process owner to lead and administer the day-to-day work of the process team.

These roles do not represent a hierarchical set. Rather, they are part of a two-level team structure where the transformation council is the strategic

entity, and the business process management teams represent the operational part.

Budget

The process implementation budget contains both expense and capital items. It is always preferable to establish a consolidated process reengineering budget, as opposed to distributing the cost implications among the affected functional operating budgets.

Costing methods used in establishing the implementation budget should be consistent with those employed to plan and measure the continuing process costs once the reengineered business process has become operational. Since most process costing methods are activity or task based, it is recommended that *activity-based costing* (ABC) be used in building the process implementation budget.

It is important to point out that all implementation activities in the project plan must contain, beyond people and technology, all the required resources: materials, energy, plant and facilities, intellectual capital, and other assets needed to ensure the successful plan implementation.

It may happen at times that the deployment of a newly reengineered business process may take more time and effort than its design. This will more likely be the case where a massive upgrading of IT, such as the implementation of an ERP system, is part of the reengineering effort.

Strengths

If implemented properly and with the committed and continuing support of senior management, BPR can yield very significant returns. It helped corporate giants like Procter & Gamble and General Motors recover after financial drawbacks due to competitive pressures; Southwest Airlines reengineered the company using IT; Michael Dell understood how and where to apply the concept of BPR to reengineer Dell Computers; Ford reengineered its manufacturing processes where *job #1* is product quality.

Through a smoothly functioning reengineering infrastructure, accountable process ownership, and effective teamwork, reengineering provides the solution to the two known dilemmas of process management:

- variation and dispersion and
- scale and complexity.

Variation and dispersion is the apparent conflict between the need for process consistency across the enterprise and the variation caused by geographic dispersion and differences in market requirements. Reengineering tries to solve this dilemma by creating different versions of the same basic process flow, each under the leadership of a working owner, and assigning an operations director to each instantiation of each version, to deal with variation in market (customer) requirements.

Scale and complexity are the combined effects of process size (amount of resources, number of operating locations) and structure (number of constituent subprocesses and activities, handoff points, multiple process versions). Scale and complexity can make a single process owner's job extremely challenging. Again, the suggested approach is subdivision of the process and the assignment of additional process management resources. Subprocess owners operate under the overall management umbrella provided by the executive process owner.

The idea of reengineering was rapidly adopted by a huge number of firms around the world, but most notably in North America. They were striving for renewed competitiveness, which they had lost due to:

- the market entrance of new foreign competitors (e.g., Japanese manufacturers in the United States);
- their inability to satisfy changing customer needs (e.g., IBM in the mid-1980s); or
- their flawed cost structure.

Reengineering was implemented at an accelerating pace and by the late 1990s more than half of the Fortune 500 companies claimed to either have actually initiated reengineering efforts, or to have firmly planned to do so.

Weaknesses

To this day, a number of critics claim that BPR dehumanizes the workplace. By the late 1990s, BPR gained the somewhat unfair reputation of being an outright excuse for "downsizing", e.g. major and irreversible reductions of the workforce.

A number of factors can cause BPR initiatives to fail to reach their initial objectives. The reasons for

4 Process Reengineering

such failure are many; however, most can be avoided or corrected:

- lack of sustained management commitment and leadership;
- unrealistic scope and expectations, or prior constraints on problem definition or project scope;
- underestimation of the cultural resistance to change;
- overconfidence in technological solutions;
- poor **project management**;
- concentration on redesign as the only alternative;
- neglecting people's values and beliefs;
- settling for partial or minor results, or quitting too early;
- making reengineering happen as a grassroots effort;
- dissipating energy and resources across an excessive number of reengineering projects; and
- failure to distinguish reengineering from other improvement programs.

The traditional process reengineering methodology evolved around the concept of physical and material workflow, its improvement and redesign. This is because the underlying methodologies have their roots in manufacturing and, in particular, in IT, whose purpose is to automate defined procedures.

This approach to BPR has led to the development of skillful ways of redesigning business processes to measurably improve performance, and to use IT for the streamlining and coordination of complex workflows. However, as companies embarked on using this approach to performance improvement, it eventually became clear that some important things were being overlooked.

The traditional process reengineering approach did not adequately address several fundamental areas [4]:

- the human dimension, including the flow of people-to-people communications in the process;
- the planned adaptability of redesigned processes to deal with further unpredictable change, and
- overall economic and business performance of the enterprise.

Overlooking the Human Dimension

Overlooking the human dimension – mobilizing and motivating people – has been viewed as a major

shortcoming of classic workflow-based process reengineering. The downsizing, done without concern for the people, has made the term *reengineering* imply that people do not matter because the approach taken is mechanical and impersonal.

Another limitation of traditional reengineering is its inherent difficulty in harnessing the great variety of intellectual, managerial, and professional communications where there are no easily definable physical work products – e.g., most work results are agreements, decisions, commitments, and the transfer of information and knowledge.

Problem with Unpredictable Change

One of the greatest challenges to BPR is that it must engender thinking that helps anticipate and deal with *unpredictable change* – the ability to quickly identify and abandon old paradigms that underlie current operations.

The Process Paradox

Success at the business process level, but no improvement or even negative overall impact at the enterprise level, is the essence of what has become known as the *process paradox*. After months or even years of careful process redesign, many companies have achieved impressive improvements in individual business processes – only to watch overall results still decline. By now, paradoxical outcomes of this kind have become almost proverbial.

- IBM's small business division won the Malcolm Baldrige Award in 1992 in recognition of the successful redesign of the development and production processes that produced one of IBM's all-time success products: the AS400. At the same time, the IBM Corporation was facing the worst crisis of its history including the downsizing of half of its workforce worldwide.
- In 1995, Corning's telecommunications products division earned the Malcolm Baldrige Award. Six years later the company was experiencing what many believed was its fight for corporate survival.

Summary

Today, defining business processes as the starting point for business analysis and redesign is a widely

accepted practice and standard part of the change management portfolio. Change, however, is sought and implemented in a less radical way than what was originally proposed by BPR.

The key task of twentieth century executives and managers was *performance improvement*. The solutions of TQM and BPR were *process solutions* – focusing primarily on key, individual business processes.

Now, management faces the compelling requirement of doing all this continuously, faster, and better than competition – and in the face of accelerating and unpredictable change. The double challenge of the twenty-first century is to be successful while continuously *adapting* and responding to a highly competitive, fast moving, information-driven, and near-turbulent environment. In fact, “continuous adaptation” is the motto for the new century.

References

- [1] Hammer, M. (1990). Reengineering work: don't automate, obliterate, *Harvard Business Review* July-August 1990, 104–112.
- [2] Hammer, M. & Champy, J. (1993). *Reengineering the Corporation: A Manifesto for Business Revolution*, Harper Business, New York.
- [3] Davenport, T. (1993). *Process Innovation: Reengineering Work Through Information Technology*, Harvard Business School Press, Boston.
- [4] Pall, G.A. (2000). *The Process-Centered Enterprise: The Power of Commitments*, CRC Press, Boca Raton.

GABRIEL A. PALL

Kaizen

Introduction

Process improvement involves the use of many tools and techniques. Always the aim is to improve some aspect of performance, quality, productivity, competitiveness, or profitability. One appealing technique is Kaizen. Kaizen is a Japanese word which may be translated as gradual change for the good or continuous improvement. Kaizen projects often feature as part of a lean **six sigma** initiative.

A characteristic of Kaizen projects is that they are typically targeted at “quick fix” improvements that can be achieved through a short and intense period of project work (generally no more than one week). Kaizen projects, therefore, differ from lean six sigma projects which are usually expected to take a period of months.

Kaizen projects are tackled during what is often termed a *Kaizen event* or *Kaizen blitz* [1]. For a project to be suitable for a Kaizen event it must be well defined and have clear boundaries. In addition, it must be possible unambiguously to measure the results, for example, the time taken to complete a task, the number of steps taken, or the number of defective units. A Kaizen event can be used to tackle one component of a larger project as long as the scope of the Kaizen subproject is well defined.

The most common objective of a Kaizen event is to improve process flow by attacking process waste. The solutions identified are generally low tech, for example, the rearrangement of process steps into a more efficient order. Kaizen is suitable for production and nonproduction areas of a business [2]. Examples of Kaizen projects include streamlining an invoicing process, or reducing waiting time in one section of a process.

Teams

Kaizen events are carried out by small teams of people. Kaizen is a powerful improvement tool because team members are taken out of their day-to-day responsibilities and allowed to have ring-fenced time to think about a problem and its solution. Price and Williams [3] state that “companies which use Kaizens have found they generate energy among those who

work in the area being improved, and produce immediate gains in productivity and quality”.

For Kaizen to be successful the team members have to be open-minded and treat everyone in the team as equal partners. The following are guidelines:

1. Encourage all team members to think.
2. Think positively and focus on “how” rather than reasons why not.
3. Think widely rather than being tied to fixed, conventional ideas.
4. Try to get to the root cause of a problem; the five why technique, for example, may be useful. Many Kaizen solutions seem obvious in hindsight.
5. Take action rather than holding out for perfection. Do what can be done now with the resources at hand.
6. Use the ideas and insights of everyone in the team. If someone suggests a solution that works well in another department, then consider it as a potential solution for the department currently in focus, even though the situations may be different.

As the implementation of new ideas is an integral part of a Kaizen event it is vital that the Kaizen team members have the authority to put process improvements in place. The team working on a Kaizen process improvement project should either include the process owner or have their authorization to implement the solution. Otherwise, there will be a delay before the results of the Kaizen event can be put into practice. “Quick wins” are often achieved at Kaizen events but it is important that whatever is changed as a result of the event, can be changed back, if necessary, to reduce the risk of unforeseen consequences. This is particularly important in highly interrelated specialties such as healthcare.

Not everyone who will be directly affected by changes will be in the Kaizen event. The concerns of other individuals should therefore be taken into account. Good communication is vital.

Types of Projects

Kaizen projects are typically aimed at reducing waste within a defined process. The following types of project can be well suited to Kaizen events:

2 Kaizen

- Projects which seek to increase throughput in a particular process.
- Projects which seek to decrease waiting times for internal customers.
- Projects that look to improve staff or resource utilization.
- Projects that aim to address processing time for tasks – or to introduce same day processing for tasks.
- Projects that aim to reduce the amount of physical space required for a process.

Each potential Kaizen project should be judged individually on its ease of implementation and potential benefit. This can be done using a payoff matrix. The matrix has two axes; the x-axis is the business impact of the subproject, the y-axis the ease of implementation of the subproject.

Each member of the Kaizen team rates each possible project for ease of implementation and potential improvement impact. Those projects falling into the upper right quadrant of the matrix should be tackled first.

Ease of implementation	Easy		
	Hard		
		Low	High
		Business impact	

Often the initial scope of a project may be too wide to make it suitable for a Kaizen event. The payoff matrix can be a simple and effective tool for sorting through possible subprojects within a larger project scope. Later on during the Kaizen event the payoff matrix can also be used to sort through and prioritize potential solutions.

Following Kaizen processes are generally considered to have improved if:

- Their duration time has decreased.
- The space required for the process has decreased.
- Fewer resources are used (e.g., people, machines, material, energy, and information).
- Outcomes improve (e.g., quality, customer satisfaction, and cash flow).

Kaizen projects are about identifying and trying possible solutions to a problem within a short space

of time. The project will run smoothly if necessary supporting evidence and data are available and accessible at the start.

Kaizen Project Worksheet

Once a Kaizen project has been selected it is a good idea to develop a Kaizen project worksheet. The worksheet should contain all the project details, including an outline of the problem to be addressed and the project objectives. The worksheet is useful for: clarifying the problem, ensuring the problem is understood by the whole team, and recording the current state of play. The worksheet is also useful for referring back during the course of the project to ensure that the original issue is still the one being addressed, in other words that the project scope has not inadvertently changed.

Kaizen Methodologies

Methods and tools used during a Kaizen event may include any of those encountered in the define, measure, analyse, improve and control (DMAIC) stages of a lean Six Sigma project. The event works best if data, such as demand, cycle time, voice of the customer information, waiting and process times are already available. Summary statistics, such as the mean cycle time are estimated. This makes sure that specific, measurable, achievable, realistic, and with a known time scale (SMART) objectives can be set. It may be possible to determine an optimal rhythm for the process in terms of available time per unit of demand (sometimes called the *takt time*). Particular emphasis is placed on understanding the process as it currently is. Often the process is walked through and the team make observations and talk to people involved in the process. The team members create a **process map** that shows the flow and any rework loops. They try to quantify the basic process steps in terms of activity times and waiting times. Value adding and nonvalue adding process stages are identified. The team then tries to come up with possible process changes and solutions; techniques such as brainstorming, theory of inventive problem solving (TRIZ), and six hats can be used to help. When solutions are identified, the results of implementing them can be investigated empirically or by simulation.

Outline of Kaizen Event

Clear examples of step-by-step procedures are given in [3] and [4]. In summary, the steps are:

- Select the project area and a small number of relevant team members.
- Develop a project worksheet.
- Assess the current situation in appropriate ways, such as looking at process data, stepping through the process, assessing the voice of the customer, identifying value adding and nonvalue adding process stages. Agree SMART objectives.
- Through brainstorming and other methods, identify factors that could be causing variation, inefficiencies, waste, or delay.
- Refine the factors to be investigated using the payoff matrix to find a subset to be addressed.
- Explore possible quick fix options to check their feasibility and check for any adverse effects. Refine the options using the payoff matrix to find a subset to be addressed.
- Trial implement of possible solutions by stepping through the alternative processes, maybe use simulation.
- Inform people affected by the changes and inform everyone involved about the changed work practice.
- Document, implement, evaluate change, celebrate, and agree follow-up procedures to ensure that the change is part of continuous improvement.

Example of a Kaizen Worksheet for a Health Care Kaizen Project [5]

Date: 27/02/06

Team: J. Richards, S. Smith, T. Jones, and D. Young

Project Location

Admission and testing facility.

Problem Statement

Patients presenting for admission are asked to wait before being called for admission. Once called, they are admitted and asked to wait in another waiting room.

Waiting time can be up to 2h resulting in poor patient satisfaction.

- Patients have no information on how long they will wait.
- Staff roles and responsibilities are not clearly defined resulting in poor patient case management.
- Patients and patient notes are moved around the facility from room to room.
- Reception and admission staff are frequently interrupted adding to wait times.

Current State Measurements and Demand Requirements

Average wait time in waiting room prior to first call = 52 min.

Average total lead time = 132 min.

Rough estimate of initial takt time for assessing process (one member of staff).

Takt time = Net available time to work/customer demand.

Shift length = 7.5 h; lunch break = 1 h; net avail time = $7.5 - 1 = 6.5$ h.

Current customer demand = 10 patients per day.

Takt time = $6.5/10 = 0.65$ h (39 min), one patient every 39 min.

Objectives

To reduce the average patient waiting time at first wait point to 20 min.

To reduce average lead time by 25%.

Process Owner

J. Richards

Process Owner Signature

J. Richards

Strengths and Weaknesses of Kaizen

The strengths of the Kaizen method include:

- If quick wins exist they are likely to be achieved.
- Involves people, but only needs a small time commitment.

- Can be done with short notice.
- Relatively low cost.
- Flexible.
- Equally applicable in production and nonproduction.
- People oriented.
- Enjoyable and strengthens team.

The weaknesses of the method include:

- High expectations and potential for disappointment.
- If problem is not well chosen it fails.
- Limited time for training in methods.
- May not lead to workable solutions.

Summary

Kaizen is a positive way to address well-defined problems with clear boundaries. It increases the chances for a quick solution by bringing relevant people together with a mutual aim of spending ring-fenced time thinking about a problem. Practitioners find the method useful and interesting. It brings a team together and is good for addressing high priority issues especially those that lead to waste, delays, and inefficiency.

References

- [1] Laraia, A., Moody, E.P. & Hall, R. (1999). *The Kaizen Blitz: Accelerating Breakthroughs in Productivity and Performance* (National Association of Manufacturers), John Wiley & Sons, New York.
- [2] Keyte, B. & Locher, D. (2004). *The Complete Lean Enterprise, Value Stream Mapping for Administrative and Office Processes*, Productivity Press, New York.
- [3] Price, M. & Williams, T. (2007). *Fast and Intense: Kaizen Approach to Problem-Solving*. Available on ISixSigma® at Address <http://www.isixsigma.com/library/content/c050601a.asp>.
- [4] Hahlen, B. & Mahoney, S. (2007). *Applying Transactional Lean at Medical Imaging Firm*. Available on ISixSigma® at Address <http://healthcare.isixsigma.com/library/content/c051012a.asp>.
- [5] van-Lottum, C.E. Example Kaizen worksheet for a health care project. available on request from isru-enquiries@ncl.ac.uk.

Related Articles

Lean Methods, Creating Flow with; Six Sigma.

SHIRLEY Y. COLEMAN

Work-Out

Introduction and Background

Jack Welch developed the GE Work-Out in the late 1980s as a unique approach to improvement in the workplace. He had recently developed the Crotonville leadership development center, and was delighted with the open atmosphere there, which allowed business leaders to share both their ideas and also their concerns about the issues facing GE. Although thrilled with the openness he saw in “the pit” – the main meeting auditorium at Crotonville – Welch remained frustrated that he did not sense a similar environment elsewhere in the business. In too many places he saw a bureaucracy filled with controlling leaders who pretended to have all the answers, and a stifled workforce afraid to open up and share their true feelings.

Realizing that he could not bring all GE employees to Crotonville, Welch decided to take Crotonville to all the employees. He created a process that mimicked what he experienced at Crotonville, and could be replicated in various locations around the world. Perhaps owing to his Massachusetts roots, Welch decided to model the process after a New England town meeting. There would typically be 40–60 people, facilitated by someone outside of GE (originally local college professors) to ensure that people felt free to speak openly. The fact that the managers of the participants were not invited to the Work-Out was critical! The facilitator would define the scope of the issue being addressed, and then get people to share their ideas and frustrations about it. Welch [1] notes that at some of these meetings people shared such passionate feelings that they began to swear at each other. However, they got past the bureaucracy to the core issues, which often involved unnecessary reports and meetings that got in the way of serving customers. Much of the focus was, and still is, on eliminating nonvalue-adding work.

The Work-Outs were a big success, and eventually involved hundreds of thousands of GE employees. Not only did they produce tangible results, but they also helped people feel empowered after some difficult years of downsizing in the company. They helped create an atmosphere where employees worked across global and organizational boundaries for the benefit of

GE as a whole, and where good ideas were respected regardless of the rank, gender, age, or ethnicity of the person submitting the idea. This atmosphere was later given a name by Welch – *boundarylessness*. Of course, GE Work-Out has morphed and evolved over the years. However, it has retained the key principles upon which it was founded, and is still used today at GE and many other companies.

The Essence of GE Work-Out

So, what are the essential elements of GE Work-Out that make it a unique approach to improvement? It was founded on the following core principles:

- an open environment where people are free to speak their mind, without fear of retribution by management;
- enlistment of the “frontline” workers – those most familiar with the issue at hand, rather than a “panel of experts”;
- use of impartial facilitators rather than management to conduct the meetings;
- short interactions (2–3 days) and immediate decisions by management for at least 75% of the team suggestions – no burying of the proposals by sending them to committees for “further study”;
- an emphasis on “busting bureaucracy”, that is, eliminating unnecessary meetings, reports, approvals, and so on.

One may immediately see some similarities between the principles of Work-Out and those of lean *kaizen* events (see **Lean Methods, Creating Flow with**), and perhaps also to those of quality circles. Through the years, GE Work-Out has evolved. Today, it may involve only a small team of 5–10 people, and may only last a day or even an afternoon, depending on the circumstances. As GE gained more experience with Work-Outs, it began using internal facilitators rather than relying on external resources. However, GE has maintained a focus on these original principles. For example, when an internal facilitator is used, this person will be impartial, and not have a clear stake in the solution selected.

The process for conducting a GE Work-Out is fairly simple, at least in principle. Typically a business leader initiates the Work-Out by stating the problem, defining the scope, and explaining any constraints on the solution (e.g., “we can’t spend

any more than \$20 000”). The business leader then departs, leaving a facilitator in charge of the process. The facilitator walks the group through a team problem solving process (see, for example, Hoerl and Snee [2]). This process provides plenty of time for the Work-Out team to discuss the problem in detail and brainstorm ideas for solutions, perhaps in subgroups. The facilitator collects the ideas and ensures that the entire team carefully considers and debates them. The best ideas may rise to the surface, or the facilitator may use some type of prioritization method to identify the most promising potential solutions.

Regardless of which method is used, the team needs to refine the ideas into a few specific suggestions for management. These cannot be vague concepts, such as “improve morale”, but must be tangible, specific suggestions, such as “reduce the approvals needed on this form to two: the shift supervisor and the plant manager”. Once the suggestions have been carefully defined, they are presented to management at the end of the Work-Out. Management reserves the right to decline suggestions. However, as noted previously, they are expected to give an immediate answer on at least 75% of the proposals. On those suggestions for which management needs more time, they must give a date by which they will make a final decision.

The Role of Work-Out within a Process Improvement System

Clearly, the Work-Out approach is not the best method to solve every problem that may occur within an organization. There will be complex technical problems that require more sophisticated, data-based approaches, such as Six Sigma (*see Six Sigma*). However, for certain types of problems Work-Out is an excellent approach, and will likely produce better results more quickly than other options. Based on experience, Work-Out is best suited to attack problems that have the following attributes:

- We believe that a relatively simple solution exists, one that will not require extensive **data collection** and analysis to find.
- We believe that knowledge of potential solutions already exists in the people working within the process.

- We believe that the wisdom of good solutions will be readily apparent to those working within the process, minimizing the need for a “test period” to evaluate the solution.

The rationale for the first two points should be clear, given the nature of a GE Work-Out. It is a process designed to get the best ideas and knowledge out of people, rather than to create new knowledge. The reason for the third point is that we want management to make immediate decisions about the proposals, rather than sending them off for further review or study. If a trial period is required to evaluate the proposal, this will slow down the implementation of ideas, defeating the purpose of Work-Out.

Some example problems for which Work-Out might be appropriate include the following:

- Why are seven levels of approval needed on some financial documents? Is there no way we could simplify the approval system, and still maintain good control?
- We seem to spend way too much time preparing reports, many of which are not used. How can we reduce the time spent on report preparation, and still provide people with the information that they need to do their jobs?
- There have been a number of defects related to human errors on one step of our assembly process. How can we revise the way we complete this assembly step to eliminate the errors?
- The metrics that we are using in our operation are driving the wrong behavior. How can we revise the metrics to be more aligned with the behavior we actually want to see?

An organization’s process improvement system should include informal, knowledge-based methods like Work-Out, as well as more sophisticated approaches involving the use of statistical methods to develop new understanding of processes, through such tools as designed experiments. No one improvement method will be best for every problem, so organizations need a diverse improvement “toolkit”. Once a problem has been identified, scoped, and critically evaluated, people can make the decision as to which improvement method is best suited to attack it.

Interestingly, GE has occasionally found that after the Define, Measure, and Analyze phases of Six

Sigma projects, solutions to the problem are readily obvious to the people who work in the process. Rather than a rigorous Improve phase of the project utilizing statistical techniques, some teams have found that they can hold a one-day Work-Out and implement an effective solution immediately. After putting controls in place to maintain the improvement (Control phase), the project becomes a DMAWC (define, measure, analyze, work-out, control) Six Sigma project, rather than the typical DMAIC (define, measure, analyze, improve, control), and is completed more quickly than anticipated. This is one example of integrating improvement methodologies to obtain synergistic benefits.

Summary

GE Work-Out is a proven improvement methodology that has produced significant results not only for GE, but also for many other companies. It is most

appropriate for problems having a solution that is already captured within the collective knowledge and experience of the people working in the process. As such, it is typically classified as a “knowledge-based” rather than a “data-based” improvement approach. In addition to tangible business benefits, GE Work-Out can also help create a “boundaryless” work environment where people feel free to express their ideas and opinions, and are anxious to work across organizational and hierarchical boundaries.

References

- [1] Welch, J.F. (2001). *Jack: Straight From the Gut*, Warner Business Books, New York, pp. 183.
- [2] Hoerl, R.W. & Snee, R.D. (2002). *Statistical Thinking for Business Improvement*, Duxbury Press, Pacific Grove, pp. 457–464.

ROGER W. HOERL

Statistical Quality Control

Introduction and History

Statistical quality control (SQC) may be broadly defined as “the use of statistical methods in the control and improvement of quality in industrial production” [1]. Historically, these methods have included **control charts**, tolerancing methods, sampling inspection, and a wide variety of other statistical methods, though those listed are most commonly regarded as being the core of SQC. Ishikawa [2] says that “Expressed simply, the relationship between statistical quality control, total quality control, and research and technology is as follows: Good quality control is impossible without proper technology. The motive power behind the search for causes is research, technology, and technical skill (i.e., experience and training). However, technology improves rapidly through the use of the S part of SQC (statistical methods), i.e., through carrying out the quality analysis and process analysis using the statistical tools of the QC approach. We must use proper technology, statistical techniques, and control techniques as tools to control quality and promote effective TQC.”

The work of Frederick W. Taylor in the early twentieth century in introducing the principles of what he termed *scientific management into mass production* may be properly regarded as a precursor to the field of SQC. Taylor pioneered the use of the scientific method in the study and improvement of manufacturing processes. Although his image has historically suffered due to the perception of his work as leading to division of work into small tasks and consequent suboptimization, he clearly pioneered the development of standardized production and assembly methods and the study of the quality of manufactured goods.

The field of SQC may be regarded as having been launched, however, on May 16, 1924, when Walter A. Shewhart of Bell Telephone Laboratories circulated a memorandum which included a modern control chart, the first one known to have existed. Shewhart continued to evolve the concept of control charts, developing new charts for a variety of applications, and in 1931 published his classic book *Economic Control of Quality of Manufactured Product* [3]. Other names involved in the early history of SQC at Bell

Telephone Laboratories include Harold F. Dodge, Harry G. Romig, Thornton C. Fry, E. C. Molina, and W. Edwards Deming.

By the end of the 1920s, Harold F. Dodge and Harry G. Romig of Bell Telephone Laboratories had developed a comprehensive system for statistically based **acceptance sampling**, which could serve as an economical alternative to 100% inspection, which was in common use at the time. By the end of the 1930s, the methods of SQC were being used broadly at Western Electric, though not in industry at large.

In 1938, W. Edwards Deming convinced Shewhart to give lectures at the graduate school of the US Department of Agriculture on the subject of SQC.

With the advent of World War II, the method of SQC found broad usage in assuring the quality of products. This usage was often demanded by the US Government as a condition of purchase. This was necessary as the quality of manufactured goods plummeted due to the unavailability of skilled personnel driven by the war effort. Training courses were established by many companies as well as the War Production Board and the armed services. These courses covered the use of SQC, with the topics of control charts and acceptance sampling being taught in depth. The result was that many companies were using the methods of SQC by 1945.

Sadly it appears that these methods fell into disuse shortly after the end of World War II. Companies were able to sell all they could produce, and volume production was the rule of the day. The use of statistical control charts and inspection (acceptance sampling) methods declined for a while in the United States, and this trend was not actually reversed until the 1980s when the competitive position of US manufacturers was clearly in decline as compared to Japanese manufacturers, particularly in the automotive and electronics industries.

In spite of the decline beginning in use of SQC methods in 1945 with the end of World War II, the American Society for Quality Control was founded to sustain the quality-improvement techniques that had been used during wartime. The organization continues to this day (with its name being changed to American Society for Quality) as an advocate of modern methods of SQC as well as a broader advocate for quality.

The Japanese experience with SQC in the 1950s provides an interesting and notable contrast. Ishikawa

[4, pp. 18–19] says that SQC was overemphasized in Japan in the 1950s, creating problems including (a) workers rebelling because they felt their experience and common sense were more important, (b) lack of standards to use, (c) lack of data, (d) sampling methods not followed, and (e) measurement devices destroyed by workers suspicious of the reason for installation of the devices. Ishikawa summarizes [4, pp. 18–19] “It is true that statistical methods are effective, but we overemphasized their importance. As a result, people either feared or disliked quality control as something very difficult. We overeducated people by giving them sophisticated methods where, at that stage, simple methods would have sufficed Top and middle-level managers did not show much interest Those who were members of the Quality Control Research Group tried to persuade top managers to join, but perhaps because of our relative youth our efforts were met with little visible success.”

In recent decades the set of methods falling under the category of SQC has expanded. Recent textbooks on the subject often include not only control charts, tolerancing methods, and acceptance sampling, but also designed experiments – full and fractional factorials as well as **response surface** methods.

Controversy in the Field of Statistical Process Control

There have often been disputes in the field about the appropriateness of the use of sampling inspection methods. Deming [5, 6] has cogently argued the viewpoint that sampling should be either 100 or 0% for an in-control process (other than for the purpose of verifying the ongoing state of control of the process).

By-products of Use of Methods of Statistical Process Control

Grant and Leavenworth [7] point out that practitioners have often noted that the use of the methods of statistical process control tends to result in a greater knowledge through engineering and other staff of the real problems of production. In addition, there is the greater reality in specified tolerances which can result from a true knowledge of actual process capabilities.

The use of statistical methods for process study has the additional benefit of taking arguments between production, engineering, and quality

departments and enabling these arguments to be transformed to discussions based on real facts provided by the use of SQC techniques. In fact, SQC can constitute a common language to be used in communication between the three groups as they work to arrive at a reasonable solution to their common difficulties.

Tools of Statistical Quality Control

No proposed list of the tools of SQC could be complete. However, the following list is a portion of those covered by Wadsworth *et al.* [8] and is suggestive of the most commonly used tools and the breadth of application of the methods.

Control Charts: *P* charts, *C* charts, *U* charts, individuals and moving range charts, *X*-bar and *R* charts, **exponentially weighted moving average** (EWMA) charts, cumulative sum (**CUSUM**) charts.

Graphical methods: checksheets, histograms, **Pareto** diagrams, Ishikawa (cause-and-effect or fish-bone) diagrams, scatterplots, flow charts, location plots.

Acceptance sampling: single and double **attribute sampling** methods, **sequential sampling** plans for attributes or variables.

Tolerancing methods. *Designed experiments:* **factorial** experiments, **fractional factorial** experiments, evolutionary operation, response surface methods, analysis of means.

Reliability: reliability modeling, accelerated life testing.

Broadening the Application of the Methods of Statistical Quality Control by W. Edwards Deming

W. Edwards Deming was responsible for a significant broadening in the scope of application of the methods of SQC. He examined the use and philosophy of the control chart as invented and advocated by Walter Shewhart, and commented [5] “Elimination or reduction of common causes can be effected **ONLY BY ACTION OF MANAGEMENT.**”

He expanded on this in [6, p. 310] to say

“Shewhart used the term **assignable cause** of variation where I use the term special cause. I prefer the adjective special for a cause that is specific to some group or workers, or to a particular production worker, or to a specific machine, or to a specific local condition.”

And in [6, p. 314]

“The term common causes for faults of the system, was used first, so far as I know, in a conversation held about 1947 with Dr. Harry Alpert (deceased) on the subject of riots in prisons.”

And Deming in [6, pp. 320–321]

“Removal of common causes of trouble and of variation, of errors, of mistakes, of low production, of low sales, of most accidents is the responsibility of management.”

This significant philosophical broadening of the applicability of the methods of SQC has led practitioners to use the control chart as an analytical tool for deciding whether the responsibility for improvement in a given context is local (workers) or global (managerial). This broadening was instrumental in forging a pathway for providing insights from statistical thinking in the practice of management, and particularly that of process management.

References

- [1] Burr, I.W. (1976). *Statistical Quality Control Methods*, Marcel Dekker, New York.
- [2] Ishikawa, K. (1989). *Introduction to Quality Control*. 3A Corporation, Tokyo. Translated by J.H. Loftus.
- [3] Shewhart, W. (1931). *Economic Control of Quality of Manufactured Product*, Van Nostrand, New York.
- [4] Ishikawa, K. (1985). *What is Total Quality Control: The Japanese Way*, Prentice Hall, Englewood Cliffs. Translated by David J. Lu.
- [5] Deming, W.E. (1982). *Quality, Productivity and Competitive Position*, Massachusetts Institute of Technology.
- [6] Deming, W.E. (1986). *Out of the Crisis*, Massachusetts Institute of Technology, Cambridge.
- [7] Grant, E.L. & Leavenworth, R.S. (1996). *Statistical Quality Control*, McGraw-Hill, New York.
- [8] Wadsworth Jr, H.M., Stephens, K.S. & Godfrey, A.B. (1986). *Modern Methods for Quality Control and Improvement*, John Wiley & Sons, New York.

WILLIAM C. PARR

Critical-to-Quality Matrices

Introduction

In the 1970s a new technique to quantify customer needs began to be applied in the shipyards of Kobe, Japan. In the 1980s, this approach began to spread to other countries, finding a particularly strong reception among automobile manufacturers in the United States. This technique became known in English by its literal translation from Japanese, **quality function deployment**, or QFD. The cornerstone of this method was the use of a sequence of matrices to translate and prioritize attributes at finer and finer levels of detail. For example, companies typically began with vague, fuzzy needs that they heard from customers, such things as “I want my car to have that new car smell,” “A good writing pen shouldn’t feel flimsy,” or “I want clear reception on my cell phone.” Companies needed to translate these needs into tangible metrics, things that could be measured in design and manufacturing, so that they could be confident that they were actually providing the characteristics that customers desired. They would further need to translate these overall performance metrics into functional requirements – performance metrics on subassemblies, then to prioritization of process steps to identify the “make or break” steps in the process where the performance metrics would be determined, and possibly to translate these into further detail, identifying the specific process control measures that could be used at these critical steps. See Cohen [1] or Clausing [2] for further details on QFD.

As should be evident from the above discussion, completing a complete, rigorous QFD study typically takes weeks, if not months. If this information is needed, then the time and effort will be well worth it. However, in many applications, a full QFD study is not needed. Rather, only the translation from fuzzy customer needs to tangible metrics, with prioritization, is required. This is often the case with **Six Sigma** projects. Many organizations have utilized only the first matrix in such cases, which also goes by the name House of Quality. In Six Sigma, the key performance metrics that teams try to improve are called *critical to quality* or CTQ metrics. Therefore, in the late 1990s the term *CTQ matrix*

began to become more popular than the term *House of Quality*. Whether used as part of a QFD study, within a Six Sigma project, or simply to clarify how to satisfy customer needs, the fundamental purpose is the same – to translate fuzzy, intangible customer needs into a set of tangible metrics, and to prioritize these metrics so that teams know on what to focus to ensure they are satisfying the customer.

The Elements of a Critical-to-Quality Matrix

Figure 1 shows one example of a CTQ matrix. This matrix is translating the needs of small business loan customers for a local bank, L&M Savings and Loan. Figure 2 shows a generic house, with each “room” labeled for reference. Room 1 lists the needs of the customer. These come from marketing research, focus groups, customer interviews, and so on. Note in Figure 1 that these needs are intangible, and stated in words that customers would use, such as “easy access to capital”. Toward the right of the graph, there is a column with relative priority ratings of these needs on a 1–5 scale, with 5 being the most important. These priorities should come directly from customers, through surveys, for example, rather than from the subjective opinions of people within the company.

In Room 2, L&M is comparing itself against key competitors, in terms of how well the customers’ needs are being satisfied. The key competitors in this case are First Union and County Bank. Note that these data are customer perceptions, obtained through customer surveys, or a similar customer-oriented method. This room does not list actual performance data on tangible metrics, which will come later. As we can see, L&M’s current offerings (the basis of the comparison) are better than competition in some cases, and worse in others. This room helps planners determine where they most need to improve with new products or services.

In Room 3 at the top of the graph, L&M has listed potential tangible metrics, that is, things they could actually measure, that might relate closely to customers’ perceptions of the needs listed in Room 1. For example, the ratio of the actual amount of the loan approved (how much was granted by the bank) to the amount requested – listed in the first column, is certainly measurable, and might relate to a customer’s perception of easy access to capital.

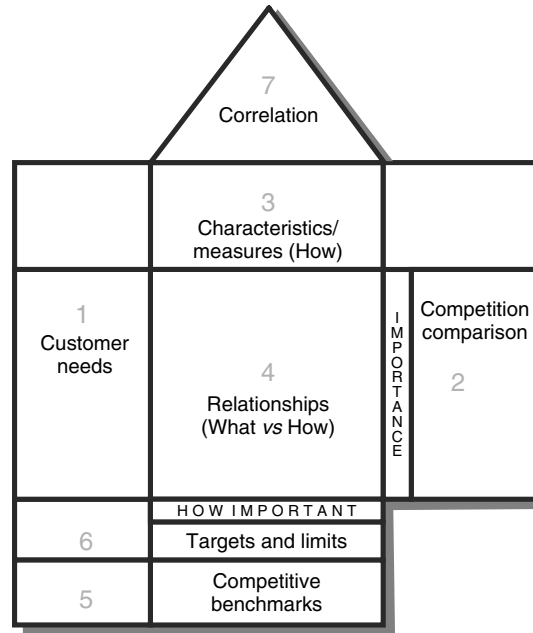


Figure 2 The rooms of the critical-to-quality matrix

needs in Room 1. This work is done in Room 4. Teams do this using the knowledge and experience that they have, which is why CTQ matrices are considered a knowledge-based tool. Of course, external data will typically be needed to identify and prioritize customer needs. Clearly, teams need to have the most knowledgeable people about these needs and metrics in the room when these decisions are made. Traditionally, teams assign one of four possible relationships to each what/how pair: strong, moderate, weak, or none/negligible not differentiating between a positive or negative correlation. As shown in Figure 1, these are often represented with symbols, which make it easier to see patterns in the graph. Note that each symbol also represents a numerical value, 9, 3, 1, or 0. These are traditionally on a nonlinear scale, as they are here, in order to highlight strong relationships. Here, the team believes that the ratio of approved to requested loan amount has a strong relationship to a customer's perception of easy access to capital.

Once the team has developed tangible metrics in Room 3, they have the ability to obtain actual performance data on competition and their own company's service, as shown in Room 5. In the case of physical products, such as an automobile, teams

can obtain data by physical tests of key metrics, such as horse power. For services, they need to obtain industry data, or data from third-party providers. In contrast to Room 2, these data represents hard metrics, not customer perceptions. Using these data, plus the relationship map in Room 4, the team can see which metrics they need to work on in order to impact customer perceptions. In fact, they can make this more formal by multiplying the need priority rating with the numerical relationship value, in order to obtain a relative importance for each metric. These values are shown at the top of Room 6 in Figure 1, and labeled CTQ priority. For example, the 56 associated with requested loan amount applied is obtained by multiplying 4.8 by 9 (strong relationship with easy access to capital), which equals 43.2, and adding $3 \times 4.3 = 12.9$ (moderate relationship with variety of terms/conditions), and rounding off to the nearest integer.

Using the above information, teams will often design the targets and possibly the specifications that they think they would need for the CTQs in order to achieve the customer perceptions they desire – in terms of needs being met – according to their intended brand positioning. That is, designers often manipulate Room 6, which they

4 Critical-to-Quality Matrices

have some degree of control over, in order to impact Room 2, for which they have no direct control. Obviously, modifying these targets in Room 6 will have cost implications that will need to be addressed.

When the team considers modifying targets in Room 6, it is often helpful to consider the potential interrelationships between these metrics. For example, if the team decides to increase the ratio of approved to requested loan amounts, that is, to give people more money, they might need to charge a higher interest rate to cover this additional risk. This would negatively impact the interest rate metric. In order to keep track of these interrelationships between the metrics, teams may choose to define and document these in Room 7 at the top. There are many scales that can be used for documenting these relationships. In this case, a plus symbol represents a positive relationship (if one metric is increased, the other is likely to also increase), a minus sign indicates a negative relationship, and no symbol indicates negligible interrelationship between the metrics. Due to the nature of this example, there are minimal interrelationships among the metrics.

Potential Benefits, Pitfalls, and Keys to Success for Critical-to-Quality Matrices

As seen in Figure 1, there is a tremendous amount of information that can be captured and easily communicated in a CTQ matrix. It lists customer needs and their relative importance, including how well the business is doing versus competition in terms of meeting these needs. It lists tangible metrics (CTQs), how these relate to the needs, their relative importance, and how well the business is doing versus competition on them. It may also include a depiction of how these metrics relate to one another. All of this information allows development of targets and possibly specifications that would be required to achieve a brand strategy. This is quite a lot of information on one graph!

Potential benefits of utilizing CTQ matrices include:

- As noted, the matrices capture a tremendous amount of information in one relatively simple graph.
- Cross-functional teams develop the matrices, enhancing organizational cohesiveness.

- Internal decisions are directly based on customer needs, and this fact can be easily communicated to others in the organization.
- Assuming that the needed information is available, the matrices can be completed fairly quickly, with minimal training of the team.

While there are several unique benefits of CTQ matrices, as with all improvement methods there are also limitations, or potential pitfalls. It is particularly important to understand such pitfalls when using knowledge-based methods, since the team may not have objective data to verify the results. Potential pitfalls of CTQ matrices include:

- Without proper information available, the team may end up guessing as to the real needs of the customer. The results of the process will never be better than the collective knowledge on the team.
- It is possible to arrive at the most popular answers, rather than the most accurate answers. This is referred to as *breathing your own exhaust*.
- The focus can divert toward filling in matrices, rather than understanding customers' needs and making important business decisions.

Given these benefits and pitfalls, we can derive keys to success that will help ensure good results. Keys to success when developing CTQ matrices include:

- Make sure that the team's opinions are based on objective data, especially as it relates to documenting customer needs.
- Fulfill the spirit of this method by having cross-functional teams work together. Have the right people in the room when their expertise is needed – avoid guessing.
- Have a strong, impartial facilitator lead the team through this process, someone who will ensure open participation, and prevent anyone from “taking over” the process.
- Verify the results *via* a “reality check” with key stakeholders. Ask if the results make intuitive sense.
- Remember that the process should work for the team, rather than the team working for the process. The focus must stay on satisfying customer needs, not on filling in matrices.

Summary

CTQ matrices are a proven method of developing products and services that will satisfy and delight customers. They can help translate fuzzy, intangible needs into metrics that organizations can measure and track. These matrices can also be used to prioritize customer needs and tangible metrics, helping organizations focus on the critical few variables. The process of having cross-functional teams complete the matrices helps organizational teamwork and communication. Properly completed, the matrix provides a tremendous amount of information on customer needs and their priority, key metrics on which the organization must focus to satisfy these needs, benchmarking of how well they are doing versus competition on both intangible needs and hard metrics,

and the targets and specifications needed to achieve the brand strategy. As with any knowledge-based method, teams must ensure that they have the right information on which to base their decisions, and that they verify their results.

References

- [1] Cohen, L. (1995). *Quality Function Deployment: How to Make QFD Work for You*, Addison-Wesley, Reading.
- [2] Clausing, D.P. (1994). *Total Quality Development: A Step-by-Step Guide to World-Class Concurrent Engineering*, ASME Press, New York.

ROGER W. HOERL

Concept Selection Matrices

Introduction

Concept selection refers to a step in a project during which several alternative approaches, design alternatives, architectures, or solutions – referred to as *concepts* – are evaluated and one of the concepts is selected for the project. The inputs to concept selection are the set of alternatives and a set of selection criteria, and the output is the selected concept – which may be one of the alternatives considered, or might be a new alternative generated during the concept selection process.

If there is only a single criterion for evaluation, then the concept selection process may involve optimization to that criterion. For example, if the decision will be based totally on a financial metric such as projected cost or profit, then the concept selection process could be based on **Monte Carlo** simulations of the financial results that would be forecast to result from each decision.

If there are both multiple alternatives and multiple selection criteria for a decision, a concept selection matrix approach may be appropriate. The goals of the selection process involve both tangible and intangible benefits. Not only should the selection process provide a superior product that will facilitate the development of a superior product, but it should also lend itself to “buy-in” from the team to go forward with the selected concept with confidence that a reasonable decision was made.

“Experience gained over many projects has led to the conclusion that matrices, in general, are probably the best way of structuring or representing an evaluation procedure the type of matrix referred to here is not a mathematical matrix; it is simply a format for expressing ideas and the criteria for evaluation of these ideas in a visible, user-friendly fashion” [1]. Let us illustrate with a simple example.

Example

Suppose that a family wants to consider doing something special together as a family. They come up with a set of selection criteria:

- It should be cost effective.
- It should have minimal impact on the children’s schoolwork.
- It should have minimal impact on mom and dad’s careers.
- Every member of the family should be able to get involved.
- It should include having the family members communicating with each other.

Each family member contributes one concept for doing something special together:

- Dad suggests a special “golf weekend”.
- Mom suggests preparing a pancake breakfast together (“breakfast of champions”).
- Susan, the elder daughter, suggests a “vacation in Rio de Janeiro”.
- Billy, the son, suggests having a “movie weekend”.
- Lisa, the younger daughter, suggests eating dinner out at a nice restaurant.

The family then chooses their current “family time” approach – watching television in the evening – as the benchmark, which Stuart Pugh referred to as the *datum*. Each of the concepts is compared against the datum, watching television, for each of the criteria.

As shown in Figure 1, the family would evaluate whether “vacation in Rio” is more cost effective (+), the same as (S), or less cost effective (–) than the benchmark, TV time, and would likely give it a (–). For the last criteria, all of the concepts except “movie weekend” would involve more communication than TV time; since the family talks a little during a TV show but are “shushed” if they talk during a movie, they give it a (–).

After filling out the matrix with +, –, and S symbols, the number of instances of each symbol is summed for each concept to provide a summary of how well each concept meets the selection criteria. In Figure 1, the first three concepts each have two pluses, but only “breakfast of champions” has no minuses.

Pugh Concept Selection Approach

The Pugh concept selection approach generally is applied in more technical situations, such as considering alternative designs for a product, or subsystem,

2 Concept Selection Matrices

	Datum/ TV time	Vacation in Rio	Breakfast of champions	Eat dinner out	Golf weekend	Movie weekend
Criteria						
Cost-effective		+ S —	+ S —	+ S —	+ S —	+ S —
Minimal career impact		+ S —	+ S —	+ S —	+ S —	+ S —
Minimal school education impact		+ S —	+ S —	+ S —	+ S —	+ S —
Every member of family can get involved		+ S —	+ S —	+ S —	+ S —	+ S —
Family time includes communication		+ S —	+ S —	+ S —	+ S —	+ S —
Sum of pluses		2	2	2	1	0
Sum of minuses		3	0	2	2	2
Sum of sames		0	3	1	2	3

Figure 1 Pugh concept selection matrix for family time example

or component of a product. It uses the method of controlled convergence, which can involve several iterations of the matrix.

Before the first iteration of controlled convergence, each of the alternative concepts should be developed to the same level of detail: if one concept involves a detailed sketch and a financial forecast, then all of the concepts should be developed to have a similarly detailed sketch and financial forecast so that the level of detail does not bias the evaluation process.

After the first iteration, the team can review the negatives of the strong concepts and propose ways to improve these negatives, generating more concepts. The team can also consider hybrids, where the negatives of a stronger approach coincide with positives of other approaches, suggesting a way to improve the stronger approach. It can also review some of the more promising of the weak concepts and propose ways to improve their negatives, generating more concepts.

The weakest concepts are then eliminated, reducing the matrix size. The best concept can now be used as the datum or benchmark, and the matrix rerun to see if the best concept continues to excel, converging on a superior concept to pursue. A second, more technical example of a Pugh matrix is shown in Figure 2,

for selection of technology for security by verifying a person's identity.

A major intangible benefit of the Pugh concept selection approach is that it is perceived by the team to be unbiased, so that the team can achieve consensus with the feeling that the best concept won.

It is also simple enough that the team does not require extensive training to use the approach, although it would be helpful if the facilitator has some training in terms of skills and understanding of the process.

Weighted Matrices

A question that often comes up with respect to concept selection matrices is whether the selection criteria should be weighted, so that a favorable result for a very important criterion carries more weight than a favorable result for a less important criterion (Figure 3). Stuart Pugh stated that weightings "... imply an attempt to impart to the procedure too much precision, and thus inhibit qualitative judgments" [1]. Furthermore, the determination of the relative weightings can involve considerable time for debates among the team, and can involve attempts to bias the weightings to favor a concept.

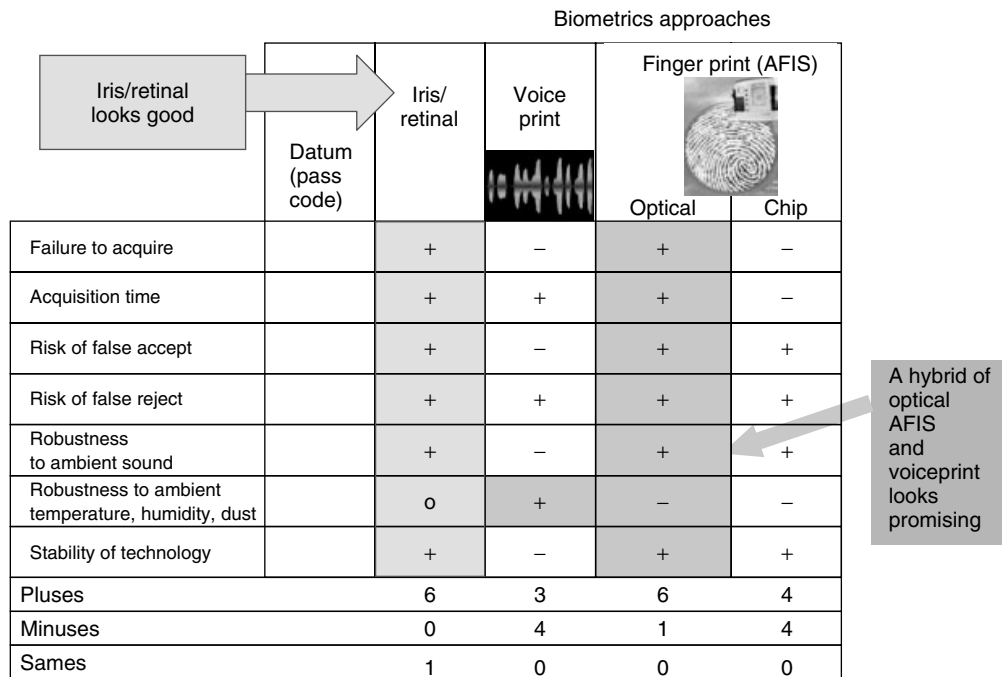


Figure 2 Pugh concept selection matrix for selecting a technology for security, verification of a person's identity. The datum is the current approach of entering a password or pass code

Criteria	Weight	Datum	Alternative 1	Alternative 2	Alternative 3	Alternative 4	Alternative 5
	1 ▲		+ ▲	+ ▲	+ ▲	+ ▲	+ ▲
	2 ▲		s ▲	s ▲	s ▲	s ▲	s ▲
	3 ▲		– ▲	– ▲	– ▲	– ▲	– ▲
	4 ▼		– ▼	– ▼	– ▼	– ▼	– ▼
	1 ▲		+ ▲	+ ▲	+ ▲	+ ▲	+ ▲
	2 ▲		s ▲	s ▲	s ▲	s ▲	s ▲
	3 ▲		– ▲	– ▲	– ▲	– ▲	– ▲
	4 ▼		– ▼	– ▼	– ▼	– ▼	– ▼
	1 ▲		+ ▲	+ ▲	+ ▲	+ ▲	+ ▲
	2 ▲		s ▲	s ▲	s ▲	s ▲	s ▲
	3 ▲		– ▲	– ▲	– ▲	– ▲	– ▲
	4 ▼		– ▼	– ▼	– ▼	– ▼	– ▼
Weighted sum of pluses			0	0	0	0	0
Weighted sum of minuses			0	0	0	0	0
Weighted sum of sames			0	0	0	0	0

Figure 3 Weighted version of concept selection matrix, with weighting for relative importance of each criterion

4 Concept Selection Matrices

Rather than weighting, Stuart Pugh suggested that “what has been found useful is the grading of criteria into three categories – important, moderately important, and desirable, or (high, medium, low).”

Reference

- [1] Stuart P. (1991). *Total Design – Integrated Methods for Successful Product Engineering*, Addison-Wesley Publishing Company, Chapter 4.

ERIC CHARLES MAASS

Process Maps, Construction of

What is Process Mapping

A process is a series of steps that converts an input(s) to an output(s). This series of steps is usually done repeatedly. As the familiar proverb says “A picture is worth a thousand words”, so a good way to explain these steps is pictorially. Therefore, process mapping is a key component in any quality effort to control, improve, document, and streamline a process. Process mapping can take on many forms but the commonality is that the specific steps or events that are required to provide a service or product to a customer are listed. The customer could be an external paying customer or an internal customer.

Process mapping can be called by several names as well as have several appearances. A generic process map may be called a *process flowchart*, *flowchart diagram*, or a *business map*. Specific maps can take on special names like SIPOC (suppliers–inputs–process–outputs–customer), functional deployment map (FDM), or opportunity map.

Reasons for Creating a Process Map

There are many reasons to create a process map. Four major reasons are to document the current state, to find opportunities for improvement in the current process, to document a change in the process, or to create a new process.

Many quality systems, such as ISO 9000, require documentation of key processes in the organization. Many companies also have their own internal requirements for process documentation. This documentation can be helpful to training employees or to get a common understanding.

One of the most common uses for maps is to identify changes to the process that would increase efficiencies, reduce cycle times, reduce costs, or decrease defects. Improvement teams often find steps that are redundant, unnecessary, and do not add value to the product or service. In the **Six Sigma** methodology define–measure–analyze–improve–control (DMAIC), mapping is used in the define phase to look for quick wins. Quick wins are improvement

ideas that can be implemented easily and quickly. In the measure phase, maps are used to identify points in the process to collect data to look for more complicated or less obvious improvements. In the improve phase, they are used to demonstrate the proposed new, improved process.

Process mapping has also been used as part of the groundwork before beginning to develop software programs to support the process. Software developers have long used mapping as a tool to help define the requirements for code. In addition, mapping is a precursor to the use of process simulation. Process simulation is the technique of creating a computer model of the process.

Process Map Formats

Process maps may take on many different looks. Some of the common ones are SIPOC, detailed process map, an FDM, and an opportunity map. Each map has its own purpose and therefore an appropriate time to use it.

SIPOC Map

A process map that has gained popularity with the increase of Six Sigma training is the SIPOC map. SIPOC is an acronym for suppliers–inputs–process–outputs–customers. The SIPOC is a very high-level map that does not use traditional mapping symbols. The first column of the SIPOC lists the key suppliers to the process. These may be internal or external suppliers. The second column lists the key inputs to the process. The inputs are usually those components that are transformed during the process to create the output. The third column is the process column where the major process steps are cataloged. Since this is a top-level map, there usually are only a few major steps listed. The fourth column is where the outputs of the process are listed. The last column would list the main customers or customer types of the outputs. Again, the customer could be an external or an internal customer.

Although the SIPOC is listed and displayed in this order, this does not imply that the SIPOC would be created in this order. Since there has been an increase in the focus on the process from the customers and their requirements, it makes sense to start mapping the process from the customer’s perspective. So, the

2 Process Maps, Construction of

mapping may start with a focus on who the customers are and what are the outputs of interest. To place this emphasis on the customer, there is a similar process map called *COPIS*. This is the SIPOC listed in the opposite way.

The SIPOC is good as a starting point for the team before adding additional detail. It is also good as a display to upper management to get a general sense of the process. It can help define who should be on the team for the more detailed mapping to follow.

An example of a SIPOC of purchasing a plane ticket is shown in Table 1. In this example, the customer could be either the traveler or a travel agent. The traveler is also a supplier in that she/he must supply information about travel requirements. The airlines also are supplier since they provide schedule information. Usually, the outputs and inputs listed are the few key ones. In this case, the key output would be the purchased ticket. The major steps of the process, listed without worrying about decisions or feedback loops, are as follows: entering information in computer system, review of flight schedules, flight selection, entering credit information, and then purchasing the ticket.

Simple Detailed Process Map

A simple detailed process map is usually done to a deeper level than a SIPOC. The detailed process map utilizes symbols to explain what happens in the process. The detailed process map is usually displayed to show the easiest flow going right to left. Many different symbols are used but there are some basic symbols that are traditionally used. The first symbol is the rectangle. The rectangle is used to indicate an activity or operation. The diamond is used to illustrate a decision point. Arrows indicate the flow of process. Ovals indicate the start and stop of the process. Circles are connectors between pages. Table 2 shows these symbols, the use of the symbols, and a simple example of what the symbols might represent.

The detailed process map is useful for finding areas for improvement, understanding redundancies in the process, to find waste, and to identify areas for collecting data to monitor the process.

An example of a simple detailed map of a customer service representative answering a customer's question is illustrated in Figure 1.

Table 1 SIPOC of purchasing a plane ticket

Suppliers	Inputs	Process	Outputs	Customers
Traveler Airlines	Traveler requirements Flight schedules	Enter information Review flight schedules Select flight Take credit information Purchase ticket	Ticket issued	Traveler Travel agent

Table 2 Traditionally used symbols for detailed process mapping including usage and a simple example

Symbol	Usage	Example
	Activity or operation	Assemble parts
	Decision	Request fulfilled?
	Flow	Go to next step
	Start or stop of process to be mapped	Receive customer request
	Connector to another page	Go to page 2

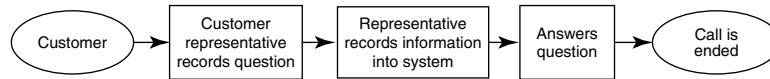


Figure 1 Detailed process map for the process of a customer service representative answering a customer's question

Besides indicating activities, decisions, and process flow, detailed process maps may include a myriad of other information. A partial list of this additional information is as follows:

- the number of people required for each step
- the machinery required at each step
- the tooling required at each step
- special inputs required
- information required
- the time to conduct the step
- the cost of performing the step
- the defects associated with specific points in the process
- delays
- queues
- setup times required at each step.

Beyond this list, there are two other types of information that can be added to a map that will make two specialized maps: a functional deployment map, also called a *swimming lanes map*, and an *opportunity map*.

Functional Deployment Map

An FDM takes a traditional detailed process map and adds an extra dimension by listing the functions that actually perform the process step. In an FDM, you would establish columns for each function, department, or person involved in the process. Then each activity or decision is placed in the associated column of the function that performs that step. The purpose of the map would be to identify the key communication or transfer points. These points are often where queues or defects occur. These maps are most useful in environments where the process involves many groups.

A very simple FDM for the process of a customer representative handling a customer call is shown in Figure 2. Three different groups are involved in this process: the customer, the customer service representative, and the engineer who handles the more difficult calls.

Opportunity Map

An opportunity map is a simple detailed process map with two additional columns. The two columns indicate whether a step in the process is either value added or non-value added. The purpose of the opportunity map is to identify opportunities to eliminate or reduce waste in the process. A value-added step is considered to be a step that is inherent to converting the inputs into the product or service. Non-value-added steps are associated with waste and can include, but are not limited to inspection, audits, queues, transportation of product, checking work, and so on. These activities just increase costs and should be primary targets of improvement.

An opportunity map for a very simple assembly process is shown in Figure 3. In this process, any assembly operation would be considered value added. These steps are turning the product into what the customer expects. The inspection and transportation steps are considered non-value added, or wasteful.

Creating the Process Map

Most process mapping sessions involve the participation of a team of people. Often, if the mapping session is designed to capture the current state of the process, the team finds differing views about what the real steps are. This is often a significant output from process mapping, getting the team to see the process from all points of view.

When doing mapping, there are some basic steps that the team should follow. These steps are as follows:

1. Establish the process to be mapped. Give the process a name.
2. Determine the objective of the mapping exercise – how will the map be used. This will influence your choice of the map that you create.
3. Invite the correct people to participate. Often it is helpful to include people who do not know

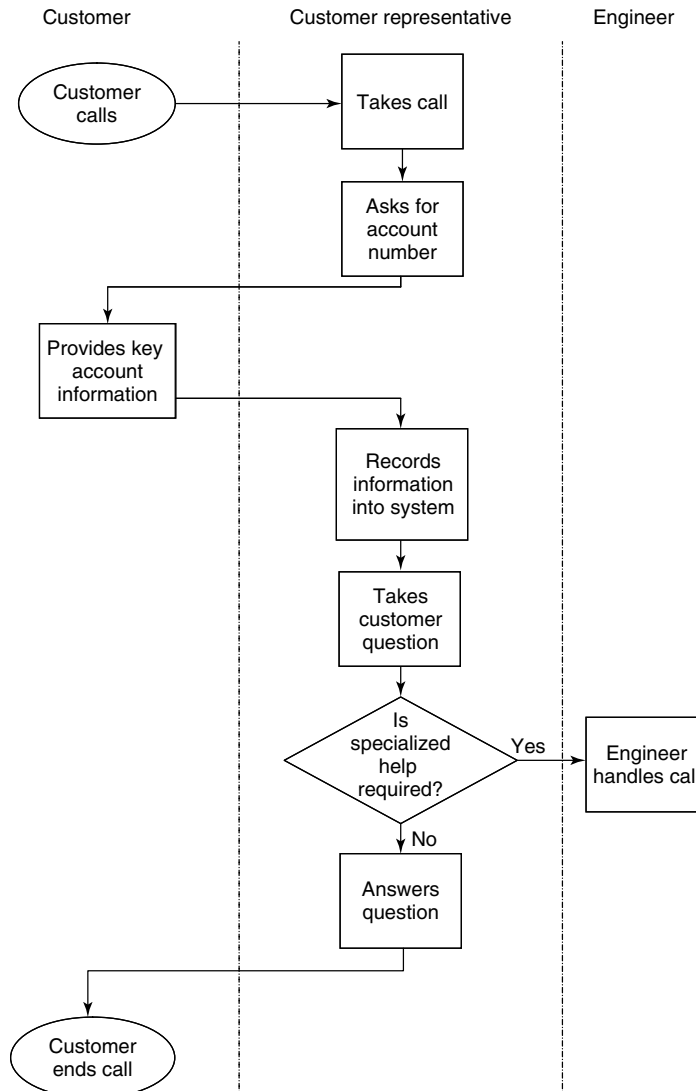


Figure 2 Functional deployment map for the process of a customer service representative handling a customer's call

- the process. These people may ask questions that people close to the process may not think to ask.
4. Define the boundaries to be mapped. These can change as the team finds that issues of concern may be influenced by upstream or downstream activities. The people involved in the mapping exercise may change as the boundaries are altered.
 5. Observe the process as it is operating. Make sure the team captures the information while observing.

6. Agree on the major steps and the order that they occur.
7. Add detail to the major steps as necessary.
8. Document in the appropriate format.
9. Validate the map with other people in the process.

If process mapping was part of a continuous improvement effort, after this “as-is” map is established, the team would start to question the status quo. Each step in the process should be reviewed to look for opportunities to eliminate unnecessary

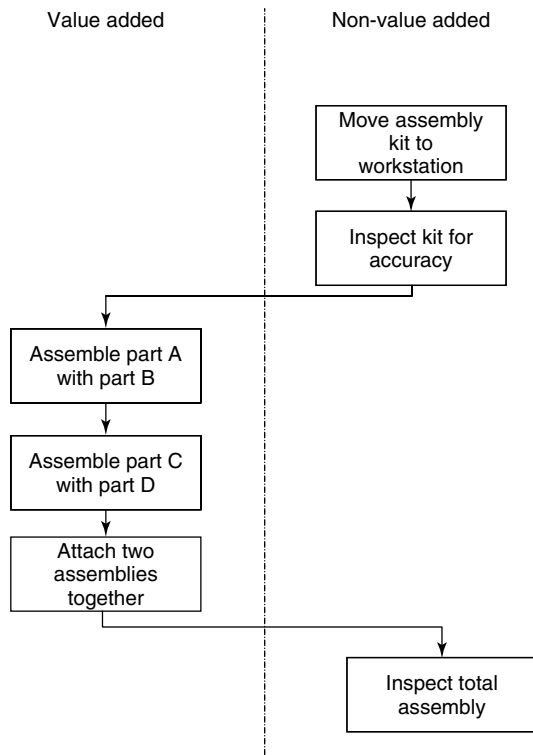


Figure 3 Opportunity map for an assembly process

or redundant steps, combine steps to improve efficiencies, or rearrange steps to facilitate a cleaner, simpler flow. The team would then create the “to-be” process map based on the changes agreed upon.

Best Practices in Process Mapping

There are some key practices that will facilitate the mapping process: involving the right people, making the map understandable, and being clear in the purpose of the map.

If a map is created by one person, the perspective may not really reflect reality. Often, there is someone who originally defined the process. Their perspective may be the wished-for state. However, the people in the process may have intentionally or unintentionally changed the defined steps. One reason may be that the process users have a misunderstanding of what the process should look like. They might not have received proper training or the

steps were not properly communicated. Another reason might be that the process users discovered a better or easier way for the process to work. The original steps may not have covered all possible contingencies.

The process may also have been a change from previous steps and the process users did not see the benefit of the change, so they reverted back to the original steps.

For all these reasons, it is important to include the process users in the creation of the process map, to capture the true state of the process. Other people who might be included in the formation of the map are the suppliers, customers, and process owners.

Since there are many different kinds of process maps, and different levels of detail, it is imperative that the creator(s) specify up front what the purpose of the map is, so that the proper format is used. Often, it makes sense to start at a very top level and add complexity as required. Knowing the purpose will also help clarify the types of information to be added.

Difficulties in Process Mapping

One of the most difficult aspects of mapping is the level of detail required. Too much detail can waste time and will not be productive. Too little detail can prevent a deep enough understanding of the process to see how to improve or to provide complete training. It is important to move from less detail to more detail. If the team tries to jump into too much detail in the beginning, they can get very bogged down in things that may be of little consequence. More detail can be added as deemed necessary.

Another difficulty of process mapping can be that the team members are in different locations. Most mapping sessions are done on a wall with a group of people using sticky notes. However, in today’s global environment, it is becoming more frequent for the mapping session to take place in an electronic environment. Electronic conference tools can help facilitate the mapping session but it can still add a level of difficulty.

To capture all the information about a process on a document can create an overly complicated document. Since one of the primary functions of process maps is as a communication tool, too much text

6 Process Maps, Construction of

can be daunting. It is important that this additional information be displayed in as easy-to-read format as possible.

Summary

Process maps are a key part of any quality program for any kind of organization. People tend to

understand pictures better than words and understanding the process is the best way to maintain and improve it. Although a process map might take on several looks or formats, the primary goal should be to keep it simple enough to be appealing to read and understandable to all people involved in the process.

LORRAINE DANIELS

Project Management, Stage-Gate Approach to

Introduction

The project process is a predetermined set of repeatable steps, or stages, that are used in sequence to plan and implement a variety of project types. The term *stage-gate* has been around for quite sometime and is defined for this article on project process as the criterion or question which must be appropriately addressed prior to proceeding to the next stage in the project.^a

There is an old saying in the project management world – “Don’t work any harder than your project sponsor – or you and your team will surely be frustrated.” Keeping the project sponsor fully engaged and working on the right things at the right times before, during, and after completion of the project is one of the key elements necessary for a project’s success. As a project manager, you might think they should already know this and it is really not your job to manage the sponsor and the project too. However, many project sponsors are not fully aware of the scope and importance of their role supporting the projects they sponsor and may not even be aware of the appropriate sequence of events necessary to perform the project successfully. As the project manager you may need to become an educator and coach the sponsor on how to effectively perform this key role. After all, why risk your project’s success if you can influence the outcome by managing the project sponsor along with the rest of the project?

Every project an organization implements, must be evaluated as an investment of resources to produce a benefit or add value to the organization. Accordingly, a good project manager knows the value of killing a project early to ensure that not a moment of effort is wasted on a project that has no merit. Most organizations are faced with more work than they have the time and resources to complete, so with limited resources even a day spent on a project that will ultimately be canceled, is a day that could have been spent more productively on another project delivering value for the investment being made. As a project manager you also may be thinking, but I do not want to be known as the *project killer* in

my organization, what is in it for me? Would it not be better to be leading projects that are hugely successful and adding great value to the organization? This is not necessarily possible if you are working on projects that ultimately will not be implemented, and if they are, do not deliver value to the organization at least several times the investment made in the project.

This article provides an overview of the disciplined approach to project process that engages the right resources at the right time, only when predetermined criteria are met at planned points along the project’s timeline. It also ties in the role a project sponsor plays, at each phase in the project process, so the project can proceed in the most efficient and expeditious manner. The project process model applies to a wide variety of projects from Six Sigma and lean projects to laboratory, manufacturing, building construction, and IT projects. Asking the hard questions and responding with honest answers will set the project up for success – or, in some cases, take it off the project list as quickly as possible.

Project Process Model Summary

The project process model is an end-to-end process comprised of four main phases and nine stages within these phases. The key premise is that all projects will be appropriately vetted and prioritized against other opportunities the organization has to invest its resources, and realize the intended benefits. To do this requires that all those involved in the project prioritization and selection effort follow the same disciplined approach for all projects.

The use of a documented process signals to everyone involved, the importance of the process. Any deviation from the process can be flagged and its potential impact on the project’s success understood. Table 1 shows the stages in each phase of the project process.

Each project moves forward in a sequence and prior to proceeding to the next stage in the project, a stage-gate criterion or question must be appropriately addressed. Each phase is described further below with an outline of the roles of key players in the project process.

Roles and Responsibilities

Individuals have specific roles and responsibilities assigned during a project because of their position

2 Project Management, Stage-Gate Approach to

Table 1 Project process phases

Planning phase	Solution design phase	Implementation phase	Operational startup phase
<i>Stage 1</i> Inception <i>Stage 2</i> Feasibility	<i>Stage 3</i> Concept design <i>Stage 4</i> Scheme design <i>Stage 5</i> Detailed design	<i>Stage 6</i> Performance <i>Stage 7</i> Commissioning and Handover	<i>Stage 8</i> Process integration <i>Stage 9</i> Closeout and Benefits Realization

and level of influence within the organization or on a specific project. At a minimum, there are four key roles that need to be established for each project: project sponsor, project manager (or project director), project team, and stakeholders.

Project Sponsor

The project sponsor is accountable for the realization of the project's business benefits and has macrolevel scope control. The project sponsor also ensures the project is appropriately resourced to achieve the intended results and undertakes regular project reviews and makes necessary adjustments. The project sponsor may also hold the executive role as well. Staying connected to the executive or top governing level in the organization is necessary to ensure that the project has the appropriate level of visibility and support in the organization to be successful. If the project sponsor does not hold the executive role, he will need to periodically inform the executive of the progress the project is making against the intended objectives.

Project Manager (or Project Director)

The project manager (or project director) is responsible for leading the planning and execution of a project in a manner consistent with the needs of the organization and the directives of the project sponsor. He is also responsible for providing feedback on the project process in support of continuous quality improvement.

Project Team

The project team is responsible for planning and executing the project in a manner consistent with the needs of the organization, the directives of the project manager, and the approved project plan. They are also responsible for looking for ways of improving the

project process and providing feedback in support of continuous quality improvement.

Stakeholders

Stakeholders are individuals and groups who have a stake in the outcome of the project and whose interests must be recognized for the project to be successful. Customers are also considered stakeholders and are members of the project team.

Project Process Phases

Planning Phase

The time spent on planning a project returns many dividends over the life of the project. The planning phase has two stages, Stage 1 – Inception and Stage 2 – Feasibility.

Stage 1 – Inception. The inception stage initiates the project. The project sponsor, working with the senior leadership, determines what is needed by the organization to meet its strategic business objectives and assigns the appropriate resources to begin the inception stage. The project manager should be assigned to work directly with the project sponsor to gather data and begin the analysis to determine the nature of the project opportunity. The initial focus is on asking the right questions to understand why the project is needed and how the project fulfills these needs. At this stage, the unmet business needs will be identified by the customer and strategic alternatives that meet the business needs, generated. This will be summarized as business objectives.

Activities in this stage include determining the nature of the project or business opportunity; identifying strategic alternatives; evaluating potential risks and benefits; estimating costs to a Class I level $\pm 50\%$; and obtaining preliminary funding. The key

deliverable in this stage is a statement of business objectives, which answers the stage-gate question “Is the project needed?”

Stage 2 – Feasibility. In the feasibility stage, the anticipated business benefits gained by executing the project will be identified and quantified where possible. Further, the project description will be written to include how it directly supports the business strategy. The feasibility stage tests the fundamental economic soundness, and business and technical viability of the project. Based on the business needs and anticipated benefits of the project, detailed project objectives will be written.

At this point a clearly defined and agreed scope of work needs to be written, based on the project objectives approved by the project sponsor. Scope definition ensures there are clear boundaries and definitions of processes, facilities, and systems to be delivered. Scope definition is a key component for developing risk analysis, funding requests, and deciding on project controls strategy. Scope will define quality, quantity, and functionality of the end result of the project. Scope will also define work specifically excluded from the project where overlapping activities may occur.

The activities of this stage include defining the project’s objective and purpose; defining the general scope of work; addressing the project’s impact in relation to facilities, processes, software, and systems; and estimating capital and operational costs to a Class II level $\pm 20\text{--}25\%$. The conclusion of this stage answers the stage-gate question “Is there a business benefit?” prior to proceeding to the next stage.

Solution Design Phase

The solution design phase initiates project design activities based on the upfront planning efforts. In this phase the project team usually expands to include the designer and any subject matter experts necessary to complete the design in a manner consistent with the business and project objectives, and scope of work. Additional project team members may come from both internal and external sources.

During this phase the risk of the scope expanding beyond the original intent is the greatest, and the project sponsor’s role is to ensure the design work proceeds according to the agreed scope of work. Any changes to the scope of work places the project at risk. The project sponsor’s role during this phase

also includes reviewing the complement of project team members and making adjustments as needed, supporting the project in clearing funding hurdles, and keeping the executive leadership updated on the project’s progress and informed of any potential risks to the project success.

Stage 3 – Concept Design. The basis of design is written at the beginning of this stage. Sometimes called the *design program*, it is the detailed listing of specifications and design requirements written prior to beginning the solution design. Preliminary design options and the associated costs of each option are then generated and reviewed to determine which designs to pursue further. The project team will seek input from key stakeholders during the review and selection of design options. At the conclusion of this stage a single design option will be selected as the basis of the scheme design stage.

The activities of this stage include developing preferred designs with associated costs; selecting a single option as the basis of the scheme design stage; preparing the project execution plan; and estimating costs to the Class III level $\pm 15\text{--}20\%$. The conclusion of this stage answers the stage-gate question “Is this the right design option?” prior to proceeding to the next stage.

Stage 4 – Scheme Design. During the scheme design stage the designer continues design of the selected option and produces design documentation to 20–35% complete. The project team will involve and seek input from key stakeholders at key points as the scheme design efforts and costs progress. The designer provides other information needed to support a full funding request with cost estimates taken to a Class IV level $\pm 10\%$. The conclusion of this stage answers the stage-gate question “Should the organization commit full funding to the project?”

Stage 5 – Detailed Design. During the detailed design stage, design documents including flow diagrams, calculations, drawings, specifications, and other information are completed and bidding or pricing packages are produced for the procurement of equipment and services needed during the implementation phase. The project team will continue to involve stakeholders as the detailed design efforts are accomplished. Cost estimates are reviewed at the Class IV level $\pm 10\%$. The conclusion of this stage

4 Project Management, Stage-Gate Approach to

answers several stage-gate questions “Are these the right services providers, contractors and suppliers? Are these the right bidding or pricing packages?”

Implementation Phase

During this phase changes to processes, systems, software, and facilities are initiated and completed, and the resulting product is transferred to the operating group. In many cases, this is the most complex and time-consuming portion of the project, requiring constant communication among the project team members, the stakeholders directly affected by all the changes, and those performing the work activities.

Stage 6 – Performance. The performance stage is where changes to the processes, facilities, equipment, software, and systems are underway and completed. The progress and quality of the work during this stage will be tracked according to the design documents, time schedules, and approved costs. Any variances will be tracked and a strategy developed to return to the planned activities.

Activities concluding this stage are focused around review and acceptance of completed and substantially completed processes, facilities, software, equipment, and systems by the customer. The conclusion to this stage answers the stage-gate question: “Is the designed change complete?”

Stage 7 – Commissioning and Handover. Commissioning is the testing of the change (facilities, equipment, software, processes, and systems) to determine if they function as designed. It also includes the documenting of the test data supporting the findings. The commissioning effort begins with a commissioning plan which includes a schedule of activities and is integrated with the overall performance stage schedule.

Handover activities during this stage include training customer operations staff to use the new or changed systems, and then systematically transferring responsibility from the project team to the customer. The conclusion of this stage answers these stage-gate questions “Does the facility, process, software, or system function as designed? Does the business/customer accept operational responsibility?”

Operational Startup

This final phase of the project process has two stages which provide evidence that the designed

and implemented changes operate in a consistent manner and that the organization realizes the benefits promised at the beginning of the project, based on its investment in the project.

Stage 8 – Process Integration. The process integration stage, also called *validation in highly regulated environments*, is the implementation of rigorous activities which provide documented evidence that the specific process, facility, software, or system consistently produces results within the design specifications; and ensures product compliance with regulatory standards and full integration with other processes. The conclusion of this stage answers the stage-gate question “Does the facility, process, software, or system produce consistent results?”

Stage 9 – Closeout and Benefits Realization. Closeout and benefits realization stage concludes the project process with activities including final adjustments and repairs, record drawings, contract closeout obligations, financial accounting, customer satisfaction surveys, operation and maintenance lessons learned, and documented evidence that the intended business benefits were realized. This final stage concludes by answering the stage-gate question: “Are the anticipated benefits being realized?” prior to final project closeout.

Summary

Viewing each project as an investment made by your organization to achieve its strategic objectives, underscores the need to gain a significant return on that investment in the form of capability, flexibility, speed, cost, and competitive advantage in the industry. By applying the same disciplined approach to each project, your organization establishes ways of working that sets the organization up to reliably produce projects and have the greatest likelihood of gaining the return on investment promised at the outset.

Keeping the project sponsor connected to the process and aware of the role he plays during each phase, brings visibility of the importance of the project to all levels of the organization and contributes to the likelihood of having the necessary resources to ensure a successful outcome.

End Note

^a. The origin of the stage-gate methodology is based on the experiences, suggestions, and observations of a large number of managers and firms in over 60 cases as observed by Robert Cooper. The term *stage-gate* first appeared in an article by

Cooper in the *Journal of Marketing Management*, 3, 3, Spring 1988. An even earlier version can be found in Cooper's book: *Winning at New Products*, 1988 (http://www.12manage.com/methods_cooper-stage-gate.html).

THOMAS F. TRABERT

Team Building

Strategic Team Alignment

Team building is being used very successfully in different environments including technical Six Sigma project teams. Historically, teams have been used in every situation and program imaginable – process improvement teams, continuous improvement teams, total quality teams, continuous quality improvement teams, and many more. Some have been more successful than others. Today, teams and team building are working in companies and on very technically focused projects. Our work at North Carolina State University focuses on **Six Sigma** teams; we are seeing some amazing successes. Success may be linked to a slight variation in the team approach. The teams and team charters are intentionally linked to the strategic plan, to the vision and mission of the organization with meaningful metrics in place, have the buy-in of leadership, and can show financial gains or ROI within the first six months to one year. With these key factors in place, we are seeing a dramatic difference in how successful teams work. For possibly the first time at many companies, we are intentionally linking our team project charters to the organization's strategic plans. This can be attributed to the scientific or methodical approach of Six Sigma training and implementation. In order to ensure that our projects will have a measurable impact, and that they are important enough to the organization to dedicate critical resources, i.e. time and money on, we must ensure that our leadership can draw a line directly from the project to the strategic plan. While serving as a project champion early in the define phase for a service company, we were able to clearly link their black belt scheduling project to their strategic plan. This action enabled the team to focus on the bigger picture – the strategic objectives of the organizations – and verbalize how this team fit into the big picture. This created buy-in and focus for the team, as well as accountability. In terms of financial return on the projects, the best measurement we have so far is a form completed by each company on every project to report the financial gains; this is verified by their CFO. Completion of the data is required at the time of project certification. As one example, a service company reported over \$1 000 000 annual reduction in cost on a single black belt project. Most

savings range from \$25 K to \$50 K per project on average, and focus on reduction in cost or time, or increase in quality or capacity. The average size of a team varies from five to eight members; we have recorded teams as large as 12 members.

Team Outperform Individuals

According to Katzenbach and Smith [1] in *The Wisdom of Teams*, teams working together will outperform individuals acting alone. It is not a natural state to work in teams; we prefer to operate as individuals and do not readily take the responsibility for the dynamic work of others – even when part of an established, working group [1]. This is verified by a quick review of most reward systems; they are established to compensate us individually and not as teams. When working with others collaboratively on a project, we can frequently outperform individuals simply by assigning tasks and breaking down the project into workable pieces to get the work done. I conducted some informal research on this factor. During our Six Sigma green belt training classes, we conduct a hands-on team building training strategy. It is a group activity that compares individual behavior to group behavior. When analyzing results over a 5-year period in multiple industries, over 90% of the results validated that the team outperformed the individual acting alone during the staged team building activity. The fact that some individuals were reluctant to participate was also verified, indicating a nonnatural performance state, while some individuals scores did outperform the group results, although a very low percentage (<0.1) per team per training class.

Team Work in Industry

The success of teams and team projects in a technically focused environment is not tied to any specific industry. According to our records at the university collected over the last 5 years, 75% originates from manufacturing-based industries, while 25% comes from nonmanufacturing (service and health-care/bio/pharmaceutical based) industries. There is no clear distinction of particular team success tied to any particular industry. The team building component is more naturally formed in an industry that is accustomed to working in natural teams, such as manufacturing, and unlike healthcare. According to

2 Team Building

our records, healthcare is experiencing success with projects in Six Sigma; the cross-functional project teams may be more of a challenge in that industry owing to the individual nature of the work [2].

Project Team Roles

The teams on our technical projects have some similarities. The members are all viewed as problem solvers. Since we instruct our black belts at a more advanced technical level, they are relied upon for their statistical analysis and are frequently called upon to be team leaders. Green belts are extraordinary problem solvers and can lead tactical teams, or can be utilized to seek out solutions to a segmented project, or a smaller part of a very technical project. The necessary leadership depends upon the nature of the project and the available resources. While working with a service industry on a technical project, we observed that the black belt served as the leader and analyst, with green belts as members and **data collection** specialists. Master black belts are trained to “train the trainer”. The master candidates may be utilized as internal trainers, as a money saving tactic to continue training the organization’s internal resources. This adds to the return on investment for the overall strategy of the company’s Six Sigma implementation plan.

Optimizing Future Resources

The plan for the future of these teams looks very positive. We are seeing an increase in Six Sigma continuing education course enrollment at our university. Our percentage of completed projects for all industries is approximately 20% for both green and black belt projects, with enrollment in green belt classes nearly double the black belts. This may be indicative that training is more meaningful than project implementation, although the Return on Investment (ROI) is measured in projects completed, but there is no conclusive data. We are seeing very few companies with dedicated resources in green or black belts; they continue to incorporate their current positions into their project work. When current black belt team members in class are asked for impediments to their job, the frequent answer is too many assignments and not enough time to complete the project work. We are allowing our best resources to be trained

at a very high level, but we are not relieving them of their usual and customary duties, as reported in team building sessions in class. In reality, many of our attendees report they are a one-person team on many projects. This may impact future problem solving capabilities and impact turnover. In one major medical care facility, we know of dedicated resources that are operating at the highest level with optimized results. How resources are leveled and optimized should be a consideration in the future, as reward and recognition continue to be a factor for high-level project teams. According to Juran and Godfrey [3], high performing project teams are difficult to build, but when maintained, can help lead the company to a good competitive advantage [3].

Summary

While technical teams aligned to a systematic approach such as Six Sigma are not the only successful teams on record, there have been some very good results with our industry partners at our university. This is not to imply that teams are the only pathways to resounding success, as many of our companies are successful without the existence of ongoing teams. For more information on industry results, please visit our web site at <http://www.sixsigma.tx.ncsu.edu>. For additional information, please see other suggested resources.

References

- [1] Katzenbach, J.R. & Smith, D.K. (2003). *The Wisdom of Teams, Creating the High-Performance Organization*, Harper Business Essentials, New York, pp. 9–26.
- [2] Institute of Medicine, Committee on Quality of Health Care in America. (2001). *Crossing the Quality Chasm, A New Health System for the 21st Century*, National Academy Press, Washington, DC, pp. 130–140.
- [3] Juran, J.M. & Godfrey, A.B. (1999). *Juran's Quality Handbook*, 5th Edition, McGraw-Hill, New York, pp. 15.10–15.28.

Further Reading

- Bens, I. (2000). *Facilitation at a Glance! A Pocket Guide of Tools and Techniques for Effective [Team] Meeting Facilitation*, Goal/QPC and AQP, Salem; Jossey-Bass, San Francisco.
- Federico, M. & Beaty, R. (2003). *Rath & Strong's Six Sigma Team Pocket Guide*, McGraw-Hill, New York.

Lencioni, P. (2005). *Overcoming the Five Dysfunctions of a Team: A Field Guide For Leaders, Managers, and Facilitators.*, Jossey-Bass, San Francisco.

Pande, P.S., Neuman, R.P. & Cavanagh, R.R. (2002). *The Six Sigma Way Team Fieldbook: An Implementation Guide for Process Improvement Teams*, McGraw-Hill, New York.

Scholtes, P.R., Joiner, B.L. & Streibel, B.J. (2003). *The Team Handbook*, 3rd Edition, Oriel Incorporated, Madison.

MARGARET G. O'BRIEN

Multi-Vari Charts

What are Multi-Vari Charts?

Multi-vari charts are a graphical form of analysis of variance. You can use them most effectively for problem solving. However, you can also use them to estimate variation in well-behaved processes.

Len Seder

Leonard Seder (1915–2004) invented multi-vari charts in the late 1940s. He published a two-part article “Diagnosis with Diagrams” in *Industrial Quality Control* in 1950 (January and March). Len gave credit to Dr Joseph M. Juran for the concept of graphing process variability in a manner similar to a stock report with daily high, low, and closing prices represented by a vertical line for the range and a point for the closing. Len described multi-vari charts as a diagnostic tool, and compared them to medical diagnostic tools like X-rays or electrocardiograms. He stressed the power of an effective graphical presentation. “The question is, then, how can these powerful techniques be made more available? One answer is to be found, in the author’s opinion, in the use of graphics in place of statistics. Much of the success of the simpler statistical quality control techniques is undeniably attributable to the forcefulness and conciseness of graphic presentation.”

Why do Multi-Vari Charts Work?

A well-planned and executed multi-vari chart will separate overall process variation into families. You will clearly see the family causing most of the change and patterns of variation within the largest family, which will provide further clues about the source of variation. The power of a multi-vari chart comes from the application of the **Pareto** principle to the sources of variation in a process.

Dorian Shainin

Dorian Shainin (1914–2000) attended MIT with Len Seder. The two men established a professional relationship with Dr Juran in the 1940s. At that time,

Dorian was in charge of quality and reliability at Hamilton Standard (a division of United Aircraft) and Len was with Gillette Safety Razor Company in Boston. Dorian was an early adopter of Len’s multi-vari charts. He was also an admirer of Dr Juran’s application of the Pareto principle to project selection.

In the 1950s, Dorian concluded that just as the Pareto principle applied to the effect individual problems were having on business performance, it must apply to the effect that individual variables had on changes in product or process performance. If a problem could be expressed in terms of variation, i.e., too much change in the values of product or process outputs, and then the Pareto principle requires that one cause must be responsible for most of the change. Dorian called this dominant cause the *big Red X*. The Red X paradigm does not suggest that the Red X is the only source of variation. In fact, it recognizes that there are thousands of potential causes. The Red X causes most of the variation. Find it and control it and you will make a significant improvement in performance.

Red X Paradigm

The Red X paradigm changes the way problems are attacked. If there is a dominant cause of variation, finding it and controlling it is the only path to improvement. The statistical law underlying analysis of variance states that independent sources of variation combine as the square root of the sum of the squares. If one source causes a range of 5 units and another independent source is causing 1 unit, the combined range will be 5.1 units (the square root of 26), not six (the sum of 5 + 1). The 5-unit cause is the Red X and it is 25 times stronger than the cause contributing 1 unit.

Once you focus on finding the Red X, the priority becomes finding it quickly, with minimal resources. The best approach is a progressive search, using a process of elimination. Eliminate everything that cannot be the Red X and whatever is left must be the Red X. This is a classic approach for solving criminal cases. Sherlock Holmes expressed it well in Sir Arthur Conan Doyle’s “*The Sign of the Four*”: “Watson, how often have I told you? When you eliminate the impossible, whatever remains, however improbable must be the truth.” Effective methods converge rapidly on the true root cause by eliminating the impossible.

2 Multi-Vari Charts

Dictionary Game Thinking

Dorian illustrated the concept with a parlor game called the *dictionary game*. A secret word is chosen. You may only ask questions that can be answered with a “yes” or “no”. And you are given a dictionary as a tool to help discover the word. Uninitiated contestants often use a 20-questions approach. The initial common questions are as follows:

- “Is it a person?”
- “Is it a place?”
- “Is it a thing?”

Since we are looking for a word, additional questions might be as follows:

- “Does it have more than five letters?”
- “Is it a noun?”
- “Does it have more than two syllables?”

The keys to the dictionary game are listed below:

1. Use a process of elimination.
2. Take advantage of the dictionary’s structure (it is alphabetical).
3. Start with broad questions (eliminate a lot).
4. Narrow the questions as the search progresses.
5. Ask questions that will narrow down the search no matter the answer.

Here is an example. Let us choose “desktop” as the secret word. Start the search by finding the center of the dictionary by page number. In my dictionary, the first “a” word is on page 43 and the last “z” word is on page 1373. The midpoint is page 708. The first word on page 709 is “lucky”. Asking, “does the word come before ‘lucky’ ”, will eliminate half the possible words. It does not matter if the answer is “yes” or “no”. In our example, the correct answer is “yes”. The word “lucky” and all words following it have been eliminated.

The next step is to divide the remaining section in half and repeat the process. The next midpoint is on page 375. The first word on page 375 is “domino”. The question is now: “Does your word come before ‘domino’?” Again the answer is yes and another quarter of the dictionary has been eliminated.

Within 10 questions, the search will be narrowed to a single page. In my dictionary, “desktop” is on page 344. The pages are divided into two columns. Applying the question to the word atop of column

two, “despair”, narrows the search to the first column. There are 26 words in the column. Word 13 is “desk”. The question is now: “Does your word come before ‘desk’?” The answer is no.

Within 17 questions, we would have eliminated all the words in the dictionary, except “desktop”. The first question eliminated half the words; the second a quarter; and the last question only eliminated one. Notice that even though the question form stays the same, the questions start broad and become narrower as the search progresses.

Dictionary Game Thinking with Multi-Vari Charts

A multi-vari chart is an important tool for finding a process Red X. By separating the overall process variation into families of variation, all but the largest family can be eliminated. The search for the Red X is quickly narrowed to a small part of the process.

The Case of the Incapable Lathe

In 1966, Dorian published an article in Industrial Quality Control. It was titled “The Case of the Incapable Lathe” and it illustrated the power of multi-vari charts in finding a Red X.

A manufacturer was making jet engine fuel pump controllers. A key component was a rotor shaft. Shaft diameter was an important dimension. The shafts were turned on a lathe. A capability study determined that the process was not capable. Management put pressure on manufacturing to fix the problem. After failing to get tolerance relief from engineering, manufacturing assigned a team of experts to examine the problem and recommend a solution. The team concluded that the lathe was old and worn out and simply not capable of maintaining the tolerance of ± 0.001 in (Figures 1 and 2). They believed the lathe had worn bushings and spindles and needed to be replaced. They recommended that management purchase a new lathe. Management was not convinced and asked Dorian to investigate.

The process produced shafts on the lathe. Finished shafts were deposited in a basket. After cleaning to remove oil and metal chips, the parts were inspected for size. By the time the parts were measured, part orientation on the lathe and order of production were lost.

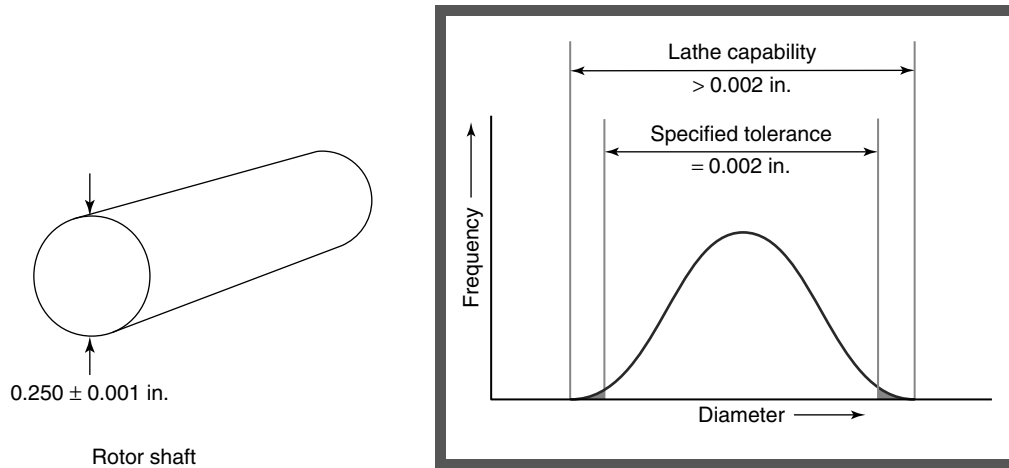


Figure 1 Tolerances of incapable of lathe

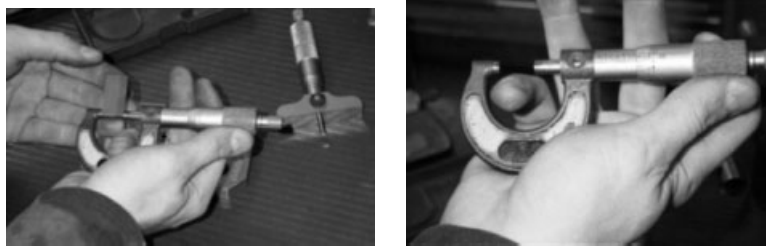


Figure 2 Measuring the diameter of the rotor shaft

Planning the Multi-Vari

To start with, identify the families of variation present within the process. Three broad families combine to produce the overall observed variation: within piece, piece-to-piece and time-to-time.

- Within piece describes deviations from the intended shape of the parts.
- Piece-to-piece captures process changes from one machine cycle to the next.
- Time-to-time captures process trends, shifts, or cycles.

Within Piece. Study the physical process the machine uses to create shape. Then consider how the shape might deviate from the design. For example, the lathe will rotate the workpiece about its central axis as the cutting tool moves in from the outside. As the cutting tool meets the workpiece a circular

shape is produced. The cutting tool then moves down the workpiece creating a cylinder. A change in the relationship between the workpiece center and the cutting tool during a rotational cycle will create an oval shape (out of round), instead of a circle. A change in the relationship between the workpiece center and the cutting tool as the tool travels down the length of the workpiece will produce a tapered shape.

To capture these potential changes, Dorian had the setup man capture parts as they came off the lathe. The setup man marked the end of the part that was held by the lathe chuck for future reference. See Figures 3, 4, and 5. He then found and measured the maximum and minimum diameters at each end of the part (chuck and tail). Four readings represented the range of diameter values within each part.

Piece-to-Piece. Study the organization of the manufacturing process to find piece-to-piece families. If multiple lathes produced the rotor shafts, then

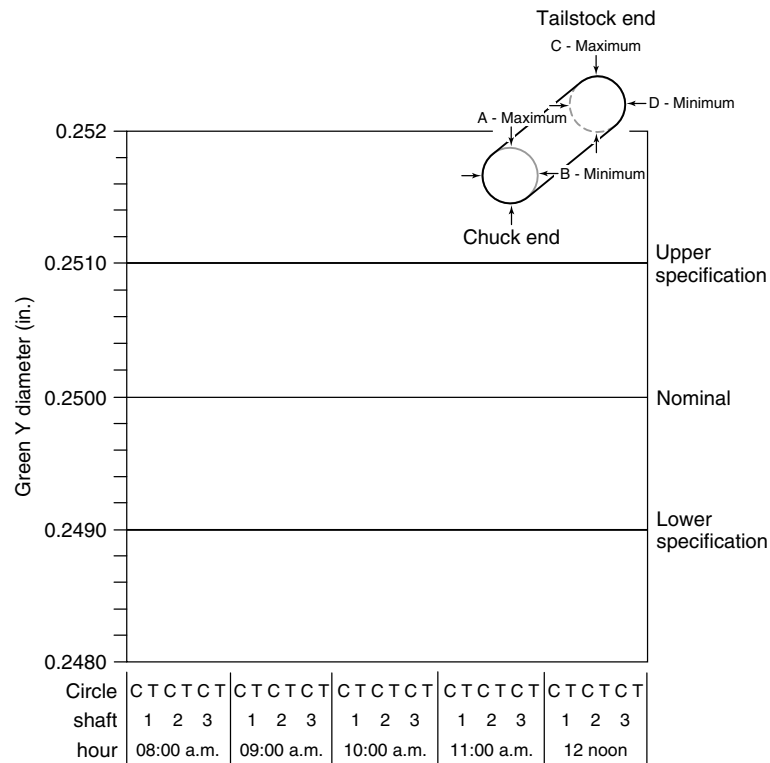


Figure 3 Components of variation for rotor shafts. C: chuck; T: tailstock

piece-to-piece would be consecutive parts from the same lathe, but a new family lathe-to-lathe would have to be captured. If the shafts had been produced on a six-spindle screw machine, then piece-to-piece would be consecutive pieces from the same spindle, not consecutive pieces from the machine. One machine cycle would produce six parts. These broader families come from parallel process paths. Create a detailed process flow diagram to find the parallel paths. Cavity-to-cavity, pallet-to-pallet and line-to-line are common examples.

Only one lathe produced rotor shafts, so there were no parallel paths. Piece-to-piece was captured by sampling consecutive parts from the lathe. Differences in size or shape from one part to the next would indicate instability in the equipment. A sample of three consecutive pieces will reveal trends, shifts, or cycles in the equipment.

Time-to-Time. Processes change over time. Operators make adjustments. They introduce new lots of material. Ambient conditions such as temperature and

humidity vary. These changes may cause changes in process outputs. Select sampling times that are likely to catch important changes. Machine start-up, tool changes, material changes, and operator changes are good candidates for sampling times.

Dorian decided to take samples each hour.

Completing the Plan

Once you identify the families of variation, develop a stratified sampling plan and lay out a graphical format to display the data. A stratified sample means the sample is in the order of production. Stratified samples are not random.

The graphical layout in Figure 3 will capture the data from three consecutive parts taken each hour from 8:00 a.m. to noon.

Executing the Multi-Vari Study

The number of samples required to find the largest family of variation cannot be predetermined.

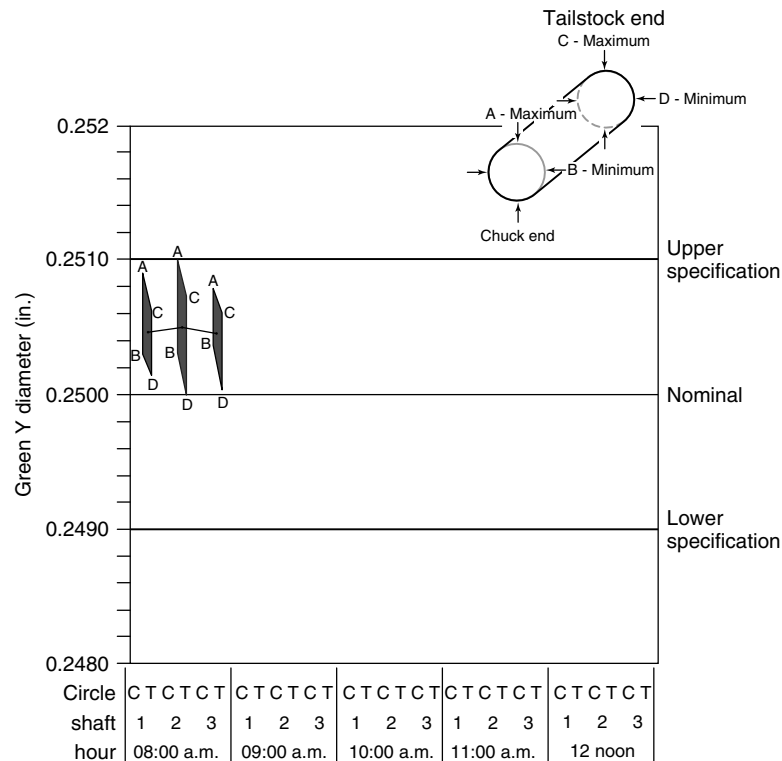


Figure 4 Sources of variation for the rotor shafts. C: chuck; T: tailstock

Continue sampling until you capture variation of Red X magnitude (80% of the overall variation should be sufficient). During the sampling, do not try to control or vary specific process parameters. Let the process vary normally. Measure parts immediately and plot the data. Stop sampling when the Red X family is revealed. It is a bit like fishing. Cast your net in a small area. If you do not catch anything, progressively cast the net wider and wider until the catch is made.

Analyzing Multi-Vari Charts

Once the data are plotted, look for the largest family. Calculations are discouraged. The magnitude should be clear from the graph. Also look for patterns in the data. These provide important insights for those who understand the mechanics of the manufacturing process.

Figure 4 shows the multi-vari chart for the rotor shaft after the first three-piece sample at 8:00 a.m.

The vertical axis captures changes in part diameters. The distance from A to B represents out-of-roundness at the chuck end of the part. The distance from C to D is out-of-roundness at the tailstock end. Notice that the first three pieces have the same out-of-roundness at each end. The vertical distance from A to C represents part taper.

The multi-vari chart has revealed considerable diameter variation within each piece. The variation uses half the tolerance. However, we have not seen variation of Red X magnitude. The process capability study revealed variation exceeding the tolerance limits. More data are needed.

Figure 5 shows the multi-vari chart at noon.

Notice that the taper has increased and the out-of-roundness has stayed about the same. More importantly, changes in the position of the hourly subgroups exceed the tolerance and are the largest family. Whatever be the Red X, it is time based. Even though there is considerable within piece variation, that is not where the Red X resides and it must be eliminated for now.

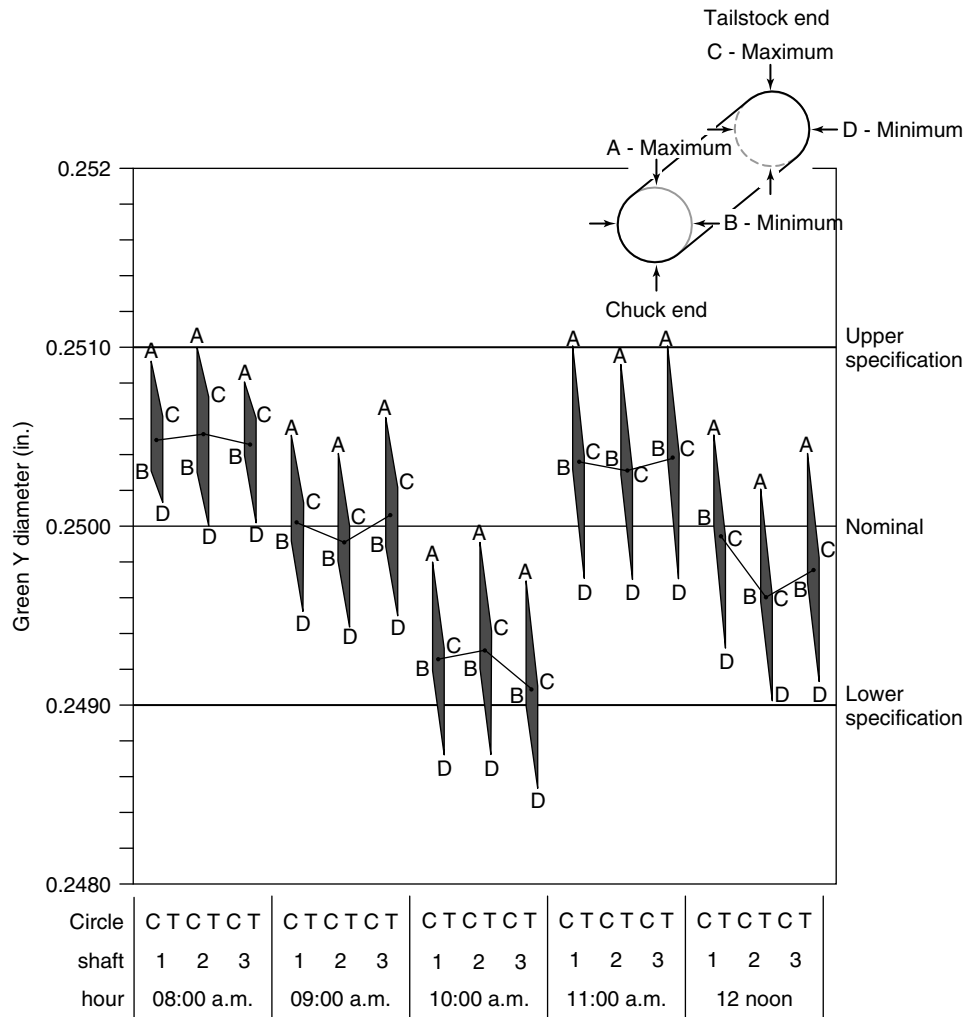


Figure 5 Multi-vari chart for the rotor shafts illustrating that the biggest variation is hour-to-hour. C: chuck; T: tailstock

Multi-vari charts reveal the largest family and they also reveal patterns of change. The time-to-time family trends downward from 8:00 to 10:00 a.m.; it shifts between 10:00 and 11:00 a.m.; and it starts another downward trend from 11:00 a.m. to noon.

The statistical portion of the investigation is over. Now process knowledge and an understanding of the physical world take over. The manufacturing supervisor thought excessive tool wear caused the trend and a new tool explained the shift between 10:00 and 11:00 a.m. The tool comes from the outside. If the tool were wearing, the parts would get larger. Our parts are getting smaller. A material

lot change could explain a shift, but is not consistent with the trend.

A few questions revealed that the operator took a break at 10:30 a.m. During the break, the process shut down. Temperature is a time-to-time variable that fits the clues. Machining parts involves friction and raises temperature. Metal expands as it heats up. Expanding metal will cause the parts to be smaller after they have cooled. Shutting down the lathe during break provides time for cooling, allowing the process to go back to its eight o'clock state. The lathe has a coolant reservoir to control process temperature. An examination of the reservoir revealed that the coolant

was well below the prescribed levels. Adding coolant eliminated the time-to-time changes and made the process capable.

Multi-vari charts do not reveal the Red X. They reveal the Red X family. Further investigative steps are always required. Well-designed and well-executed multi-vari charts eliminate broad areas of the process. They allow the investigator to narrow down the search.

Kraft Paper Strength Case

Kraft paper is used for packaging. Paper strength is an important material property. A paper mill was losing market share because their paper strength was inconsistent. The paper process is more complicated than making rotor shafts on a lathe.

Trees are cut and the logs shipped to the mill. A chipper converts the logs to wood chips. A digester cooks the chips in a chemical soup that breaks down the glue that holds the fibers together. After rinsing, the wood fibers are suspended in water, creating a batch called a *furnish*. The furnish is fed into the paper machine where the fibers are distributed across a moving wire. The material dries as it moves down the machine and reels take up the finished paper at the end. The experts blamed the variation in paper strength on the trees. Some trees simply have stronger fiber than others and the paper mill cannot be held responsible.

A multi-vari chart was planned to reveal variation

- across the paper machine;
- in the machine direction; and
- from furnish-to-furnish.

The stratified sampling plan is illustrated in Figure 7. See Figures 6 and 7.

Figure 8 shows the resulting multi-vari chart:

There are several important observations and conclusions:

1. The largest family is furnish-to-furnish.
2. There is a nonrandom pattern in paper strength across the machine.
3. The furnish-to-furnish variation eliminates the paper machine and points to the digester.
4. The nonrandom pattern eliminates the trees. How would the weak fibers know to align on the right side of the paper machine?
5. There is more across machine variation with the weak furnish than with the strong furnish. This signals an interaction.

If every furnish could be made as strong as the best furnish, the across machine variation would not matter and the customers would be delighted. We do not know the name of the Red X, but we know where it resides. It is in the digester.

An investigation into the digesting process uncovered four variables that were supposed to be controlled:

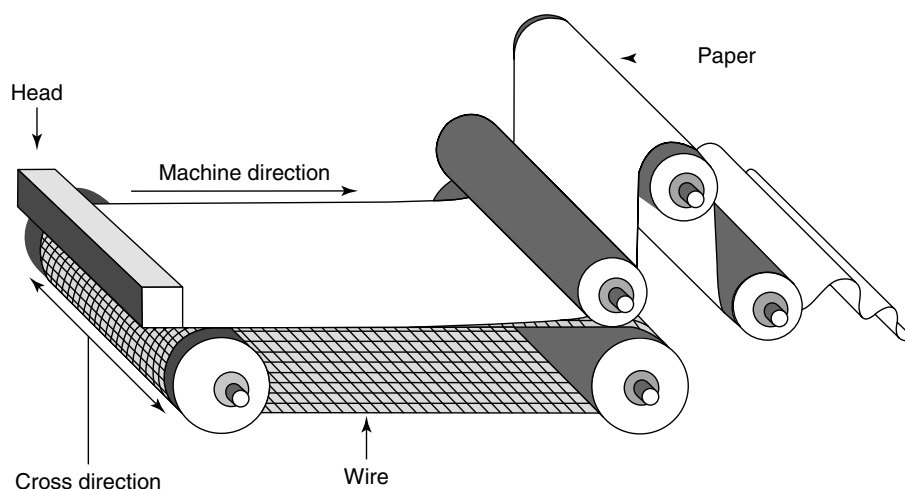


Figure 6 Kraft paper process

8 Multi-Vari Charts

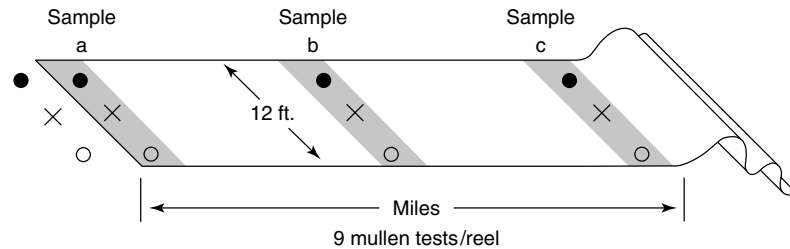


Figure 7 Stratified sampling plan for Kraft paper study

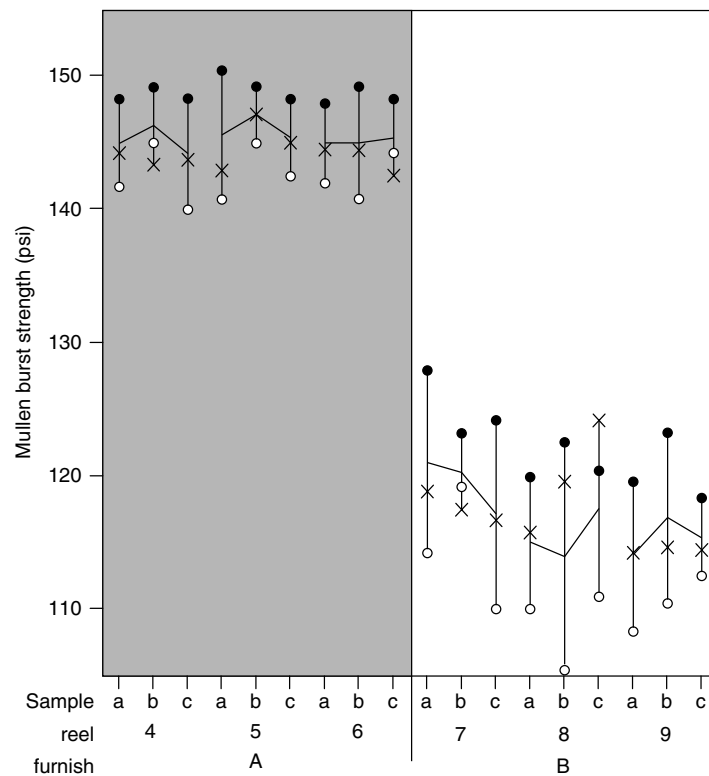


Figure 8 Multi-vari chart for the Kraft paper investigation showing that the biggest source of variation is furnish-to-furnish

1. Cooking temperature
2. Cooking pressure
3. Cooking time
4. Loading ratio (chips per gallon of water).

A full **factorial** designed experiment revealed an interaction between cooking pressure and loading ratio. See Figure 9. Once the interaction was understood and the levels controlled at their optimum, all paper reels were strong, exceeding customer

expectations. The lost market share was recovered and the share was raised to new heights.

A molder of plastic parts was having trouble maintaining the proper material packing. Packing is measured as part weight (grams). The families of variation were:

- Cavity-to-cavity
- Shot-to-shot
- Hour-to-hour

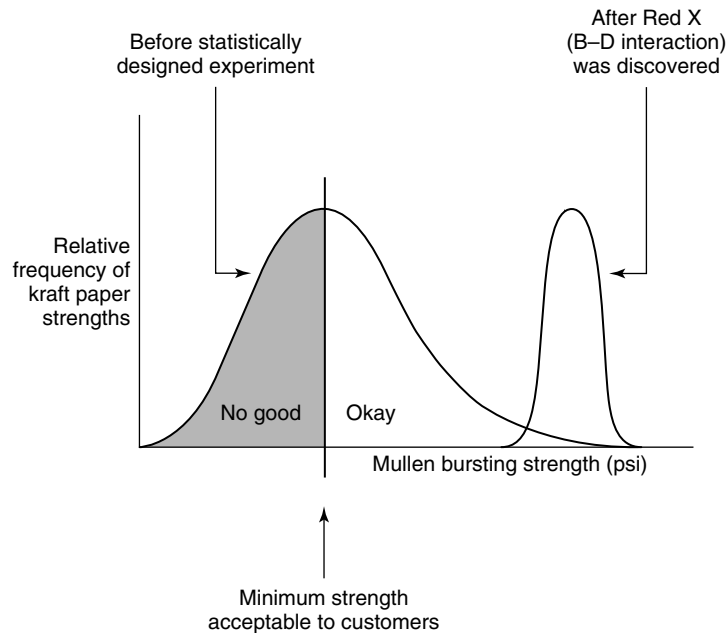


Figure 9 Kraft paper study: Using the interaction of two control variables to reduce variation

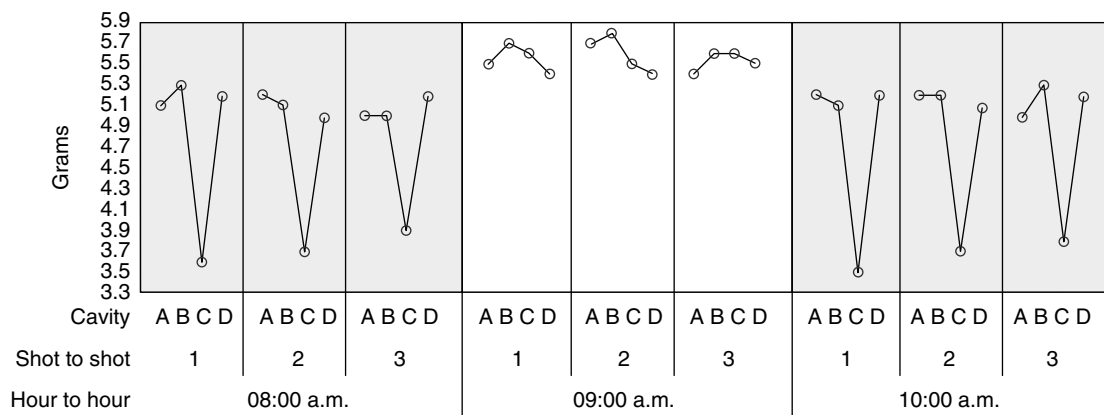


Figure 10 Kraft paper study: Multi-vari chart showing that the biggest source of variation is hour-to-hour. Also notice that cavity-to-cavity variation is the smallest at 9:00 a.m.

The multi-vari chart revealed an interesting pattern (See Figure 10). There is little variation in the shot-to-shot family. There is substantial variation in cavity-to-cavity at 8:00 a.m. and 10:00 a.m., but no cavity-to-cavity variation at nine o'clock. Another interaction has been uncovered.

There is an interaction between the cavity-to-cavity and hour-to-hour families. The investigative

team studied the cavities and could not find anything different about cavity C. When looking at the runners, they found C to be slightly smaller. The hour-to-hour variation was by the machine operator. When the operators saw a problem with the parts from cavity C being too small, they increased pack pressure. Notice that all the parts at nine o'clock are heavy. After a while, flash would increase, so they would decrease

pack pressure, thus turning on the Red X in cavity C. Increasing the runner size for cavity C solved the problem.

Using Multi-Vari Charts Effectively

The multi-vari chart is a powerful statistical tool without the statistical calculations. Properly planned and executed, the chart will reveal the largest family of variation and make a significant, early contribution to the search for the Red X.

Selecting the proper families requires an understanding of how the process creates the shape or

property desired. It also requires walking the process to see the parallel paths.

Be sure to take stratified samples and plot the data as soon as they become available. Stop sampling when variation of Red X magnitude is captured.

The chart will reveal the largest family of variation and patterns of change. The patterns will provide guidance on potential next steps.

Remember the multi-vari chart will not identify the Red X. However, as Sherlock Holmes might say, it will eliminate a great number of variables that cannot be the Red X.

RICHARD D. SHAININ

Total Productive Maintenance

TPM Definition

The total productive maintenance (TPM) definition was established by the Japan Institute of plant maintenance (JIPM) in 1971 as TPM was first introduced in industry. The initial emphasis of TPM was primarily on the production floor, but as the focus shifted to the company-wide implementation, a new TPM definition was introduced in 1989. While the new definition emphasizes “company-wide TPM”, the old definition focused on “TPM for the production sector”. The TPM definition, as it is used today, is as follows. TPM aims at:

- establishing a corporate culture that will maximize production system effectiveness;
- organizing practical shop-floor systems to prevent losses before they occur throughout the entire production system life cycle, with a view to achieving zero accidents, zero defects, and zero breakdowns;
- involving all the functions of an organization including production, development, sales, and management;
- involving every employee, from top management down to frontline operators, and;
- achieving zero losses through the activities of “overlapping small groups”.

Basic Concepts of TPM

The basic concepts of TPM are as shown in the following five items. Each is explained in relation to the above mentioned TPM definition, as follows:

1. Establish a corporate structure that will be profitable

TPM aims at establishing a corporate climate that will maximize production system effectiveness, as indicated in the first item of the TPM definition. To attain this, efforts should be made so that all possible losses may be prevented (the second item of the TPM definition) thus establishing a corporate structure, which will be profitable.

2. Prevention – to keep undesirable things from happening

A system should be built whereby all possible losses occurring in the entire production system life cycle will be prevented (the second item of the TPM definition). For this purpose, maintenance prevention (MP), preventive maintenance (PM), and corrective maintenance (CM) will be implemented. In common, among these, is the idea of prevention.

3. Total involvement – participative management and respect for humanity

As indicated in the third and fourth items of the TPM definition, every function and every member of an organization, from top management down to frontline operators, should participate, through overlapping small-groups’ activities, as indicated in the fifth item of the TPM definition. Thus, TPM is literally participative in nature, respecting humanity as advocated by Rickart.

4. “Genba” (work site) and “Genbutsu” (equipment itself) principles

In organizing a shop-floor system to prevent losses (the second item of the TPM definition), the focus should be on creating a loss-free system in equipment itself, that is, on “Genba” and “Genbutsu”, rather than on creating a system of control. In other words, efforts should be made to realize equipment “the way it should be”, to make “management by vision” possible and create a clean workshop to work at.

5. Zero-oriented approach

This approach inherently has no tolerance for loss, rather than allowing for a margin of loss in production systems. It requires a shift in attitude toward operations, away from “reducing” loss to “eradicating” loss. This is a breakthrough that causes a change in the way of thinking, from the notion that “some failures are unavoidable” to that “failures must be avoided”.

Visualization of Loss

In TPM loss needs to be expressed in numerical form. To this end, concepts like overall equipment efficiency (OEE) and overall plant efficiency (OPE) have been developed, as tools for identifying loss owing to equipment. (For more information on OEE and OPE, see Figures 1 and 2.)

2 Total Productive Maintenance

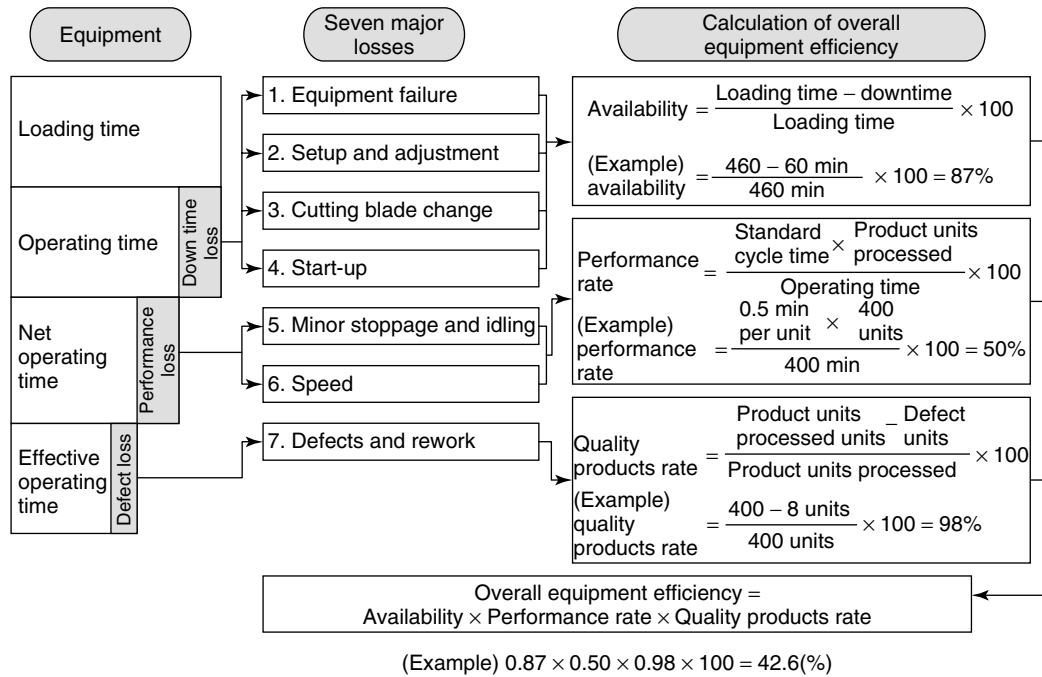


Figure 1 Overall equipment efficiency (OEE)

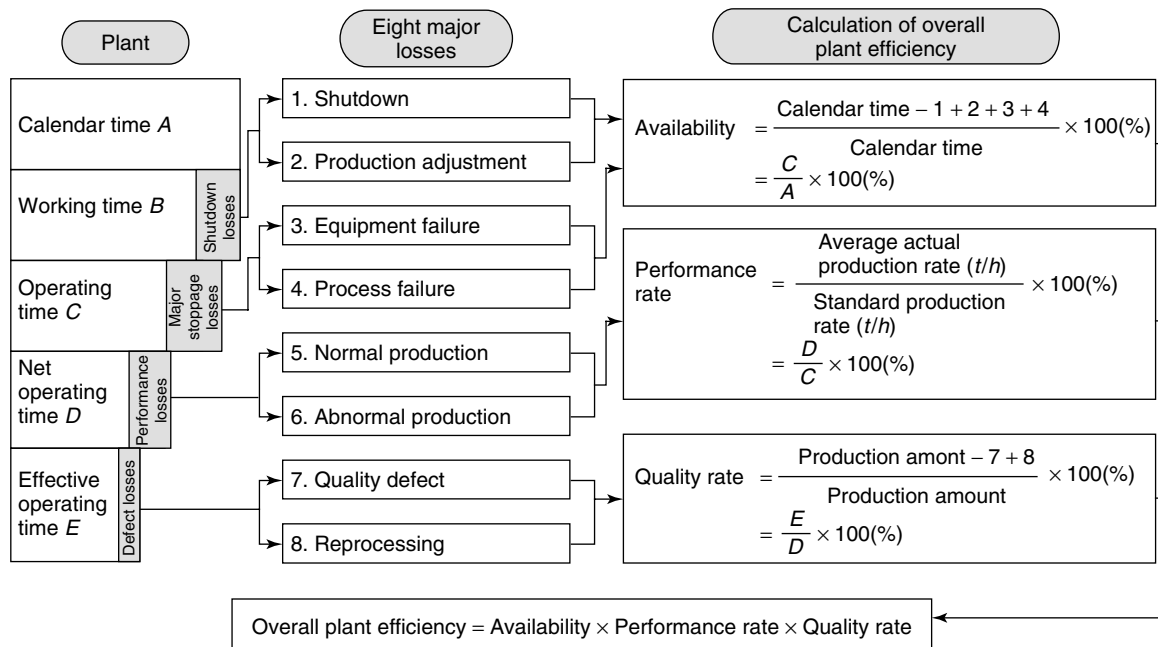


Figure 2 Overall plant efficiency (OPE)

Table 1 Case of results

	Target item	BM FY2004	Result as of September 2006
Cost	Total productivity improvement	100	112
	Processing cost loss amount	100	67
Delivery	Inventory asset turnover days	100	55
	Delivery punctuality rate	100	100
Quality	Number of installation defects	100	0
Productivity	Number of failures	100	7
	Equipment total efficiency	100	102
Safety	Number of accidents requiring absence * (Number of accidents not requiring absence)	0	0
Morale	Number of submitted proposals	100	119

Note: Mitsuba Corporation, Akagi Plant: "Let's make an evolving strong plant by TPM activities for making profits and winning", TPM awarded company presentation, 2006 *
Accidents accompanied by lost worktime

Result by TPM: from case show example of tangible effects of the company which presented TPM.

Table 1 shows indices before and 2 years after the actual implementation of TPM at a plant. Even a plant that has worked on various improvement activities before the TPM implementation proves that the TPM concept achieves sustainable operation with less loss.

Further Reading

- Bosman, H. (1999). Volvo cars Europe-gent: TPM activities for world class, in *TPM World Japan, Conference*.
Eertink, H. (2006). Human resource training in unilever Europe – foods, in *Monodukuri CEO Japan, Session*.

- Mitsuba Corporation. (2006). Akagi plant: let's make an evolving strong plant by TPM activities for making profits and winning, in *TPM Awarded Company Japan, Presentation*.
Murase, Y. (2006). Report "Three Zero Practice – Zero Accidents, Zero Defects and Zero Failures" (Japanese Version), Japan Institute of Plant Maintenance.
Nakajima, S. (1982). *Tpm Development Program – Implementing Total Productive Maintenance (Japanese Version)*, Japan Institute of Plant Maintenance.
Rowland, D.H. (1999). TPM in Milliken & Co., in *TPM World Japan, Conference*.
Sakaguchi, M. (2002). *Trend of the 21st Century's TPM. Concept of TPM Part 1-2-3 (Japanese Version)*, Japan Institute of Plant Maintenance.

TSUTOMU NAKAMURA

Lean Methods, Creating Flow with

Introduction

Lean is a management philosophy founded on the just-in-time (JIT) principles conceived by Japanese automobile manufacturers in the 1930s and perfected in the 1970s. In 1990, the book, *The Machine That Changed the World* [1] kindled interest in JIT and lean thinking in the United States. Subsequently, in 1996, Womack and Jones wrote the book, *Lean Thinking* [2]. In this book, they indicated that the goal of lean thinking is to eliminate *muda* (a Japanese word for wastefulness) and prescribed a specific course of action to implement lean concepts within an organization.

The Evolution of Lean Thinking

The principles in lean thinking were initially applied to the automobile industry in the United States. This was a natural consequence of the fact that JIT principles were specifically developed for the automobile industry in Japan. However, lean thinking is now being applied to a wide variety of industries. These concepts and principles have successfully been applied in industries as diverse as paints, furniture, electrical switchgears, aerospace, aircraft maintenance, electrical appliances, and office products.

As US organizations began to understand and embrace lean principles, the body of knowledge also started to grow and get disseminated. It was observed that lean principles, in particular, the concept of flow, could be extrapolated beyond internal operations to entire supply chains. Lean became a new model for making, distributing, and selling products that evolved from a marriage of new capabilities in manufacturing, information technology, and logistics.

Lean thinking focuses on eliminating wasteful activity at all levels in the organization (or the supply chain for that matter). Eliminating wasteful activity achieves two purposes. First, it reduces *lead times*^a and makes the organization more *flexible* and *responsive*. From the supply chain perspective, the

organization is more responsive to the downstream customers and provides smoother, more predictable demand for the upstream suppliers. For example, the Denso starter and alternator plant in Maryville, Tennessee, operates on a product cycle in which every product is made multiple times per shift. The short product cycles make it possible for both Denso and its upstream suppliers to produce virtually in lockstep with the automobile manufacturers that use these starters and alternators. Only a day or two elapses between the time a casting or a coil of wire is made at the Denso facility, and the time the casting or coil becomes part of a starter or an alternator on a completed car. The benefits of this short flow time – short quality feed back loop, responsiveness to the customer, elimination of supply chain costs related to inventory tracking – are primarily realized outside the four walls of Denso's operation. (In fact, if all elements of the supply chain in the automobile industry were as responsive as Denso, the automobile industry would be poised to adopt Dell Inc.'s direct sales model.)

Second, the elimination of wasteful activities frees up resources that could be deployed to grow its business. An organization that exemplifies such a growth model is Michigan-based Freudenberg NOK, the American partnership of German-based Freudenberg and Japan-based NOK, the world's largest producers of oil seals and custom molded rubber products. Created in 1989, this organization has relentlessly applied lean concepts to squeeze out the waste in its processes, freeing up capacity for further growth, while at the same time becoming more flexible and responsive. Freudenberg NOK had quadrupled its revenues by 2001 and its current revenues are more than \$1 billion.

Lean is More than Just Effective Operations Management

The preceding discussion should clearly convey to the reader that lean provides a set of concepts, tools, and management *prescriptions* aimed at strengthening the competitive advantages of the organization [3]. Lean is not just something within the domain of the operations function; rather, the business case for lean applies to the entire organization. To repeat a statement made earlier, lean allows organizations to pursue a *growth strategy*, a strategy aimed at uncovering additional capacity that could be deployed for

further growth. The greatest benefit of lean processes lies in the fact that organizations can grow their business *without having to spend money now* to build new plants, invest in a new warehouse, and so on. A lean implementation has the ultimate goal of delivering products that the customer values, as quickly as possible, further enhancing the organization's ability to grow its business.

Organizations that successfully implement lean concepts can, in fact, respond to customer demands faster and more reliably than their competition. In turn, that encourages the customer to place orders with the organization that are more reflective of the actual demand, rather than one that the customer has inflated in order to build a safety cushion, or deflated to reduce excess inventories. The organization is thus able to work with a smoother production schedule and consequently communicates a smoother material supply schedule to its suppliers. This facilitates strong relationships with suppliers, providing lean organizations with another competitive advantage. Lean organizations are thus able to build custom products for individual customers much more easily, without having to carry large amounts of inventory.

The Reasons for Working with Reduced Inventory Levels

Organizations often focus on reducing inventory levels to save on inventory carrying costs. However, there are arguably three other reasons that are more important – reasons that are more customer focused. First, reduction in inventory levels reduces lead times, enabling the organization to ship the product faster to the customer. For instance, if the organization has 2 weeks of work-in-process (WIP) inventory on the shop floor, then a new order that is processed using the First-Come-First-Served principle will take at least 2 weeks to complete processing. The organization is also more responsive to changing customer preferences. Second, the presence of large finished goods inventory of certain products is symptomatic of the organization having built the *wrong kind* of products, perhaps because it builds products in large batches to exploit scale economies. Finally, reducing inventory exposes problems such as poor quality, unreliable suppliers, scrap, large changeover times, and so on.

Tools for Creating Flow

To enable organizations to reduce inventory levels and create a smooth flow of products throughout the organization, or the supply chain for that matter, all the elements involved should work in concert. In the context of a supply chain, all the supply chain partners must work at the same pace. Similarly, within an organization, the most effective method of production is to have all processes work at the same rate. There is no point in having some of the resources work faster than the others, because that will invariably pile up inventory in front of the slower processes. Such a smooth flow of products is better facilitated if all resources respond to *pull* signals. The signal for a resource to produce could be generated simply when the downstream resource draws a unit of WIP from the buffer placed between the two resources.

Some of the important tools and techniques used by lean thinking to work with lower inventory levels and promote flow are: 5-S, total productive maintenance, *takt* time, standard work, mixed model scheduling, cellular layout and one-piece flow, point of use material storage, pull replenishment, and *kanbans*. Although these lean tools were originally developed for manufacturing organizations, many service organizations are finding them very useful as well. It is not possible to discuss these techniques within this short article. For more details, the reader can refer to *Streamlined* [4, Chapter 5], or numerous other books that treat this topic.

We note that before applying these tools, we must also address the scope of the lean implementation. Organizations often create lean cells to put out products at a rapid rate, only to find these products queuing up at a downstream work center. The questions that must first be addressed are as follows: should the organization implement lean on *all* its processes and activities or should it focus on a *subset* of its processes? Similarly, should it implement lean on all its products or on a subset of them? It is recommended that the organization choose one product family at a time and implement lean on *all* the processes and activities of this product family, before moving on to the next product family.

Continuous Improvement and the Pursuit of Perfection

Complacency can become the toughest challenge in any lean transformation process. Lean is not a

one-time implementation effort, nor is it a quality program of the month. It is an ongoing journey that requires a sustained effort at continuous improvement. The transformation to lean is not just an application of a few techniques; rather, it is a whole new way of looking at the operations of the organization. Because lean tools and techniques are often quite different from traditional tools, there must be a sustained effort to operate with these tools. Organizations that do not aim to continue an upward momentum will falter and fall behind their competitors. It is essential to continuously reexamine processes and look for ways to take out waste and nonvalue-added activities if organizations are to gain significant financial performance improvements.

To sustain lean implementations, it is therefore necessary for organizations to constantly initiate *kaizen* (continuous improvement) events that promote continuous improvements. At the same time, there is a need to promote *kaikaku*^b, the radical redesign of processes and methods geared for achieving breakthroughs in performance and growth. Once a *kaikaku* step is applied, *kaizen* becomes a powerful follow-up drive to perfect the processes and methods, and to continue to adapt and be relevant.

As Womack and Jones suggest, “dreaming about perfection is fun” [2, p. 26]. However, lean thinking provides ways to make these dreams a reality. Applying lean thinking and implementing the tools and techniques to create flow, result in a fundamental change in the way the organization thinks about its operations. The employees soon realize there is no end to the pursuit of perfection – efforts aimed at reducing effort, time, space, cost, and mistakes in the process of producing and delivering a product. When products flow faster through the organization, they expose hidden waste in the **value stream**. The harder you pull, the more the obstacles are revealed so they can be removed.

While lean implementations must have commitment and support from the top management, the shop floor personnel are also critical to the success of lean. Many organizations have initiated their lean efforts from the bottom up. Continuous improvement (*kaizen*) events play a vital role in getting them engaged in the lean journey, and that pays dividends. Managers at all levels became lean thinkers and change agents. A truly

lean organization makes it much easier for everyone, including shop floor employees, supervisors, lean champions, subcontractors, or first-tier suppliers, to discover better ways to create value. Because the feedback loops are significantly shortened, there is a faster feedback to employees, providing a more conducive environment for employees to pursue perfection.

Conclusion

Thomas J. Watson Jr. once said, “Whenever an individual or a business decides that success has been attained, progress stops.” Lean is a journey – an ongoing journey that requires a sustained effort to maintain the momentum. At the same time, once the initial resistance is overcome, it becomes much easier to maintain the tools and techniques of lean. Once there is a sharpened awareness among the employees, of the waste present in the system, there will be a concerted effort to maintain the momentum if the right incentives are provided. All types of organizations can benefit from lean thinking, regardless of whether they are involved in manufacturing, process, distribution, software development, or financial services. Some important points organizations should keep in mind when embarking on a lean journey are as follows:

1. While the goal of lean has often been identified as the removal of *muda* (waste), it is important to note that removing *muda* is just a means to an end. A major focus for lean is to reduce lead time. *Lean is all about lead-time reduction and creating flow.*
2. Some important steps organizations can take to create flow are, *takt* time, standard work, pull replenishment, and 5-S. Incidentally, all of these are steps that readily apply in a service setting as well.
3. Lean is a *journey* – one that really never ends.
4. Organizations that embark on the lean journey, typically should start by identifying a product family they can apply lean principles to. Lean thinking must be applied to all the processes in the organization that work on the selected product family(ies). The idea is *not* to simply lean out some of the process steps and create a few *islands of excellence*. While some waste may be removed in the process

of creating such islands, the products flowing out of these islands will end up waiting elsewhere in the organization. In particular, they will queue up in front of the constraint resources.

End Notes

- ^a. The lead time is defined as the elapsed time from the moment an order is received, until the time the order is executed and delivered correctly to the customer.
- ^b. This step is often referred to as a *kaizen blitz*.

References

- [1] Womack, J.P., Jones, D.T. & Roos, D. (1991). *The Machine that Changed the World*, HarperCollins, New York.
- [2] Womack, J.P. & Jones, D.T. (1996). *Lean Thinking*, Simon & Schuster, New York.
- [3] Miller, A. (2007). *Training Manual*, The Lean Organization Systems Design Institute, Center for Executive Education, The University of Tennessee, Knoxville.
- [4] Srinivasan, M.M. (2004). *Streamlined: 14 Principles for Building and Managing the Lean Supply Chain*, Thomson-Textere, Mason.

MANDYAM M. SRINIVASAN

Error Proofing, Healthcare

Introduction

Human errors often have great effects on quality, safety, and efficiency in healthcare: e.g., giving a medication to the wrong patient, wrong-side surgery, overlooking patient's allergies, not connecting the monitoring system when the patient is moved, or filling prescriptions with the wrong drug [1]. To prevent these errors and their undesired effects, "error proofing" solutions are effective. Successful examples include distinguishing patients with a bar code, having the surgeon's initials placed on the surgery location, using color-coded wristbands on patients with known drug allergies, using a checklist in hand, and not using drugs having similar names/forms. The basic idea of error proofing is very simple. Although human beings are flexible and creative, they also often make errors. This characteristic cannot be changed. Therefore, improving the other element of the work system, i.e., "work operations" including medications, equipment, and documents is an effective way for tackling errors. Error proofing has many different names such as foolproofing, *poka-yoke*, mistake proofing, and so on. All can be explained by this simple concept, "changing work operations to fit human beings".

Although error proofing solutions are successfully applied in healthcare, most of them are generated individually as needed to suit each particular occasion without the benefit of the knowledge of the wider body of error proofing solutions. This paper extracts the error proofing principles and solution directions existing behind these proven solutions and provides a simple methodology for systematically generating workable solutions to reduce healthcare-related human errors. We build on a long-established practice of error proofing in manufacturing industries [2–6].

Error Proofing Principles for Healthcare

In spite of the simplicity of the basic idea of error proofing, there are various types of solutions. Error proofing is often taught by having people learn many examples. There, however, exist only a few basic principles to be known. Figure 1 illustrates the five principles of error proofing that have been extracted based on the investigation of more than 500 error proofing solutions implemented in healthcare and published in 30 articles and books [7–9]. The top part of this figure shows a process where human errors occur and cause incidents/accidents. Every operation has its own tasks and risks. Caregivers have to perform the required functions to accomplish these tasks, creating the potential for human error. These errors cause abnormalities, and if necessary actions are not taken, finally cause undesirable results such as

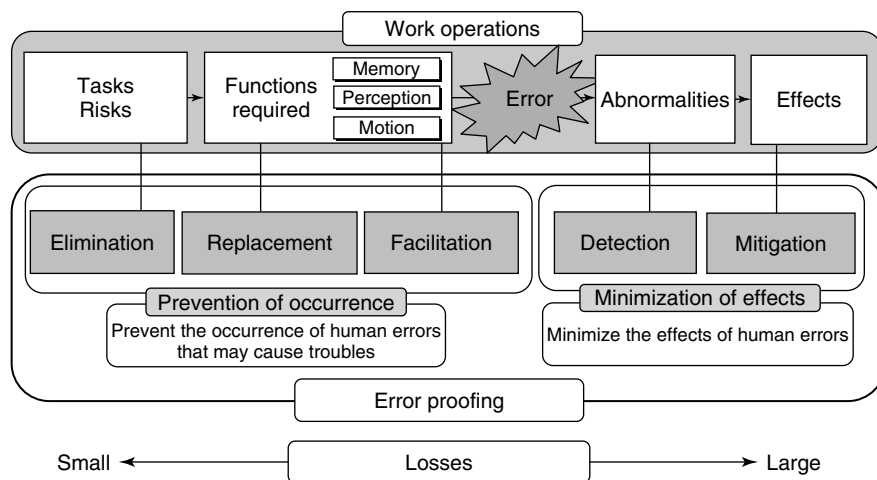


Figure 1 Principles of error proofing in healthcare

2 Error Proofing, Healthcare

human injury or death. The five principles correspond to the process.

The first three principles aim to prevent the occurrence of human errors. The most effective method of preventing errors is to eliminate operations susceptible to error from the process. The next most effective method, if the first cannot be applied, is to replace the human operations with more reliable machines/methods. The third most effective is to make the operations easier for caregivers to perform. These three principles belong to the first group: “prevention of occurrence”, and are called *elimination*, *replacement*, and *facilitation*, respectively.

On the other hand, the other two principles aim to suppress the development of errors into incidents/accidents. One method for attaining this is to detect abnormalities caused by human error and enable caregivers to take suitable corrective action. Another method is to mitigate the effects of human error. These two principles belong to the second group: “minimization of effects”, and are called *detection* and *mitigation*, respectively.

These principles have good consistency with those obtained in manufacturing [2]. However, the ratio of examples corresponding to each principle is different between healthcare and manufacturing as shown in Table 1. The application of “facilitation” in healthcare has a higher percentage of use than the application of “detection”.

Elimination

The principle “elimination” aims to change work object properties so as to remove operations susceptible to human error from the process. For example, an “elimination” solution against errors in the transfer of medication orders, charts, or specimens is to eliminate as many handoffs as possible.

The principle “elimination” can further be classified into two subprinciples:

1. **Task elimination** Eliminate error-prone tasks.
2. **Risk elimination** Eliminate risks inherent in objects or objects themselves.

The previous example corresponds to “task elimination”. On the other hand, removing concentrated potassium chloride solutions from the floor stocks in emergency rooms is an example of “risk elimination” to prevent misadministration of the medications.

The solutions based on the principle “elimination” can perfectly eliminate the possibility of human errors. In many cases, however, process/equipment design must be changed drastically, and it can have great side effects on cost, productivity, and performance.

Replacement

The principle “replacement” aims to replace the human operations with more reliable machines/methods. For example, a “replacement” solution against the error of forgetting to switch on a humidifier in a respirator is to link the switch of the respirator with the switch of the humidifier so that the humidifier is automatically switched on.

The principle “replacement” can further be classified into two subprinciples:

1. **Automation** Completely replace a particular human operation with machines.
2. **Support system** Give support tools such as checklists, reminders, guides, or samples to help caregivers to perform the operations more reliably.

The previous example corresponds to “automation”. On the other hand, using a partitioned cart that can contain only a certain number of medications is an example of “support system”. This makes counting by nurses unnecessary and prevents the error of counting medications incorrectly.

Various “replacement” solutions can be considered depending on how many functions are replaced. To replace all functions leads to large-scale and unrealistic solutions. Therefore, it is essential to focus on highly error-prone functions in the operations and replace them.

Table 1 Ratios of error proofing solutions

	Elimination	Replacement	Facilitation	Detection	Mitigation
Healthcare	7% (35)	26% (135)	48% (251)	14% (73)	5% (24)
Manufacturing	5% (56)	27% (276)	23% (234)	40% (402)	5% (46)

Facilitation

The principle “facilitation” aims to make the operations easier for caregivers to perform. For example, a “facilitation” solution against the error of mistaking the wrong concentration of morphine sulfate stocked in identical cartridges is to have the pharmacy label the high-concentration cartridge with bright orange tape.

The principle “facilitation” can further be classified into three subprinciples:

1. **Simplification** Decrease the number of changes or differences.
2. **Distinction** Distinguish changes or differences from each other.
3. **Adjustment** Adjust objects or operations to suit the individual’s capabilities.

The previous example corresponds to “distinction”. On the other hand, establishing standard drug administration times is an example of “simplification” to prevent forgetting to administer drugs at the ordered time. Making the size/form of the letters used in prescriptions easy to read also reduces the chance of misreading. This is an example of “adjustment”.

Individual “facilitation” solutions are not so effective, but their cost and side effects to operations are minimal. Therefore, it is possible and practical to use multiple solutions to increase overall effectiveness.

Detection

The principle “detection” aims to ensure that abnormalities caused by human error are detected and suitable corrective action is taken. For example, a “detection” solution against the errors in ordering medications is to use order entry systems that provide real-time alerts if a medication order is out of range for weight or age.

The principle “detection” can further be classified into three subprinciples:

1. **Recording and verification of motion** Record precedent motions and verify them at a subsequent point in time.
2. **Restriction of motion** Restrict motions to allow caregivers to notice abnormalities.
3. **Verification of result** Check the resultant medications/equipment/documents at a certain point in time.

The previous example corresponds to “verification of result”. On the other hand, using trays for surgical instruments having indications for each instrument, and nesting all of the instruments for a given operation in the tray so that it is clear if the surgeon has not removed all instruments from the patient before closing the incision is an example of “recording and verification of motion”. Redesigning anesthesia machines so that the connector for the oxygen tank cannot fit into the nitrous gas tank to prevent errors in setting the machines is an example of “restriction of motion”.

Because late detection leads to large correction costs, it is essential to develop techniques for detecting abnormal motions or their results as soon as possible. Another key success factor is to prevent hardware failures. Moreover, it is also essential to apply “facilitation” solutions at the same time because human beings sometimes ignore obvious abnormalities and continue procedures according to their understanding.

Mitigation

The principle “mitigation” aims to mitigate the effects of error by incorporating functional redundancies or using shock-absorbing materials/equipment. For example, a “mitigation” solution against errors in blood banks is to take two separate specimens from a patient and have two different technicians to test them independently. This solution introduces **redundancy** and ensures that the problem only occurs if both technicians make the same mistakes.

The principle of “mitigation” can further be classified into three subprinciples:

1. **Redundancy** Make same functions parallel.
2. **Fail-safe** Incorporate material/equipment to mitigate harmful conditions.
3. **Protector** Make harmful conditions not produce losses.

The previous example corresponds to “redundancy”. On the other hand, changing the design of the insufflators so that they are not capable of delivering pressure far exceeding any requirement for its intended use in surgery is an example of “fail-safe” to mitigate the effects caused by errors in operating gas insufflators. Requiring laboratory technicians to wear glasses and training them to immediately flush their eyes with water if they have been splashed with

4 Error Proofing, Healthcare

a toxic substance can mitigate the effects of being splashed with a toxic substance. This is an example of “protector”.

Systematic Generation of Error Proofing Solutions

Although the five error proofing principles described in the previous section are useful for understanding various error proofing solutions, we need some tools for systematically generating error proofing solutions that can attain the change required by each principle for the individual case.

Solution Directions for Error Proofing in Healthcare

When observing the process where each error proofing solution is generated, we find that there are some common directions of thinking to reach the solutions [10]. These directions were investigated and classified into the following 11 types:

1. *Trimming*: example, eliminate manual data entry.
2. *Self elimination*: examples, broken pills do not roll; pill trays.

3. *Standardization*: examples, make asymmetrical parts symmetrical; one size fits all; standard forms.
4. *Unique shapes/geometry*: examples, electrical outlets; pill shape; symbols.
5. *Copying*: examples, duplicate forms; blood ID barcodes; emergency power generator.
6. *Prior action*: examples, premeasured dose; supplier supplies 100% inspected parts.
7. *Flexible films or thin membranes*: examples, medicine bottle safety seal; prepackaged doses; coffee packs; rubber gloves.
8. *Color*: examples, color-coded charts; zones; medicines.
9. *Combining*: examples, patient records; medicine capsules with fast and slow release.
10. *Counting*: examples, tally sheets; count repetitive actions.
11. *Automatic inspection*: (presence/absence; property variation): examples, electronic verification (bar code; patient record validation; patient monitoring).

Relationships between the Principles and the Solution Directions

Table 2 shows the relationships between the error proofing principles and the solution directions. From

Table 2 Error proofing principles and solution directions for healthcare^(a)

Solution direction Principle		Trimming (7%)	Self elimination (7%)	Standardization (20%)	Unique shape/geometry (2%)	Copying (redundancy) (5%)	Prior action (23%)	Flexible films or thin membranes (2%)	Color (5%)	Combining (23%)	Counting (1%)	Automatic inspection (automation) (7%)	Total	Number of error proofing solutions
Elimination (7%)	Task elimination	17	12	3	0	0	5	0	0	4	0	1	42	23
	Risks Elimination	12	0	1	0	0	7	0	0	0	0	0	20	12
	Subtotal	29	12	4	0	0	12	0	0	4	0	1	62	35
Replacement (26%)	Automation	0	0	0	0	4	0	0	0	13	0	26	43	27
	Support system	1	0	17	2	5	108	0	1	53	0	0	187	108
	Subtotal	1	0	17	2	9	108	0	1	66	0	26	230	135
Facilitation (48%)	Simplification	0	0	91	0	0	21	9	4	73	0	0	198	150
	Distinction	18	0	10	0	16	32	0	31	31	0	0	138	83
	Adjustment	10	0	3	3	0	2	0	0	3	0	0	21	18
Detection (14%)	Subtotal	28	0	104	3	16	55	9	35	107	0	0	357	251
	Record and verification of motions	0	0	0	1	0	0	2	0	3	6	1	13	6
	Restriction of motions	0	14	0	8	0	0	0	0	0	0	6	28	14
Mitigation (5%)	Verification of results	0	30	20	0	0	2	0	0	0	0	17	69	53
	Subtotal	0	44	20	9	0	2	2	0	3	6	24	110	73
	Redundancy	0	0	13	0	17	0	0	0	0	0	1	31	17
Mitigation (5%)	Fail-safe	0	0	0	0	0	4	0	0	0	0	0	4	4
	Protectors	0	0	0	0	0	3	1	0	0	0	0	4	3
	Subtotal	0	0	13	0	17	7	1	0	0	0	1	39	24
Total		58	56	158	14	42	184	12	36	180	6	52	798	518

^(a)One error proofing solution may be reached using more than one solution direction

this table, we can see that there exist typical combinations of the principles and the solution directions, which are frequently used in many solutions: e.g., elimination by trimming, replacement by prior action, facilitation by standardization, detection by self elimination, and mitigation by copying.

Consulting Questions for Generating Error Proofing Solutions

Table 3 shows the consulting questions, which have been developed by focusing on the typical combinations between the principles and the solution directions. By answering each of the listed questions, we can generate error proofing solutions systematically for the individual case.

Let us consider the following assumed accident in a hospital: there are two patients per room; the medical charts are hung on the end of the patient's bed; the cafeteria worker unintentionally knocks both charts on the floor while removing the dinner trays and quickly replaces the charts, but on the wrong

beds; the nurse on the new shift gives the prescribed medicine to the wrong patient. This accident was caused by the combination of three human errors:

1. The cafeteria worker knocks both charts on the floor.
2. The cafeteria worker replaces the charts on the wrong beds.
3. The nurse did not notice the mismatch between charts and patients.

Therefore, the questions in Table 3 are applied to each of the above errors for generating error proofing solutions. Table 4 shows a portion of the generated solutions.

Conclusions

In our research, more than 500 error proofing solutions successfully implemented in healthcare were investigated and the principles of error proofing and the solution directions existing behind them

Table 3 Consulting questions for generating error proofing solutions

1. Eliminate tasks/risks
• Trimming – Can we eliminate the error-prone process or harmful objects? • Self elimination – Can harmful action or object eliminate itself?
2. Replace human operations
• Automation (Automatic inspection) – Can we automate the process to solve our problem? • Prior action – Can we do something beforehand to support human operations? • Combining – Can we combine (bring together/closer) two or more things to automate or support human operations?
3. Facilitate human operations
• Trimming – Can we trim similar or confusing things to facilitate human operations? • Standardization – Can we standardize the process to facilitate human operations? • Copying – Can we use redundancy to facilitate human operations? • Prior action – Can we do something beforehand to facilitate human operations? • Flexible films or thin membranes – Can we use flexible films or thin membranes to facilitate human operations? • Color – Can we use color to facilitate human operations? • Combining – Can we combine (bring together/closer) two or more things to facilitate human operations?
4. Detect abnormalities
• Counting – Can we count something to detect abnormalities in the human operations or their results? • Self elimination – Can we let people notice abnormalities by themselves? • Unique shape/geometry – Can we use shapes (1D, 2D, or 3D) to detect abnormalities in the human operations or their results? • Automatic inspection (automation) – Can we automatically inspect something to detect the abnormalities in the human operations or their results?
5. Mitigate effects
• Copying – Can we use redundancy to mitigate the effects? • Prior action – Can we do something beforehand to mitigate the effects? • Flexible films or thin membranes – Can we use flexible films or thin membranes to mitigate effects?

6 Error Proofing, Healthcare

Table 4 An example of generating error proofing solutions

Error	Principle	Consulting question	Generated solution
The cafeteria worker knocks both charts on the floor	Elimination	Trimming – Can we eliminate the error-prone process or harmful objects?	Ensure charts are not dropped by touching
		...	...
	Replacement	...	...
	Facilitation	...	...
		Prior action – Can we do something beforehand to facilitate human operations?	Move chart to new location.
The cafeteria worker replaces the charts on the wrong beds		...	...
	Elimination	...	...
		...	...
	Replacement	...	...
		Prior action – Can we do something beforehand to support human operations?	Paste a name tag at the end of the bed, on which the chart is hung
		Combining – Can we combine (bring together/closer) two or more things to automate or support human operations?	
	Facilitation	...	...
		Color – Can we use color to facilitate human operations?	Color ends of beds and charts to differentiate two patients
		...	...
	Detection	Self elimination – Can we let people notice abnormalities by themselves?	Change the shape of chart hanger so that the wrong patient chart cannot be hung.
...		Unique shape/geometry – Can we use shapes (1D, 2D, or 3D) to detect abnormalities in the human operations or their results?	
		...	...
	Mitigation	...	...
	...	...	...

were extracted. Using them enables us to systematically generate error proofing solutions to prevent healthcare-related errors and implement the activities for ensuring patient safety more effectively and efficiently.

References

- [1] Kohn, L.T., Corrigan, J. M., & Donaldson, M. S (ed) (2000). *To Error is Human: Building a Safer Health System*, National Academy Press.
- [2] Nakajo, T. & Kume, H. (1985). The principles of foolproofing and these application in manufacturing, *Reports of Statistical Application Research, JUSE* **13**(2), No., 10–29.
- [3] Shigeo, S. (1986). *Zero Quality Control: Source Inspection and the Poka-Yoke System*, Productivity Press.
- [4] Hirano, H. (1989). *Poka-Yoke: Improving Product Quality by Preventing Defects*, Productivity Press.
- [5] The Productivity Press Development Team (1997). *Mistake-Proofing for Operation: The ZQC System*, Productivity Press.
- [6] Martin Hinckley, C. (2001). *Make No Mistake: An Outcome-Based Approach to Mistake-Proofing*, Productivity Press.
- [7] Chase, R.B. (1994). Make your service fail-safe, *Sloan Management Review* **35**(3), 35–44.
- [8] Johnstone, P.A.S., Hendrickson J. A., Dernbach A. J., Secord A. R., Parker J. C., Favata M. A., & Puckett M.

- L., (2003). Ancillary services in the healthcare industry: Is Six Sigma reasonable? *Quality Management Health Care* **12**(1), 53–63.
- [9] Leape, L.L., Kabacoff, A., Berwick, D. M., & Roessner, J. (1998). *Reducing Adverse Drug Events*, Institute for Healthcare Improvement.
- [10] Terninko, J., Zusman, A. & Zlotin, B. (1998). *Systematic Innovation: An Introduction to TRIZ*, CRC Press LLC.

TAKESHI NAKAJO, TIMOTHY G. CLAPP
AND A. BLANTON GODFREY

Cost of Quality

Introduction

Cost of quality is a concept that has likely been around since the first caveman discovered his spear was broken just before the tiger pounced. Our understanding of these costs has evolved tremendously over the years driven by examples of costs eliminated and the stunning gains made through statistical **quality control**, continuous quality improvement, **reengineering**, or – more recently – **lean** and **Six Sigma** initiatives. The millions of cost reductions and substantial savings achieved by these activities have refined our thinking about the opportunities in managing the cost of quality. As examples of a \$4 million design change reducing \$20 million of warranty costs, of a \$1000-measurement system eliminating \$500 000 of waste, of a few million dollars in computer-assisted invoice checking saving tens of millions of dollars and many other examples were reported. Senior managers began to look for new ways to understand and find these costs.

For many years, people argued over whether we should discuss cost of quality or cost of poor quality. Most of these arguments were resolved by Joseph M. Juran many years ago with his clear explanation of the difference between features and failures. Adding features to a product or service almost always adds costs. In this definition of quality increased quality means higher costs. Failures almost always cost far more to the organization than they do to prevent, this poor quality is where many organizations found great opportunities for great returns on investment. Reducing failures almost always reduces costs.

But for a full understanding of the costs of our business we must consider both. Features are often a market decision. If the customer is willing to pay more for air conditioning and a superb sound system in an automobile than the cost to provide them, we have made a good decision in improving quality. If we offer features that add more costs than the customer sees as value, we have made a poor decision. If we are missing features that the customer finds absolutely necessary and refuses to purchase the product or service, we have made an even poorer decision and our organization has incurred an enormous cost due to our poor quality. It is hard to imagine a television being sold without

a remote control device. This points out one of the difficulties of feature quality. Many features that add value and increase market soon become necessary standard features. Noriki **Kano** has described this beautifully in his articles on attractive quality and necessary quality.

The Hidden Factory

Failure costs are where most organizations have focused their cost of quality efforts. These costs can be a substantial part of the organization's total operating costs. Both W. Edwards Deming and Joseph M. Juran estimated these failure costs to be as high as 30% of even well-managed companies. Juran's hidden factory example made this clear. He stated that every organization has two factories, one produces product that we sell, the other produces product that we throw away. If we could just find the hidden factory and close it down, we could save 20 or 30% of our costs.

These costs in a real factory are often quite visible: scrap, rework, machine downtime, recycled materials, and so on. The costs in white-collar work are harder to see, but usually far larger: reports that are not read, analyses never used, meetings with no value added, and so on. Even the basic traditional thinking of quality costs identified many of these costs and companies made substantial gains by reducing or eliminating these costs. Unfortunately, many of these quality costs were assumed to be part of the costs of doing business and unmeasured and hidden. Far too many companies only measured the "variance to standard costs" meaning that certain amounts of scrap, rework, machine downtime, work-in-process inventory, and other costs were necessary. Only the amount the standard costs were exceeded were considered a problem. Waste was built-in.

Traditional Quality Costs

For many years we had a limited view of the cost of quality. We defined it as four areas that were reasonably easy to measure: prevention costs, appraisal costs, internal failure costs, and external failure costs.

2 Cost of Quality

Prevention Costs

These costs are often the subject of much debate and confusion. Many people would include a number of design activities, reliability modeling, prototype testing, pilot runs, and other costs as prevention. Others call these standard design and development costs and only include the activities of the quality department in these costs. Some companies do not include these activities at all and just focus on the appraisal and failure costs.

Appraisal Costs

Appraisal costs are also measured differently in organizations. Some organizations include all checking, inspection, measurement, testing, field trials, auditing, and other appraisal costs no matter where they are done or who does them. Many of these costs are attributed to part of the time of assembly workers, supervisors, packagers, and even automated test equipment. Other companies just count the costs of people actually assigned to quality departments and their equipment, space, and other costs. There is also uncertainty about which costs should be included. If the tests are required by the customer, regulatory agencies, or others, are these really costs of quality or costs of selling the product. Most companies include all these costs in their definition of appraisal costs.

Internal Failure Costs

These costs include all costs incurred during the production of the product or service before delivery to

the customer. Thus all rework, machine downtime, delays, scrap, failure analysis, repair, retesting, bottlenecks, excessive work-in-process inventories, review meetings, and replacement parts are all counted.

External Failure Costs

These are the costs incurred after the product or service is delivered to the customer. Returned goods, warranty costs, repair costs, inventories of parts kept to make repairs, service staff, rebates, complaint handling, fines, penalties, lost revenue, and bad will are all counted in external failure costs.

Rethinking Optimum Quality

We also thought for many years that we could find the “optimum quality level”. This theoretical level would be exactly where the costs of preventing defects was equal to the cost of these defects. This is illustrated in Figure 1(a). This thinking led to concepts such as acceptable quality levels (AQLs) and other quality targets that have been discarded by most leading organizations in the past decades. There were many problems with the optimal quality level. The underlying assumption was flawed. We assumed that defects could only be eliminated by increasing costs – more appraisal or higher-quality parts or more careful assembly. Therefore, as quality levels were improved, costs necessarily went up. But there are many ways to improve quality. An automated test set may actually be cheaper and far more effective than inspection by operators thereby

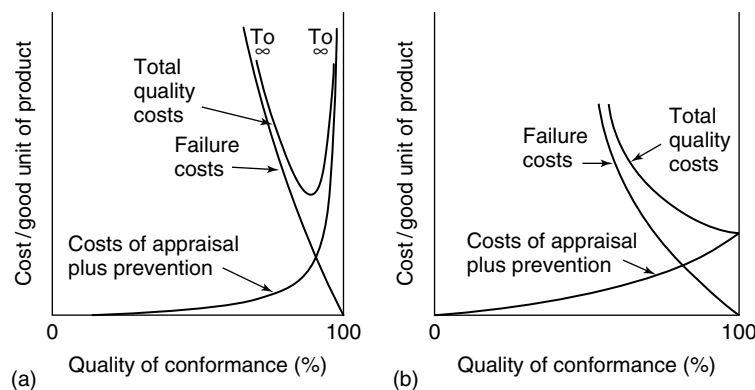


Figure 1 (a) Old thinking of cost of quality and (b) New thinking of cost of quality

reducing inspection costs and improving quality. Statistical process control, error-proofing, and other methods may eliminate defects and reduce internal costs. Again quality is improved and costs are reduced. Better designs or parts from suppliers that are higher quality and lower cost may also change the curve.

There are many other problems with this model. Quality is dynamic. The costs of defects today may be far higher than a few years ago due to more complex systems, customer demands, or repair costs. Quality is relative. An optimum quality level calculated by the producing organization may be worse than what a competitor is offering at even a better price. We may lose market share or the entire market changing the cost calculations dramatically. New methods for checking and assuring quality are continually being developed changing the cost curves far faster than we change our model and calculations.

Today most organizations use a different model, Figure 1(b). The optimum quality level is thought to be zero failures or 100% conforming products. Pursuing perfection is a basic tenet of Six Sigma quality and other philosophies. Each reduction in defects is done on a return on investment basis. Improvements are implemented whenever the cost of the change is less than the savings. For some products or services, we may never actually achieve perfection, but each step toward perfection brings additional savings.

Quadratic Loss

A further understanding of optimum quality is often attributed to Genichi Taguchi although many others were also using this concept. This is the understanding that in many situations the terms defects, nonconformances, or defectives is inappropriate. For example, in producing copper wire for telecommunications or electric power grids, resistance is a key variable. Setting upper and lower limits and calling wires outside the limits defective and all wires inside the limits good misses the point. Any deviation from the target causes loss, the further away from the target the higher the loss. Machined parts, gears, bearings, electronic components, administration of medicines, and many other products and even services should be measured with respect to deviations from a target not in binary descriptions such as conforming and nonconforming.

Newer Thinking of Quality Costs

In doing the research for his Ph.D. thesis at Sweden's Royal Institute of Technology, Lars Sörqvist studied a number of companies' best practices in managing quality costs. Although none of the companies used all the methods, he found that all had expanded their quality cost management beyond the traditional understanding. He summarized quality costs in five categories.

Traditional Poor Quality Costs

These are the costs described above. These costs often just include the sporadic problems that disrupt operations or are measured by the company's traditional accounting system. Chronic problems – often thought to be the normal costs of doing business – are not measured and are not seen as opportunities for improvement.

Hidden Costs

These quality costs have been hidden in our usual accounting methods or have not been defined as quality costs by the organization. Failure costs in nonmanufacturing operations (service, white-collar work, sales, legal, marketing, and distribution) are often well hidden. Many companies use the iceberg analogy to show this, the traditional quality costs are measured and visible, but below the surface of the water are hidden numerous other quality costs (see Figure 2).

The largest hidden costs of quality are non-value-added work. These are work processes that are done without thinking, usually “because we have always done it this way”. Unnecessary forms, wasteful meetings, moving items from place to place, multiple approvals, unused databases, many surveys, and unread reports are all examples of non-value-added work. Reengineering activities were often quite successful by just eliminating or reducing these items. Other hidden costs include areas not usually thought of as quality costs such as employee turnover, expediting, and answering customer questions due to poor communications or instructions. Employee turnover can create huge costs. Recruiting, training, and lost productivity can result in some of the highest quality costs to organizations with positions where skilled workers are needed.

4 Cost of Quality

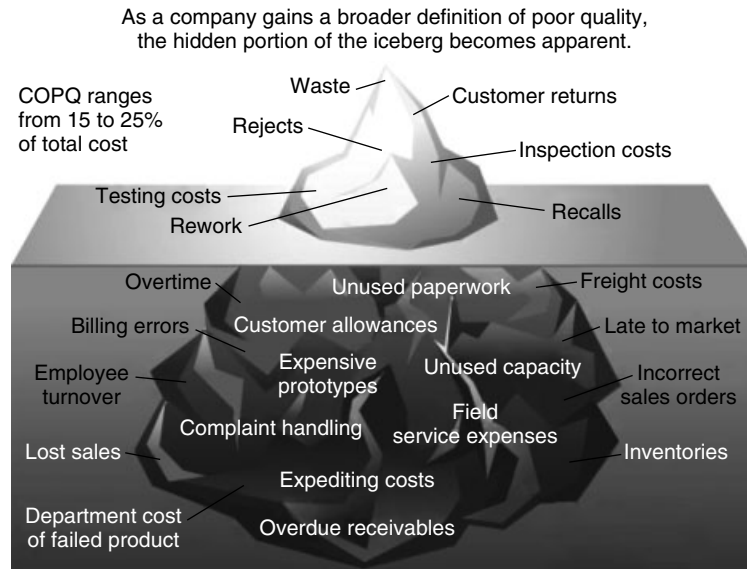


Figure 2 The hidden costs of quality

Lost Income

Another quality cost is income lost due to poor quality or the reputation for poor quality. Customers may perceive quality problems and not purchase, or purchase and return the product due to problems. Customers may demand their money back for services not at the quality level they expected. Companies may voluntarily rebate part or all of the purchase price due to product or service failures. Customers may change brand loyalties and switch to a competitor's product permanently creating long-term income losses. Customers may tell many others about their dissatisfaction leading to even greater lost income.

Customers' Costs

Leading companies are now beginning to understand that quality costs extend beyond their own costs to the costs the customers endure. The failure of a product may be totally covered by a warranty, but the customer has other costs due to time lost, the unavailability of the product, and other costs. Business customers may lose time, production, and even their customers due to failures of components, products, or services they have purchased. Sophisticated customers are now thinking in terms of total

life-cycle costs and product failures, lost time, extra inventories, delays, and other costs are considered in addition to initial price when making their purchase decisions. Customers who wait for hours in dentist's or doctor's offices or emergency rooms, or are put on hold on the telephone for many minutes, or wait for hours at home for deliveries are changing service providers.

Costs to Society

These are the costs to the community at large. Some of these are obvious: oil tanker spills in Alaska, disastrous side effects of drugs, nuclear power plant accidents. Other costs are harder to measure, but we are now beginning to be more sophisticated in estimating environmental costs, safety costs, time losses, security costs, and other costs incurred by various quality failures. Sometimes these costs are partially recovered through fines, lawsuits, and other means but usually they are borne by society as a whole or indirectly paid by higher taxes, regulatory fees, or other means. This new understanding of quality costs is including the producers' costs, the customers' costs, and the costs to society at large through the entire product life cycle.

Making a Cost of Quality Study

Frank Gryna in *Juran's Quality Handbook*, 5th Edition, gives a useful guide for studying quality costs and identifying opportunities for major savings. He suggests the following sequence:

1. Review the literature and consult others who have done similar studies.
2. Select part of the organization to serve as a pilot site.
3. Discuss the objectives of the study with key people in the organization.
4. Collect whatever cost data are conveniently available from the accounting system and use this information to convince senior management of the need for a complete study.
5. Make a proposal to senior management for the complete study.
6. Publish a draft of the categories defining the quality costs.
7. Finalize definitions and secure management approval.
8. Secure agreement on responsibilities for **data collection** and analyses.
9. Collect and summarize the data, involve the finance department as much as possible.
10. Present the costs to management along with specific recommendations and the results of an improvement project using the data.

Further Reading

- ASQC. (1987a). *Quality Costs: Ideas & Applications*, ASQ Quality Press, Milwaukee, Vol. 1.
- ASQC. (1987b). *Quality Costs: Ideas & Applications*, ASQ Quality Press, Milwaukee, Vol. 2.
- Campanella, J. (1990). *Principles of Quality Costs*, ASQC, Milwaukee.
- Groocock, J.M. (1974). *The Cost of Quality*, Pitman Publishing, New York.
- Gryna, F.M. (1999). Quality and costs, in *Juran's Quality Handbook*, J.M. Juran & A.B. Godfrey, eds, McGraw-Hill, New York, pp. 8.1–8.26.
- Sörqvist, L. (1998). *Poor quality costing*, Doctoral Thesis, Royal Institute of Technology, Stockholm.

A. BLANTON GODFREY, AND JOSEPH D.
MOORE

Failure Modes and Effects Analysis

Introduction

Failure modes and effects analysis (FMEA) is a method of investigation for determining how a product, process, or system might fail and the likely effects of particular modes of failure. For example, a product, machine, or structure might fail physically as a result of a faulty part or user behavior. A manufacturing, business, or administrative process might fail operationally as a result of poor personnel training, faulty control, faulty design, or faulty equipment. In all of these cases and others, FMEA can be used to help assess the possible ways in which failures might occur, assess the magnitude of the effects of failure, uncover the possible cause or causes of failure, and understand what can be done to prevent such failures or mitigate the likelihood of their occurring.

Having originated in the U.S. military in the 1940s, FMEA today is an analytical tool widely employed in approaches to quality like ISO 9000, ISO/TS 16949, Six Sigma, and Design for Six Sigma (DFSS). FMEA is both a qualitative and quantitative technique. Ideally, a team undertaking an FMEA of an existing product or process would have historical data or data from monitoring systems that would indicate possible modes of failure and their most important causes. However, in many cases there may be no valid measurement system in place and, in the case of new products or processes, no historical data. In those cases, the team must determine subjectively, based on their knowledge and experience, the likely modes of failure, the practical consequences of failure, and potential causes of failure.

The team then evaluates the degree to which current control mechanisms or procedures are likely to detect the cause of failure before it occurs or before its effects occur, thus allowing time to take corrective action. For example, bald automobile tires indicate that the tires should be replaced, thus preventing a failure. But for other modes of failure, such as a blowout due to a sudden puncture, it is impossible to detect the causes prior to failure.

The team numerically rates each of those three elements: effects, causes, and detection. On a scale

of 1–10, an effect is rated for severity, where 10 is “catastrophic”. For causes of failure, the team rates likelihood of occurrence, with 10 being “highly likely”. For detection, the team rates the unlikelihood of detecting a cause or failure mode far enough in advance to prevent its effect, with 10 indicating that it is virtually impossible to detect before the effect occurs. The three numbers are then multiplied for each cause–effect–detection combination, of which there may be several for each failure mode, to arrive at a *risk priority number* (RPN):

$$\begin{aligned} &(\text{Severity of effect}) \times (\text{Likelihood of occurrence}) \\ &\times (\text{Unlikelihood of detection}) = \text{RPN} \end{aligned} \quad (1)$$

For example, an effect rated 10 for severity, 10 for likelihood of occurrence, and 10 for unlikelihood of detection (where 10 indicates *no* likelihood of detection) yields an RPN of 1000 – the worst possible case. Conversely, ratings of 1, 1, and 1, respectively, yield an RPN of 1, indicating a cause–effect–detection combination of no concern. RPNs can fall anywhere along this scale of 1–1000, and the higher the RPN, the greater is the cause for concern. Therefore, the FMEA team usually first addresses the potential failures of greatest concern. Corrective actions are identified and implemented. Subsequently, the FMEA ratings can be repeated to determine if the corrective actions have been effective, enabling the team to reduce levels of risk for the most critical factors over time.

Applications

Existing Products or Processes

The most common application of FMEA is in the improvement of existing products and processes. The analysis may be undertaken as part of a process improvement effort such as a Six Sigma program or as an *ad hoc* piece of work. For example, a team of investigators considering an existing product would typically begin with the question: *how could this product fail?* They might consider possible modes of failure both in the factory prior to release of the product as well as failure in the customer’s hands, and the effects of failure, known as *liability*. Here the RPN is helpful in attacking the areas of most likely failure and least available control or detection.

2 Failure Modes and Effects Analysis

Design of New Products or New Processes

FMEA can also be used during the design stage of a new product or process to determine the possible failure modes, their effects, and the most efficient approaches to reducing the risk of those failures. In this context, FMEA is often one of the analytical tools used in DFSS, a structured, risk-management approach to product and process development that aims to design-in a Six Sigma level of performance from the inception of a product or process.

To analyze the potential design of a new product, design engineers can use FMEA to determine the ways in which the product could fall out of specification and become unsuitable for release. They can also analyze the ways in which the product could fail in the hands of the customer. For example, possible failure modes, such as overheated material or improper production line speeds that could drive the product out of specification during production, are compiled and their causes are determined. Because the product is not in full-scale production, a relevant set of preventive controls and detection mechanisms can be audited or created. In many organizations, as there are only limited resources for these development efforts, RPNs can be used to determine which areas to attack and how much effort to invest in preventing each failure mode and the related out-of-specification condition.

Post-Failure Analysis

Following a failure, FMEA can be used to perform a “postmortem” of the event. In this case, FMEA provides a format for compiling the failure modes present, linking them to the known effects, and then auditing the effectiveness of controls. Because the effects are known and have already occurred, the RPN is not relevant in this instance.

FMEA as a Statistical Tool

Although FMEA is not a purely statistical tool, it is consistent with what may be more broadly characterized as statistical thinking, the three elements of which are process, variation, and data. Processes produce variation, and data are used to understand the sources of variation and eliminate it or keep it within acceptable limits. In 1996, the American Society for

Quality defined statistical thinking as a philosophy of learning and action based on the following fundamental principles:

- all work occurs in a system of interconnected processes
- variation exists in all processes
- understanding and reducing variation are keys to success.

It is those principles that underpin the FMEA approach to analyzing and improving processes. However, because FMEA depends on the FMEA team’s subjective knowledge of the process under investigation, the results of the analysis are only as good as that process knowledge.

The FMEA Session

Prior to undertaking FMEA, it is necessary to complete a process map of the process under study, which could include a specific product under investigation. A process map is a step-by-step pictorial sequence of a process, showing process inputs, process outputs, and processing steps. In addition, any specifications or fitness-for-use issues should be available for study. For example, for a product under study, it is helpful to understand the quantitative customer specifications and the more qualitative fitness-for-use criteria like reliability and durability. In assembling the team, it is advisable to include members representing all aspects of the process. For example, if a product is under investigation, the team should include members from product design through development and, finally, to operations.

For maximum efficiency, each session of an FMEA should be restricted to no more than a few hours. Typically, two to four sessions are required to complete an FMEA. At the beginning of a session, the team should develop scoring keys, including clear definitions of the numbers 1 through 10 for each of the factors of the RPN. These definitions should be reviewed each time the FMEA is restarted.

If there are numerous inputs for the process, the cause and effect (C&E), customer, or the super-focused approach, not the comprehensive approach, should be used. The C&E and the super-focused methods utilize a smaller number of inputs, which have been prioritized as to their relationship with various output measures. The customer method also

has a smaller number of inputs – those that have a direct relationship with very specific customer needs. The comprehensive approach analyzes each input with no prioritization, which increases the time devoted to the FMEA and leads to many inputs of lesser risk being studied.

When analyzing each individual failure mode, it should be remembered that a failure mode may have more than one effect and it may have more than one cause. Assuming a one-to-one relationship between failures and effects leads to missed causes of failure modes. (Because effects are generally easier to compile than causes, there is a far less chance of such assumptions leading to missed effects.) The failure to uncover causes can contribute to poor results, since, obviously, no action can be taken against them.

As controls or detection mechanisms are put in place, the RPN should be rescored. When the rescoring takes place, the RPN should be lower as a result of the controls, the lower RPN means that there is less risk to the process. Reduced process risk leads to more stable and improved performance, which is manifested in various output metrics. These output metrics can then be used as part of another process improvement initiative, and these metrics will be monitored within the project life and potentially, afterward as well. Improved metric performance can be measured more easily than iterative RPN rescoring over time, and they should be inversely related.

FMEA in Action: A Business Process Example

The mechanics of performing an FMEA for a business process can be illustrated by following its steps through a typical process improvement project, in this case, a faulty accounts payable process. In an accounts payable process, an invoice accompanies goods that are sent to the purchaser. The invoice is matched with the goods and a method of payment is set in motion to pay the invoice within the prescribed terms.

As with any transactional process, there are multiple possibilities for failure. Consider, for example, just one subprocess in the overall accounts payable process: invoice reconciliation. The goal of this subprocess is to reconcile invoices on a one-to-one basis with the goods received. As depicted in the typical FMEA table and rating chart in Figure 1, the FMEA

team begins by asking what can go wrong with the input for the process. In this case, one possibility is the arrival of an incomplete invoice. The team then asks what the effect might be of that particular mode of failure. One possible effect is the incurring of penalties for late payment.

Next, they ask what process variables could cause this failure mode – in this case, they believe that it is the incorrect placement of vendor information, such as shipping addresses, in the master vendor list. They then consider the current controls that are relevant to the prevention of this cause of failure and discover that although a process exists for signing up new vendors through a web-based tool, there is no comparable oversight for data on older vendors.

The team then scores on a scale of 1–10 the severity of the effect, the likelihood of its cause occurring, and the unlikelihood that current controls can detect the problem before its effects ensue. It should be noted that the rating of each of these three elements is performed at the end of the FMEA process, not at the time each element is first encountered in the investigation. Further, the highest effect of a failure mode drives the rating score or RPN. When there is more than one effect for the failure mode, the process of reducing causes of the failure mode will make all effects less frequent.

In this example, the team concludes that a late payment and the resulting penalties rate a 10 for severity. The team then rates the likelihood of an incomplete invoice occurring and assigns it a 6. They then calculate the unlikelihood of detection of the problem by current controls, with 10 indicating that there is no likelihood at all. Because there is virtually no real process in place to detect the problem, it is assigned an 8. The team then calculates the RPN:

$$\begin{aligned} (\text{Severity} = 10) \times (\text{Likelihood of occurrence} = 6) \\ \times (\text{Unlikelihood of detection} = 8) \\ = \text{RPN of 480} \end{aligned} \quad (2)$$

The team might then similarly investigate other possible combinations of effect–cause–detection, calculate the RPNs for each, and pursue corrective actions in the descending order of the RPNs.

4 Failure Modes and Effects Analysis

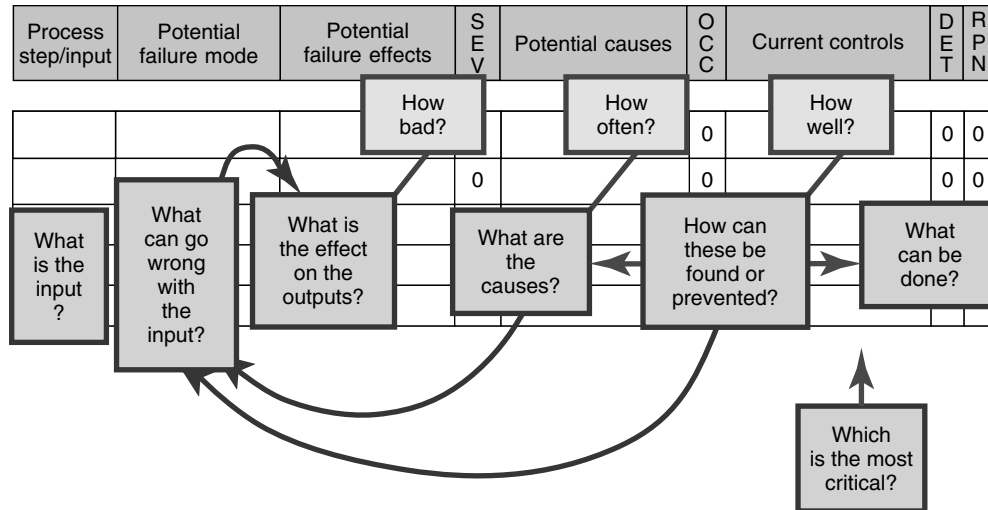


Figure 1 FMEA table and rating chart

FMEA in Action: A Manufacturing Example

The mechanics of performing an FMEA in manufacturing can be illustrated by following its steps through the solution to an equipment problem in a chemical manufacturing process.

A key production area is consistently hampered by a piece of equipment having pressure build-up issues as a result of severe fouling. The equipment stands at the end of a multi-step process where chemical solutions are made, then deposited on another material, and finally dried before final handling. The problem occurs at the drying stage. The team begins by conducting a Six Sigma investigation with many variables being studied in a hope that this equipment can be optimized to reduce the pressure build-up issues. After establishing a context with an overall process map and clarifying some of the output variable impact through cause-and-effect analysis, the team undertakes an FMEA.

In studying areas throughout the process during an FMEA session and feeling confident about the ability to detect and control problems in many of those areas, the team determines detection/control scores of 2 for them, leading to low overall RPNs. However, one major process area that has not been previously analyzed leads to highly technical discussions. The team realizes that it has little confidence that these issues can be detected and determines a score of 8

for unlikelihood of detection. Because the likelihood of occurrence of the problem at this level of analysis is unknown, the team, as a precautionary measure, scores this element as 10. The severity of the effect is determined to be 8, resulting in an RPN of 640, as depicted in Table 1.

As a result of this FMEA session, the team changes focus to an upstream technical problem versus a downstream optimization issue. After some study the phenomenon causing the pressure issues is located in a preparation process. Controls are put in place and the subsequent rise in the capacity of the process, system availability, and product quality confirms that the team has found the root cause of the pressure problems.

History

FMEA has its origins in the United States military, which first formalized it on November 9, 1949, in MIL-P-1629, "Procedures for Performing a Failure Mode, Effects and Criticality Analysis" (FMECA). Its stated purpose was to classify failures "according to their impact on mission success and personnel/equipment safety". FMECA provided the foundation for the military's subsequent refinements of the technique, embodied in MIL-STD-1629 (November 1, 1974), which encompassed FMECA, FMEA, and variants including damage

Table 1 Determining RPNs in a manufacturing process

Process step/input	Potential failure mode	Potential failure effects	SEV ^(a)	Potential causes	OCC ^(b)	Current controls	DET ^(c)	RPN
Raw material addition	Raw material (RM) flow rate is too high	Off-specification material	10	Weigh cell failure	1	Digital control system (DCS) program/lab verification	2	20
Raw material addition	RM flow rate is too high	Downstream equipment plugging	4	Weigh cell failure	1	DCS program/lab verification	2	8
Raw material addition	Moly flow rate is too low	Off-spec material	10	Weigh cell failure	1	DCS program/lab verification	2	20
Additive preparation	Concentration is too high	Precipitation leading to downstream equipment plugging	4	Weigh cell failure	1	Basic labs and appearance checks	2	8
Additive preparation	Concentration is too low	Off-spec product	10	Weigh cell issue	1	Basic labs and appearance checks	2	20
Additive preparation	Wrong or other concentration	RM does not dissolve	5	Operator loads wrong amount of additive bags	2	Basic labs and appearance checks and marked additive bags	2	20
Process setup (method A or method B)	Operational issues with method A	Downstream equipment plugging and product quality issues	8	RM physical property issues	10	Basic labs and appearance checks and equipment pressure	8	640
Process setup (method A or method B)	Method B operates as method A	Product quality	8	Incorrect time intervals	2	Sampling, operator checks on DCS, DCS alarms	2	32

^(a)Severity of Effect^(b)Likelihood of Occurrence^(c)Unlikelihood of Detection

modes and effects analysis (DMEA). MIL-STD-1629A was released on November 24, 1980, and was updated by MIL-STD-1629A change note 1 on June 7, 1983, and MIL-STD-1629A change note 2 on November 28, 1984. MIL-STD-1629 was canceled by MIL-STD-1629A change note 3 on August 4, 1998.

FMEA subsequently spread to other fields. In the 1960s, it was adopted by the aerospace industry and applied to the Apollo missions. In the 1970s, its utility was broadened to include the analysis of software-based systems. Also in the 1970s, Ford Motor Company used FMEA to troubleshoot potential design and safety failures following highly publicized lawsuits about the company's Pinto model. In the early 1980s, the automotive industry as a whole adopted FMEA when Ford, the Chrysler Corporation, and General Motors Corporation jointly developed the QS 9000 standard, known in other industries as ISO 9000, to standardize supplier quality systems. To comply with QS 9000, automotive suppliers were required to use FMEA. This initial effort culminated in SAE J-1739, an industry-wide FMEA standard adopted by the Society of Automotive Engineers in 1994.

Although FMEA is a relatively new tool in the field of process failure analysis, its use today includes manufacturing and transactional areas; process, product and software development; general process improvement and specific product failure work, and strategic planning review and capital project implementation.

Further Reading

- Automotive Industry Action Group. (1995). *Potential Failure Mode and Effects Analysis Reference Manual*, 2nd Edition, AIAG, Southfield.
- Dyadem Pres. (2003). *Guidelines for Failure Mode and Effects Analysis (FMEA), for Automotive, Aerospace, and General Manufacturing Industries*, Dyadem Press, Ontario.
- Krasder, G.D. (2004). *Failure Modes and Effects Analysis: Building Safety into Everyday Practice*, HCPro, Marblehead.
- McDermott, R., Mikulak, J. & Beauregard, M.R. (1996). *The Basics of FMEA*, Productivity Press, New York.
- Stamatis, D.H. (2003). *Failure Mode and Effect Analysis: FMEA from Theory to Execution*, 2nd Edition, Quality Press, Milwaukee.

RONALD D. SNEE AND WILLIAM F.
RODEBAUGH

Quality Function Deployment

Introduction

Quality function deployment (QFD) is a planning tool that allows the flow down of high-level customer needs and wants through to design parameters (DPs) and then to process variables (PVs) critical to fulfilling the high-level needs. By following the QFD methodology, relationships are explored between quality characteristics expressed by customers and substitute quality requirements expressed in engineering terms, also known as *critical-to-satisfaction (CTS) characteristics* within the terminology of Design for **Six Sigma**. In the QFD methodology, customers define their wants and needs using their own language, which rarely carries any actionable technical terminology. The voice of the customer (VOC) can be grouped (using affinity diagram) into a list of needs and wants, which can be used as input to a QFD's relationship matrix house of quality (HOQ).

Knowledge of customer's needs and wants is paramount in designing effective products and services with innovative and rapid means. Utilizing the QFD in the context of design for Six Sigma (DFSS) methodology allows the developer to attain the shortest development cycle while ensuring the fulfillment of customer needs and wants as depicted in Figure 1.

History of QFD

QFD was developed in Japan by Dr. Yoji Akao and Shigeru Mizuno in 1966 but was not westernized until the 1980s. Their purpose was to develop a quality assurance method that would design customer satisfaction into a product before it was manufactured. One of the first published applications of QFD happens to be in Bridgestone's Kurume factory by Kiyotaka Oshiumi who followed the steps of Dr Akao. For 6 years the methodology was enhanced from the initial crude usage of Kiyotaka Oshiumi of Bridgestone Tire Corporation. Following the first publication of *Hinshitsu Tenkai*, i.e. quality deployment in Japanese, by Dr. Yoji Akao [1] in 1972 the pivotal development work was conducted at Kobe Shipyards

for Mitsubishi Heavy Industry. Stringent government regulations for military vessels coupled with the large capital outlay per ship forced Kobe Shipyard's management to commit to upstream quality assurance. The Kobe engineers drafted the QFD matrix, which relates all the government regulations, critical design requirements, and customer requirements to company technical controlled characteristics of how the company would achieve them. In addition, the matrix also depicted the relative importance of each entry, making it possible for important items to be identified and prioritized to receive a greater share of the available company resources. The Kobe application resembles an application that is very similar to how QFD is defined today.

In 1978 the detailed methodology was published [2, 3] in Japanese and was translated to English in 1994.

QFD Methodology

QFD is accomplished by multidisciplinary teams using a series of matrices, or phases of the HOQs, to deploy critical customer needs throughout the phases of the design development. Figure 2 depicts the rooms and the generic HOQ. Going room by room, we see that the input is into Room 1 in which we answer the question "what". These "whats" are either the results of VOC synthesis for HOQ 1 or a rotation of the "hows" from Room 3 into the following HOQs. These "whats" are rated in terms of their overall importance and placed in the Importance column. This column contains the importance rating for each "what" or voice listed in Room 1 from the customer perspective.

Next we move to Room 2 and compare our performance and the competitors' performance against these "whats" in the eyes of the customer. This is usually a subjective measure and is generally scaled from 1 to 5. A different symbol is assigned to the different providers so that a graphical representation is depicted in Room 2. Next we must populate Room 3 with the "hows". For each "what" in Room 1 we ask "How can we fulfill this?" We also indicate which direction the improvement is required to satisfy the "what": maximize, minimize, or target. This classification is in alignment with the **robust design** methodology and indicates an optimization direction. In HOQ1, these become "How does the customer

2 Quality Function Deployment

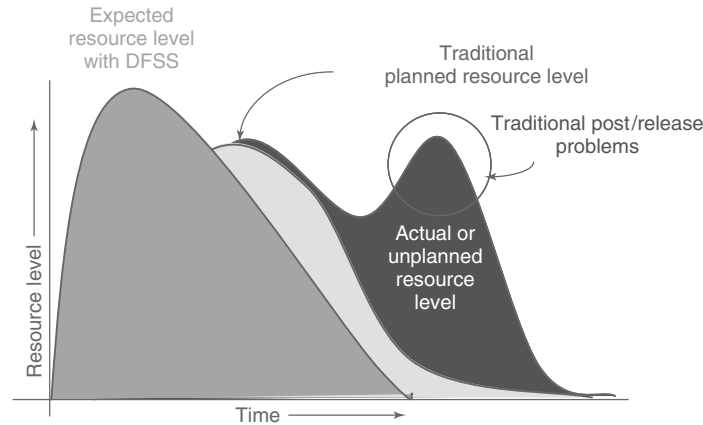


Figure 1 The time-phased effort of QFD *versus* traditional design

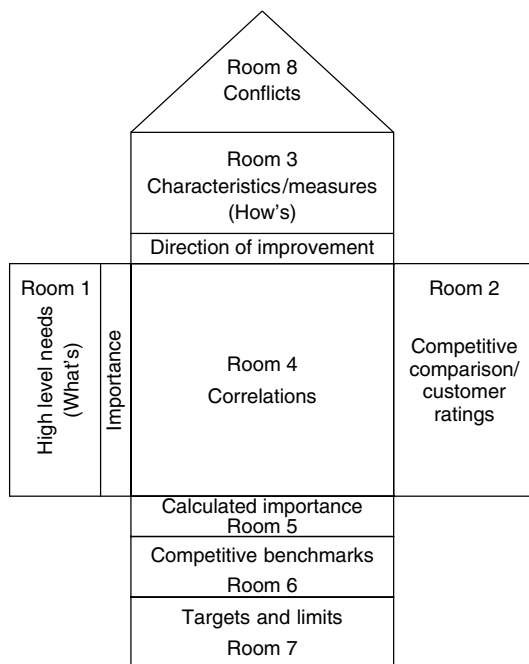


Figure 2 QFD house of quality

measure the what?” In HOQ1, we call these *CTS measures*. In HOQ2, the “hows” are measurable and solution-free functions that are required to fulfill the “whats” of CTSs. In HOQ3 the “hows” become DPs and in HOQ4 the “hows” become PVs. A word of caution: teams involved in designing new services or processes often times jump right to specific solutions

in HOQ1. It is a challenge to stay solution free until HOQ3. There are some rare circumstances where the VOC is a specific function that flows straight through each house unchanged.

Within Room 4 we assign the weight of the relationship between each “what” and each “how” using 9 for Strong, 3 for Moderate, and 1 for Weak. In the actual HOQ these weightings will be depicted with graphical symbols, the most common being the solid circle for Strong, an open circle for Moderate and a triangle for Weak (Figure 3).

Once the relationship assignment is completed, by evaluating the relationship of every “what” to every “how”, the calculated importance can be derived by multiplying the weight of the relationship and the Importance of the “what” and summing for each “how” or CTS characteristic. This is the number in Room 5. For each of the “hows” we can also derive quantifiable **benchmark** measures of the competition and our company; in the eyes of industry experts; this is what goes in Room 6. In Room 7, we can state the targets and limits of each of the “hows”. Finally, in Room 8, often called the *roof*, we assess the

Strong	●	9
Moderate	○	3
Weak	▼	1

Figure 3 Rating values for the strength of a “what”–“how” relationship

interrelationship of the “hows” to each other. If we were to maximize one of the “hows”, what happens to the other “hows”? If both “hows” improve, then we classify their relationship as a synergy, i.e. improving one of the “hows” results in improvement of the other “how”, both in their respective direction of goodness, whereas if it were to move away from the direction of improvement, their relationship would be classified as a compromise. Wherever a relationship does not exist it is just left blank. For example, if we wanted to improve customer intimacy by reducing the number of call handlers then the wait time may degrade. This is clearly a compromise. Although it would be ideal to have **correlation** and regression values for these relationships, often they are just based on common sense or business laws. This completes each of the

eight rooms in the HOQ. The next step is to sort on the basis of the importance in Room 1 and Room 5 and then evaluate the HOQ for completeness and balance.

The QFD methodology is deployed through a four-phase sequence shown in Figure 4. Each phase is a replica of the QFD HOQ captured in Figure 2 with some incorporated changes due to a structured cascading process. The four phases are

- Phase 1 – critical-to-satisfaction planning – House 1
- Phase 2 – functional requirements planning – House 2
- Phase 3 – DP planning – House 3
- Phase 4 – PV planning – House 4

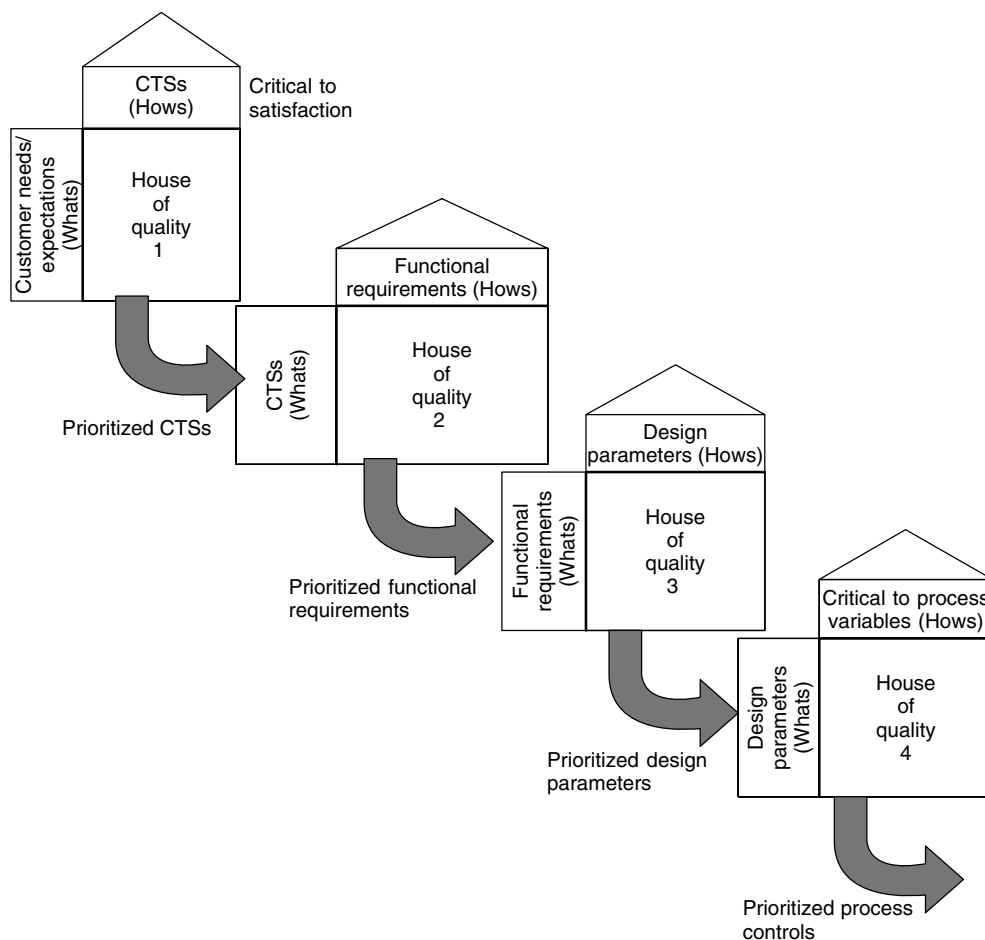


Figure 4 The four phases of QFD

4 Quality Function Deployment

Each of these four phases deploys the HOQ with the only content variation occurring in Room 1, the VOC room, also known as the “whats”, and Room 3, the design substitute quality characteristics that answer the “whats”, also known as the “hows”.

QFD Example

There was a business that purchased \$3 billion of direct materials (materials that are sold to customers) and \$113 million of indirect materials and services (materials and services that are consumed internally in order to support the customer facing processes) and yet the business had no formal supply chain organization [4]. The procurement function existed in each functional group so operations purchased what they needed as did human resources, legal, and information systems. This distributed function resulted in a lack of power in dealing with suppliers and mostly supplier favorable contracts and some incestuous relationships (lavish trips and gifts given to the decision makers).

The corporate parent required that all procurement be consolidated into a formal supply-chain organization with sound legal contracts, high integrity, and year-over-year deflation on the total expenses. The design team was a small set of certified Black Belts and a master Black Belt who were being repatriated into a clean-sheet supply-chain organization.

The design team tasked with the project identified the customers and established their wants, needs, delights, and usage profiles. In order to complete the HOQ, the team also needs to understand the competing performance and the environmental aspects that the design will operate within. The design team realized that the actual customer base would include the corporate parent looking for results, a local management team that wanted no increased organizational cost and the operational people who needed materials and services. Through interviews laced with a dose of common sense, the following customer wants and needs (the “whats”) were derived:

- Price deflation
- Conforming materials
- Fast process
- Affordable process
- Compliance
- Ease of use

The next step was to determine how these would be measured in the customer’s eye. The CTS measures (the “hows”) were determined to be

- Year-over-year price deflation (%)
- Defective deliveries (%)
- Late deliveries (%)
- Cost per transaction (\$)
- Cost to the organization per order (\$)
- Speed of order (hours)
- Number of compliance violations
- Number of process steps

Next the team used the QFD application, Decision Capture,^a to assign the direction of improvement and importance of the wants/needs, the “whats” or VOC, and the relationship between these and the “hows”, or CTS characteristics. Figure 5 shows the results. We can see that price deflation is moderately related to percentage of defective deliveries and number of compliance violations, it is weakly related to percentage of late deliveries and strongly related to year-over-year price deflation. We see that price deflation is extremely important and the direction of improvement is to maximize it.

In Figure 5, we can see that the importance of the CTSs has been calculated. The method of calculation is as follows for the top sorted CTS:

Year-over-year price deflation = 9 (strong) × 5 (extremely important) + 1 (weak) × 5 (extremely important) + 3 (moderate) × 4 (very important) + 1 (weak) × 3 (somewhat important) = 65. Thus, the CTS importance score is a product sum of the relationships in Room 4 (in the respective CTS column) by the corresponding “whats” importance rating listed in Room 1. In effect, it is a weighted sum by the importance to the customer. The other columns are similarly calculated.

The upper and lower halves of phase 1 HOQ are shown in Figures 6 and 7, respectively. Figure 6 shows the competitive ratings of the wants/needs (Room 2) and the technical ratings of the CTSs (Room 6). For this purpose the team used Competitor 1 as the incumbent process and used a sister business with a fully deployed supply-chain organization for Competitor 2. The synergies and compromises have also been determined.

The next step would be to assess the HOQ for any “whats” with only weak or no relationships. In this example no such case exists. Now we look for

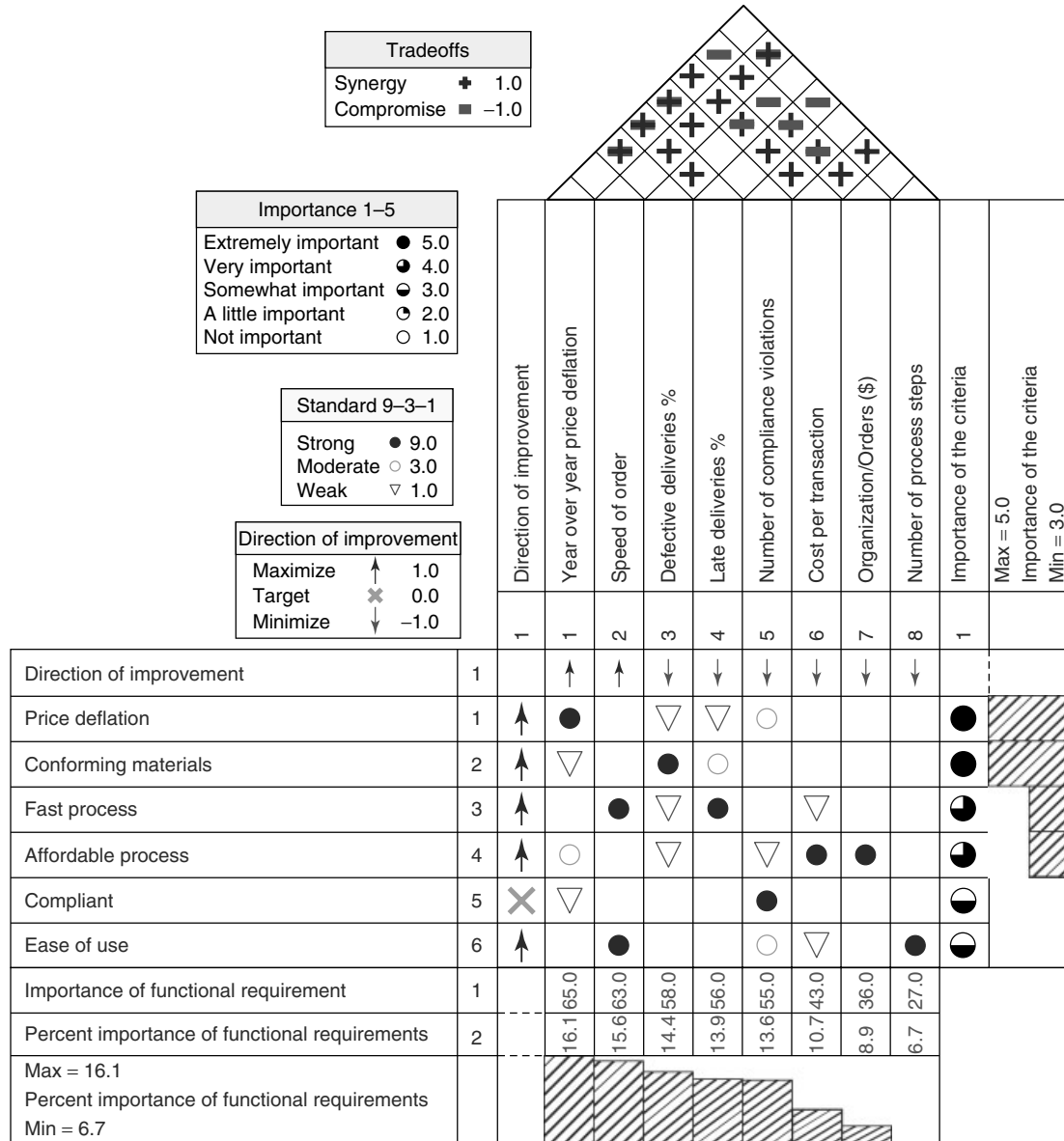


Figure 5 Partially completed phase 1 QFD HOQ

“hows” that are blank or weak and also see there are none of these.

Summary

QFD is a planning tool used to translate customer needs and wants into focused design actions. This

tool is best accomplished with cross-functional teams and is key in preventing problems from occurring once the design is operationalized. The structured linkage allows for rapid design cycle and effective utilization of resources while achieving Six Sigma levels of performance.

To be successful with the QFD, the team needs to avoid “jumping” right to solutions and process

6 Quality Function Deployment

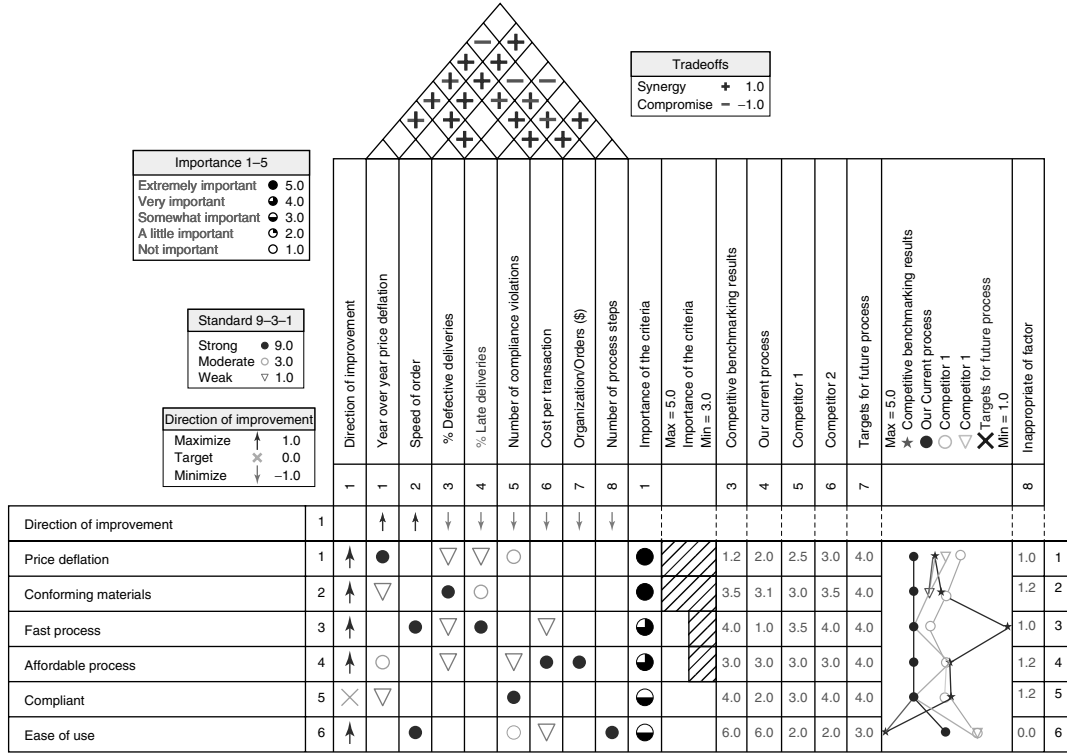


Figure 6 Upper half of phase 1 QFD HOQ

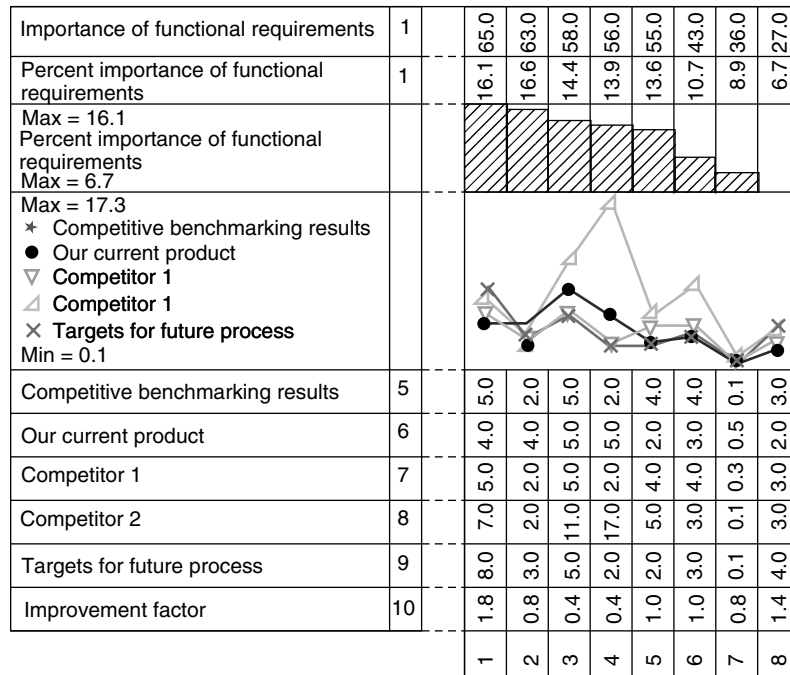


Figure 7 Lower half of phase 1 QFD HOQ

phase 1 QFD HOQ and other phases (Figure 4) thoroughly and properly before performing detailed design.

It is important to have the correct VOC and the appropriate benchmark information. Also a strong cross-functional team willing to think out of the box is required in order to truly obtain Six Sigma-capable products or processes. From this point the QFD is process driven, but it is not the charts that we are trying to complete; it is the total concept of linking the VOC throughout the design effort.

End Note

^a. Decision Capture is a software application from International TechneGroup. www.decisioncapture.com

References

- [1] Akao, Y. (1972). New product development and quality assurance – quality deployment system (in Japanese), *Standardization and Quality Control* **25**(4), 7–14.
- [2] Shigeru, M. & Akao, Y. (eds) (1978). *Quality Function Deployment: A Company Wide Quality Approach (in Japanese)*, JUSE Press.
- [3] Shigeru, M. & Akao, Y. (eds) (1994). *QFD: The Customer-Driven Approach to Quality Planning and Deployment* (Translated by Glenn H. Mazur), Asian Productivity Organization, ISBN 92-833-1122-1.
- [4] El-Haik, B. & Roy, D.M. (2005). *Service Design for Six Sigma: A Roadmap For Excellence*, John Wiley & Sons.

BASEM EL-HAIK

Cause-and-Effect Diagrams

Introduction

Statisticians have always been careful not to confuse **correlation** with causation (see e.g., Cox [1]). A famous example derived from statistics on the population of Oldenburg in Germany and the number of observed storks in the city in 1930–1936, demonstrates why correlation does not imply causation [2]. Figure 1 is a scatter plot of population size *versus* number of storks (see **Correlation**).

If we look at the data in Table 1, we quickly realize that time is a lurking variable and that both variables in the scatter plot grow with time, hence the correlation. Since storks need chimneys to make their nests, an increase in the number of buildings, due to an increase in population, explains why more storks are nesting in Oldenburg.

Correlation is therefore not causation, and scatter plots cannot be assumed to describe cause-and-effect

relationships. **Causality** is, however, a basic element of the scientific method and management skills. Establishment of causality relies on a combination of axiomatic thought and empirical evidence derived from **observational studies** or designed experiments (see **Assignable Cause**; **Measures of Association**; **Causality**). Cause-and-effect diagrams are used to present such relationships. We will present three such diagrams: Ishikawa diagrams, structural equation models, and **Bayesian networks**.

Ishikawa Diagrams

In the summer of 1943, at the University of Tokyo, Dr Kaoru Ishikawa was explaining to engineers from the Kawasaki steel works how various factors can be sorted out and related in specific ways. This was the origin of the cause-and-effect diagram later called the *Ishikawa diagram*. The Ishikawa diagram, also called a *fishbone diagram*, came into wide use in the Japanese industry and became a critical tool for **quality control** and quality improvement all over the world. It is one of the seven essential tools for quality

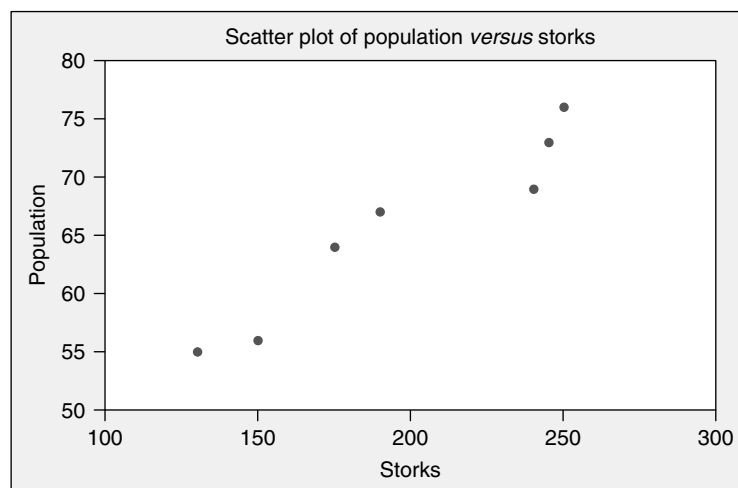


Figure 1 Scatter plot of number of storks *versus* population size

Table 1 Population and number of storks, by year

Year	1930	1931	1932	1933	1934	1935	1936
Population in thousands	50	52	64	67	69	73	76
Number of storks	130	150	175	190	240	245	250

2 Cause-and-Effect Diagrams

and process improvement also called the *magnificent seven* [3]. The Ishikawa diagram is used to identify, explore, and display possible causes of a problem or event. A typical diagram will have four to six main branches for causes affecting one type of event. Traditional classifications of these branches are five Ms: Manpower, Machines, Materials, Measurements, and Methods or, in administrative areas, the four Ps: Policies, Procedures, People, and Plant. A modern variation on the five Ms uses six categories: People, Machines, Materials, Measurements, Methods, and Environment. These are only suggestions typically used to jump-start the brainstorming session used to complete the diagram. Figure 2 is an example from an improvement project initiated by the dean of a school of management to improve the operation of his office, and, in particular, the scheduling of appointments [4]. The improvement team, consisting of the dean, two professors at the school and the two secretaries at

the dean's office, brainstormed to list causes for an unusually high number of "collisions" in the dean's schedule. Many meetings with the dean were delayed or even canceled at the last minute. This problem led to complaints by the faculty and staff about the poor service level of the dean's office. The possible causes that were listed included the open-door approach, lack of meeting agendas, interruptions of all sorts, and no scheduled end time to meetings. Following some more analysis, the dean's office instituted a procedure where all meetings were scheduled with an agenda and an end time. A confirmation note, with this information, was sent ahead of time to the meeting participants. "Collisions" dropped by 50% and the dean's office became a model of service quality on campus. The secretaries even started planning teleconference meetings for specific topics at everyone's satisfaction. For more on Ishikawa diagrams see Ishikawa [3], Brassard [5], and Kenett and Zacks [6].

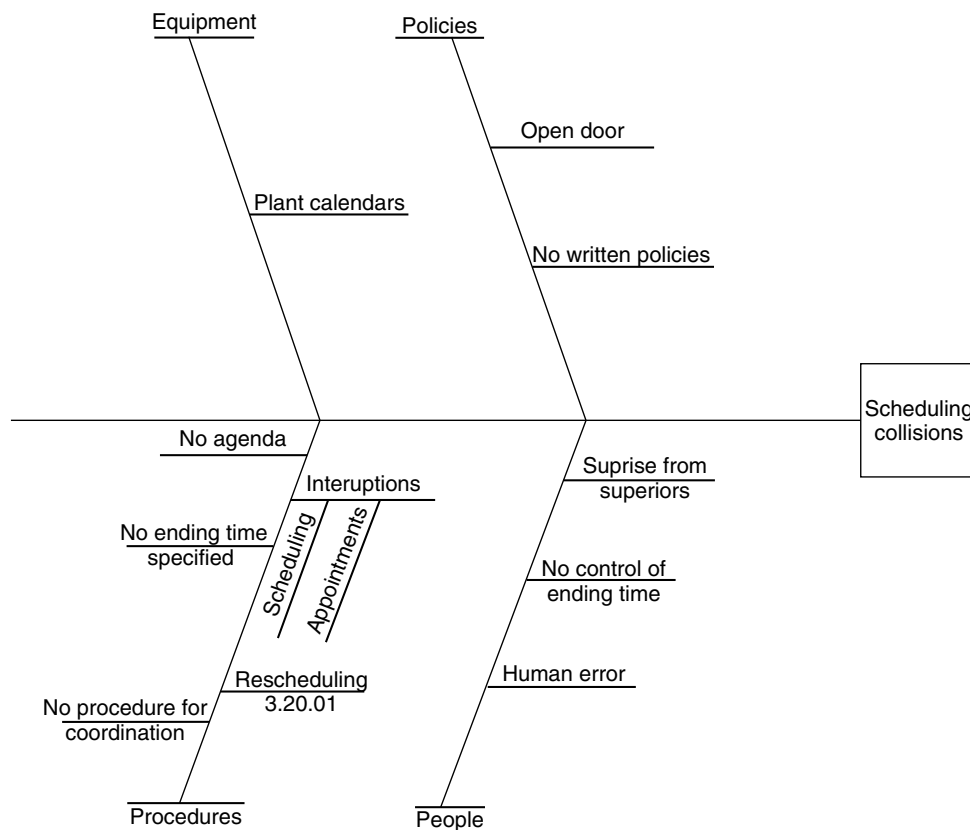


Figure 2 Ishikawa diagram of causes for scheduling collisions at a dean's office

Structural Equation Models

The geneticist Sewall Wright developed a “method of path coefficients” which is partly graphical and partly algebraic to explain genetic phenomena. Path analysis allows researchers to compute the magnitude of cause-and-effect relationships from correlation measurements, provided the path diagram represents correctly the causal process underlying the data [7, 8]. Later, economists developed a similar approach labeled structural equation modeling. We use here both terms interchangeably.

Structural equation modeling involves the specification of an underpinning linear regression-model incorporating the structural relationships between unobserved or latent variables, together with a number of observed or measured indicator variables.

Structural equation models assume there is a causal structure among a set of latent variables, and that the observed variables are indicators of the latent variables. The latent variables may appear as linear combinations of observed variables, or they may be intervening variables in a causal chain.

As latent variables are, by definition, unobservable, their measurement must be obtained indirectly. This is done by linking one or more observed variables to each unobserved variable. A fully specified structural equation model is potentially a complex interplay between a large number of observed and unobserved variables, and residual and error terms. For more on structural equation models see Bollen [9], Hoyle [10], and Kline [11].

The American customer satisfaction index (ACSI) is a prime application example of a structural equation model. ACSI uses telephone customer interviews, based on a structured questionnaire, as inputs.

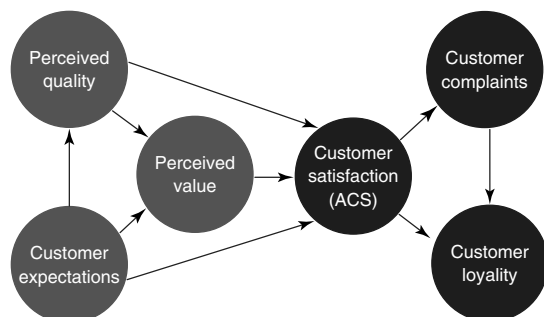


Figure 3 The American customer satisfaction survey index model

The latent variables in the ACSI model are presented in Figure 3 as a cause-and-effect model with indices for drivers of satisfaction on the left side (customer expectations, perceived quality, and perceived value), satisfaction (ACSI) in the center, and outcomes of satisfaction on the right side (customer complaints and customer loyalty, including customer retention and price tolerance). The indices are multivariable components measured by several questions that are weighted within the model. The questions assess customer evaluations of the determinants of each index. Indices are reported on a 0–100 scale. The survey and modeling methodology quantifies the strength of the effect of the index on the left to the one to which the arrow points on the right. These arrows represent “impacts”. The ACSI model is self-weighting, to maximize the explanation of customer satisfaction (ACSI) on customer loyalty. Looking at the indices and impacts, users can determine which drivers of satisfaction, if improved, would have the most effect on customer loyalty. For more on ACSI see Fornell *et al.* [12–15].

Structural equation models are also applied to integrated management models. An example of such an integrated model was implemented by Sears, Roebuck and Company into what they call the *employee-customer-profit model*. The cause-and-effect chain links three strategic initiatives of Sears: (a) a compelling place to work, (b) a compelling place to shop, and (c) a compelling place to invest. In order to push forward these initiatives Sears introduced an integrated structural equation model linking how employees felt about working at Sears, how employee behavior affected customers’ shopping experience and how customers’ shopping experience affected profits. The model identifies the drivers to employee retention, customer retention, customer recommendation, and profits. Specifically, the model predicts that five units increase in employee attitude drive a 1.3 unit increase in customer impression, which creates a 0.5% increase in revenue growth [16]. For more on integrated management models see Kenett [17] and Godfrey and Kenett [18].

Bayesian Networks

Bayesian networks combine Bayesian theory with graph theory for inferring probabilistic relationships among variables. They are used to handle complex

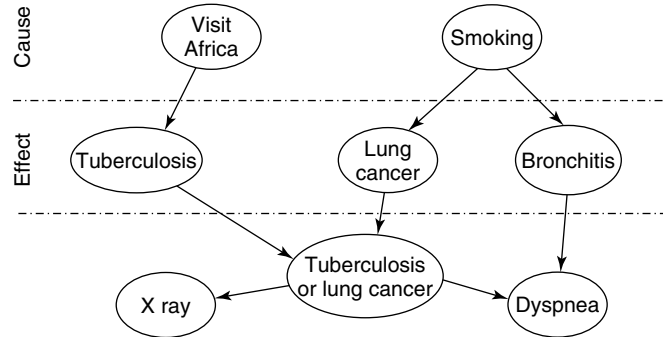


Figure 4 An example of causal network (adapted from Lauritzen and Spiegelhalter [19])

problems where variable interactions are too intricate for description by an analytic model. Bayesian networks represent current knowledge by graphically linking cause-and-effect variables. For example, consider a patient that visits a clinic with health related problems. The physician examining her wants to identify the disease and prescribe the correct therapy. Figure 4 is a graph representing cause-and-effect variables linked to tuberculosis and lung cancer. At the top are variables that describe the domain (“visit Africa”, “smoking”). The domain knowledge experts have linked variables describing an effect (i.e., “tuberculosis”) with possible causes (i.e., visiting Africa may cause tuberculosis).

Bayesian networks are directed acyclic graphs (DAGs) in which nodes represent random variables and arcs represent direct probabilistic dependencies among them. The structure of the directed graph provides insights into the interactions among the variables and allows for prediction of effects of external manipulation. Prediction consists of conditioning a causal variable to determine its effect on variables depending on it on the graph. Diagnosis is obtained by conditioning on an effect, in order to ascertain the profile of the causal variables linked to that effect. Conditional dependence probabilities, estimated from data, are used to generate a network representing learned causality. Bayesian networks represent the quantitative relationships among the modeled variables. Let $X_1, \dots, X_n$ be a set of variables and $P(X_1, \dots, X_n)$ be the joint probability defining the knowledge about the problem domain.

Numerically, a Bayesian network represents the joint probability distribution among the n variables. This distribution is described efficiently by

exploring the probabilistic dependencies among the modeled variables and the model size is reduced by conditional independence. Each node, X_i , is a random variable either with discrete or continuous states and is described by a probability distribution, $P(X_i | \text{parents}(X_i))$, conditional on its direct predecessors and quantifying parents’ effects on the node. Nodes with no predecessors are described by prior probability distributions. Then, by the Markov property and conditional independence, the joint distribution is simplified through conditional distributions to:

$$P(X_1, X_2, \dots, X_n) = \prod_{i=1}^n P(X_i | \text{parents}(X_i)) \quad (1)$$

A set of oriented arcs links a pair of nodes and identifies probabilistic dependencies. For more on Bayesian networks see Pearl [20–22], Cowell [23] and **Bayesian Networks**.

References

- [1] Cox, D.R. (1992). Causality: some statistical aspects, *Journal of the Royal Statistical Society. Series A* **155**, 291–301.
- [2] Box, G.E.P., Hunter, W.G. & Hunter, J.S. (1978). *Statistics for Experimenters: An Introduction to Design, Data Analysis, and Model Building*, John Wiley & Sons, New York.
- [3] Ishikawa, K. (1968). *A Guide to Quality Control*, Asia Productivity Organization, Tokyo.
- [4] Kelley, T.F., Kenett, R.S., Newton, E., Roodman, G.M. & Wowk, A. (1991). Total quality management also applies to a school of management, in *Proceedings of the 9th IMPRO Conference*, Atlanta.
- [5] Brassard, M. (1989). *The Memory Jogger Plus*, Goal/QPC.

-
- [6] Kenett, R. & Zacks, S. (1998). *Modern Industrial Statistics: Design and Control of Quality and Reliability*, 2nd Edition 2003, Chinese Edition 2004, Duxbury Press, San Francisco.
 - [7] Wright, S. (1921). Correlation and causation, *Journal of Agricultural Research* **20**, 557–585.
 - [8] Wright, S. (1934). The method of path coefficients, *Annals of Mathematical Statistics* **5**, 161–215.
 - [9] Bollen, K.A. (1989). *Structural Equation with Latent Variables*, John Wiley & Sons, New York.
 - [10] Hoyle, R.H. (1995). *Structural Equation Modeling: Concepts, Issues, and Applications*, Sage Publications, Thousand Oaks.
 - [11] Kline, R.B. (1998). *Principles and Practice of Structural Equation Modeling*, Guilford Press, New York.
 - [12] Fornell, C.G. (1992). A national customer satisfaction barometer: the Swedish experience, *Journal of Marketing* **56**, 6–21.
 - [13] Fornell, C.G., Johnson, M.D., Anderson, E.W., Cha, J. & Bryant, B.E. (1996). The American customer satisfaction index: nature, purpose and findings, *Journal of Marketing* **60**, 7–18.
 - [14] Fornell, C.G. (2001). The science of satisfaction, *Harvard Business Review* **79**, 120–121.
 - [15] Fornell, C.G., Mithas, S. & Krishnan, M.S. (2005). Why do customer relationship management applications affect customer satisfaction? *Journal of Marketing* **69**(4), 201–209.
 - [16] Rucci, A., Kim, S. & Quinn, R. (1998). The employee-customer-profit chain at sears, *Harvard Business Review* **76**, 83–97.
 - [17] Kenett, R.S. (2004). The integrated model, customer satisfaction surveys and six sigma, *Proceedings of the First International Six Sigma Conference*, Center for Advanced Manufacturing Technologies, Wroclaw University of Technology, Wroclaw.
 - [18] Godfrey, A.B. & Kenett, R.S. (2007). Joseph M. Juran, a perspective on past contributions and future impact, *Quality and Reliability Engineering International* **23**(6).
 - [19] Lauritzen, S.L. & Spiegelhalter, D.J. (1988). Local computations with probabilities on graphical structures and their application to expert systems, *Journal of the Royal Statistical Society, Series B: Methodological* **50**(2), 157–224.
 - [20] Pearl, J. (1988). *Probabilistic Reasoning in Intelligent Systems: Networks of Plausible Inference*, Morgan Kaufmann.
 - [21] Pearl, J. (1995). Causal diagrams for empirical research, *Biometrika* **82**, 669–710.
 - [22] Pearl, J. (2000). *Causality: Models, Reasoning, and Inference*, Cambridge University Press.
 - [23] Cowell, R.G., Dawid, A.P., Lauritzen, S.L. & Spiegelhalter, D.J. (1999). *Probabilistic Networks and Expert Systems*, Springer, New York.

RON S. KENETT

Factor Relationship Diagrams

The Fundamental Elements of DOEs

Statistical design of experiments (DOEs) are commonly applied to supplement engineering and process knowledge regarding product compositions and manufacturing processes. The fundamental purpose of DOE techniques is to manipulate product or process factors (X 's) across certain levels to understand their impact on critical performance parameters (Y 's). However, DOEs are often inappropriately resourced, applied, and implemented owing to lack of consideration of key elements of DOE planning. Simultaneously, software models of DOE results are often naively accepted based on a technical consideration of statistical significance without understanding the risks of extrapolating the DOE results into future process, service, or product. Many of the issues with supporting, planning, and investigating DOE results are not technical or statistical in nature; they are based on the need to appropriately question DOE plans and subsequent results prior to making decisions based on a proposed DOE strategy. "Most mistakes in the use of DOE deal with investigating the same factors that are always investigated, rederiving basic first principles that should already be understood, underestimating the impact of noise, and focusing on the tool rather than the questions driving the need for the tool" [1].

A DOE consists of a set of process or product parameters, the *design factors*, deliberately manipulated as factors in the experiment. Additionally, the DOE is influenced by other sources of variation that are not deliberately manipulated but that can impact the experimental results. These nonmanipulated factors are often referred to as *noise*. How these noise factors are managed during the experiment provides insight into the *inference space* of the DOE. How the noise and nonmanipulated sources of variation are managed during an experiment is typically referred to as the *experimental unit structure*. The unit structure consists of the experimental units to which experimental treatments will be applied, uncontrollable factors that cannot be manipulated yet are believed to have an impact on the experimental response(s), as well as sources of variation that are deliberately held

constant during the DOE. Together, the design factors and the unit structure comprise the experimental design.

The Role of the Factor Relationship Diagram

Much of the planning and analysis of DOE focuses on the factor effects, the levels of the factors, and the interactions amongst the design factors, while work to understand the contributions to changes in Y due to the unit structure sources of variation is given much less consideration. However, both statistical and practical conclusions, such as a test for statistical significance, depend as much upon the management of the unit structure during an experiment as it depends on the manipulated factor effects. The factor relationship diagram (FRD) is a schematic tool used to visually display the relationships between the potential sources of variation in a DOE. The FRD shows the relationship between the design factors and the unit structure and provides insight into how this relationship influences the information obtained from the DOE. The FRD was introduced by Wendy Bergerud to assist the experimenter in clarification of the relationship of design factors to their appropriate experimental units in order to increase the likelihood of correct statistical tests [2]. Bergerud's tool was modified and expanded by Sanders and Coleman to provide a visual means for understanding the identification of the experimental inference space and the potential risk involved with restrictions on **randomization** [3]. Additionally, the FRD is coded by font or color to visually differentiate between the design factors, the unit structure, and any lines of restriction in the experimental design.

The basic elements of the FRD are illustrated in the oversimplified example of Figure 1. The example is based on a batch, mixing process without reaction. The three design factors, mix time, blade speed, and mix temperature, are each manipulated in the DOE at two distinct levels, denoted by $-$'s and $+$'s. These design factors are shown in boldface type. Additional sources that can change from one treatment combination to another, thereby influencing the value of the response Y , are shown in italics below the design factor combinations. These sources of variation changing between treatments, such as measurement error and within batch inconsistencies,

2 Factor Relationship Diagrams

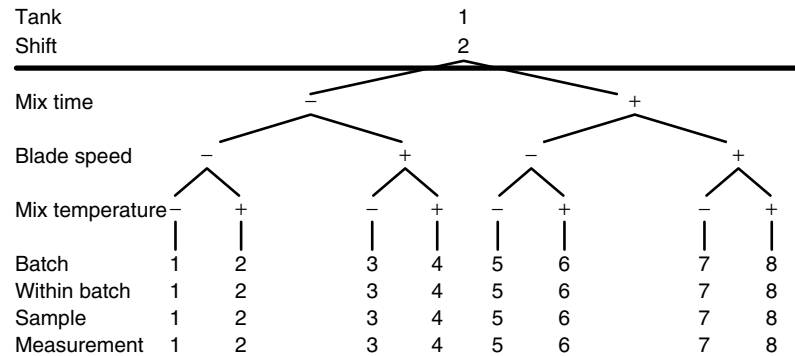


Figure 1 Simple FRD for a DOE on a batch mixing process

which are not a part of the design structure of the DOE describe the experimental error for this experiment. Hence, if this experiment is analyzed under the assumption of scarcity using normal **probability plots**, the unit structure at the bottom of the FRD provides insight into sources of variation that can also be creating movement in the experimental response. Also, based on the FRD, one quickly interprets that each experimental treatment requires a batch of product; hence, for a full **factorial** DOE with eight runs, eight batches of product are required. These batches represent the *experimental units*, or the division of material such that any two experimental units receive different combinations of design factors.

Hence, in planning any experiment, it is important to understand the inference space within which the DOE is to be run. Since DOEs necessarily take into account only a small subset of all factors potentially impacting the response of interest, the experimenter must understand those unit structure factors that are

held constant and impact the ability to extrapolate experimental results and those which are changing from one treatment to another impacting the values of the experimental response variable.

The FRD and Restrictions on Randomization

As previously noted, the FRD contains information on the design factors, the unit structure, and any lines of restriction in the DOE. A line of restriction exists when material, personnel, or equipment is not randomly allocated across the treatments or when the treatment combinations are not run in random order. Figure 2 shows an alternative DOE strategy where only four batches are required to complete the DOE. However, the resulting DOE is a **split-plot** design, where mix temperature and blade speed are whole-plot factors and mix time is a split-plot factor. The

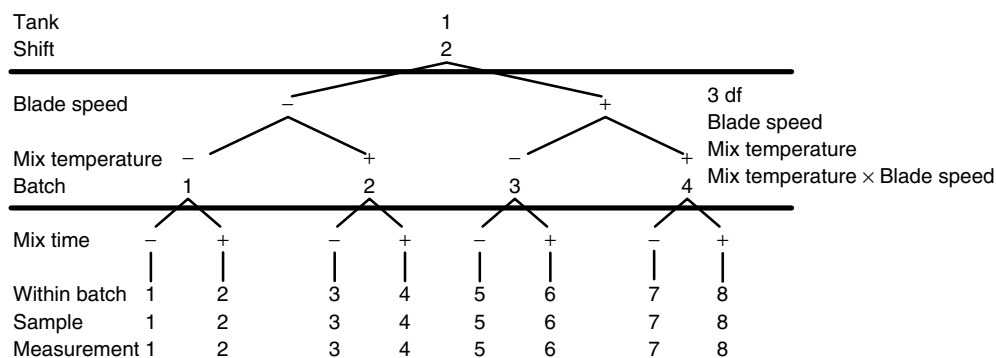


Figure 2 Alternative FRD showing a split-plot experiment

experimental units for the whole plot are the four batches, where the experimental unit for the split plot is within batch. The line of restriction (LOR) on the FRD indicates this partitioning of the unit structure in the experiment.

A common mistake is to ignore subsampling and the partitioned unit structure in the DOE of Figure 2 and, thus, to analyze the experimental results as if the DOE were completely randomized. When the restriction is not taken into account, the experimenter would assume seven **degrees of freedom** available for effect **estimation**. However, the FRD in Figure 2 quickly shows that there are only four whole-plot treatment combinations; hence, only three degrees of freedom are available for estimation of whole-plot effects. The remaining four degrees of freedom estimate the split-plot effects – mix time and its interactions with the whole-plot factors. Hence, the LOR provides a visual insight into the appropriate error structure of the experiment, assists in identifying the degrees of freedom available at each hierarchical level of the unit structure, and therefore encourages correct analysis and statistical testing.

Conclusions

The FRD provides a visual means for evaluating alternative DOE plans *a priori* to experimentation to compare potential cost of the proposed experiment with the information available from the experiment. The FRD also helps the experimenter to understand and communicate the inference space of the experiment. Since DOEs only take into account a subset of all of the factors potentially impacting the response of interest, the remainder of the factors are either intentionally neglected or not recognized as contributors to the experimental outcomes. These sources of variation in the unit structure are either held constant or they change across the experimental study, and their effects may be more important than those of the design factors being manipulated in the DOE [4].

Finally, the FRD is a powerful tool for helping the experimenter to understand the outcomes of a DOE. Possible outcomes of a DOE include, but are not limited to:

1. one or more factor effects are identified as statistically significant, or
2. no design factors or their interactions with other design factors are statistically important.

Any of these possible outcomes are not only impacted by the factor selection and the boldness of the levels at which the factors are set, but they are also determined by the management of other sources of variation during the DOE. For instance, if no factors are indicated to be statistically important, then there is evidence that the unit structure changing during the experiment created more variation than did the design factors. The FRD assists in understanding what sources did not vary and which sources of variation may potentially be more important to investigate in further studies.

In summary, the FRD is a visual tool for promoting the effective use of DOEs.

1. The FRD assists the experimenter in understanding the relationship between design factors and the unit structure. It therefore encourages correct statistical tests.
2. Multiple FRDs are useful for comparing multiple alternative DOE plans and communicating the resources required and the information available with each alternative plan.
3. The FRD helps to identify the inference space in which an experiment is run, and it helps to prevent the extrapolation of experimental results without appropriate validation.
4. The FRD visually illustrates the impact of restrictions of randomization in the experimental units, and it thereby encourages proper statistical analysis of DOE results.

References

- [1] Farrington, M. (2002). *Robust Design Training Manual*, Whirlpool Corporation.
- [2] Bergerud, W.A. (1996). Displaying factor relationships in experiments, *American Statistician* **50**(3), 228–233.
- [3] Sanders, D. & Coleman, J. (2003). Recognition and importance of restrictions on randomization in industrial experimentation, *Quality Engineering* **15**(4), 533–543.
- [4] Lynch, J.G. (1982). On the external validity of experiments in consumer research, *Journal of Consumer Research* **9**, 225–239.

CHERYL HILD AND DOUG SANDERS

Process Maps and Statistics

Introduction

A process map is an invaluable tool for directing the statistical study of a process. When we use statistical methods to gain knowledge of a process with the intention of improving that process, we are working with a process that is already in current operation. Consequently, the methods, materials, work practices, machine settings, and other operating conditions are already in place. At issue then is selecting what set of methods, practices, or parameters to change to improve operation. Our statistical study of the process can only be effective when it is coupled with knowledge of current process behavior. The process map is designed to capture, in the most effective manner, this kind of process knowledge.

In particular, the process map should include knowledge on the process that will guide process study and improvement. We should look to the process map to address the following required information:

1. What process output characteristics are critical for providing customer value?
2. What has been the past experience in maintaining appropriate variation and levels of critical characteristics?
3. What processing factors affect critical characteristics and are they or can they be managed?

The answer to the above questions will clearly require a mixture of organization, engineering, and statistical knowledge to answer. Organizing this information into a manner useful for process study is the aim of the process map.

How to Draw a Process Map

First, it is necessary to determine the scope of the process map. This will be a function of the traditional view of the process, the process owner's sphere of influence, and an expectation regarding the work necessary to improve the process. As a general comment, individuals constructing maps tend to make the scope

of the process too narrow [1]. Before finally deciding on the scope, we should consider those factors, which we are considering as initial inputs to the process. Are some of these inputs internal to the organization? If so, we should consider expanding the process map to include the activities that generated those inputs.

The second step in drawing a process map is to draw a flowchart or other graphical representation of the process. Instructions on flowcharting can be found in the literature [2]. Then, categorized inputs and outputs of the process are added. These inputs and outputs are similar to the kind of information shown on a **cause and effect** (C/E) diagram. The advantage of placing the inputs and outputs on a process map is that, where they effect the process is displayed. Consequently, the nature of how they impact the process can be better understood and **data collection** plans devised to evaluate their anticipated effects. Constructing the process map as described requires that the process be observed several times before the map can approach comprehensiveness and as work is done to improve and maintain the process, the process map will need to be revised to reflect these changes.

Further Characterizing the Inputs and Outputs

The critical outputs of a process (usually referred to as *the big Y's* in **Six Sigma** terminology) can be used to determine how the process is performing. Big *Y's* relate to the overall process and are typically those characteristics of prime importance to the customer. "Little *y's*", again in the parlance of Six Sigma, are intermediate outputs of a process, often measured at the end of process stages. The process map will capture our knowledge of what these intermediate outputs are and indicate whether these "little *y's*" are currently measured or not.

On a process map the inputs, or process "*X's*" (which include measurements made on incoming materials to the process, process methods, parameter settings, environmental factors, etc.) are categorized as controllable (*C*) or noise (*N*). Choosing to control an input and thereby labeling it with a "*C*" indicates that, given the current technology, resources, and process knowledge, these are the variables that need to be managed. Typically, these variables will have target values and requirements on their limits (possibly communicated as specifications). "Noise" refers to variables that either cannot be controlled

2 Process Maps and Statistics

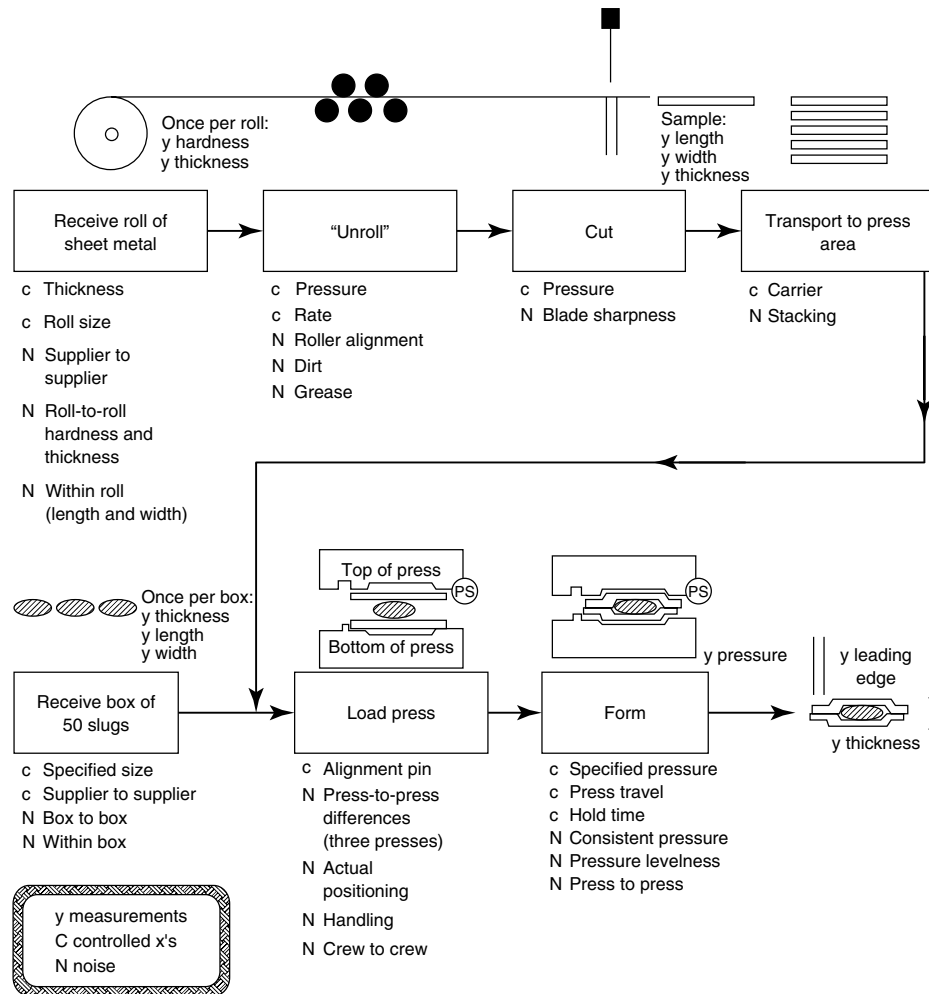


Figure 1 Process map of a crimp operation

(without excessive resources and technology) or are thought to have little influence on the output, so are not controlled. Noise could be intermittent disturbances or common cause, consistent, sources of variation.

Example

A better understanding of a process map is gained by looking at an actual map of a process. The process map shown in Figure 1 was used to describe a crimp operation. It was very helpful in thinking about how the variability in the crimp could be

set to a target with minimal variation. A successful crimp required considering the incoming raw material properties, standardizing the operator procedures, and determining correct machine settings. Examples of other project work supported by a process map can be found in Leitnaker and Cooper [1].

Summary Comments

As previously discussed, a process map is an invaluable tool for directing the statistical study of a process. Some specific examples of uses of the process map are listed, and discussed, below.

1. The process map exists to document current process knowledge. As such it should serve the following purposes.
 - (a) Capture current knowledge as to which inputs need to be controlled.
 - (b) Day-to-day management of processes often focuses on a few, conveniently measured variables, but effective process management requires retaining a more comprehensive picture.
 - (c) A common failure in process improvement work is to make changes to a process without regard to what currently exists and why it is there. The process map helps direct attention to important process characteristics that must be considered when process changes are contemplated.
2. The process map will also serve to enhance improvement activities. Because of the wealth of information captured on a process map, it can be used to generate and support actionable, manageable projects.
3. Most importantly, the process map should be used to provide guidance in planning statistical studies.
 - (a) When the objective of our process study is to evaluate the consistency of process outputs, either intermediate or final, the map provides specific information about the processing conditions across which the study should occur. For example, if an incoming raw material is known to be critical to a process output, a study evaluating the behavior of that output will need to cover a sufficient time across which that material will change.
 - (b) If our process study is to include a statistically designed experiment then the process map provides immediate candidates for possible factors to be included in the designed experiment.
 - (c) A more difficult task in designing an experiment is to understand the effect that process variation will have on experimental error. The process map provides not only a list of sources of variation that affect the process, but also provides guidance about where and at what frequency this variation will occur in the process.
 - (d) By necessity, the conditions seen during any statistical study will be a subset of the conditions that will occur in the future. In the same way that it is necessary to identify the frame and population in statistical surveys, an understanding of the process being studied is necessary for making inferences based on a statistical study. The inference space of a statistical study can be better understood by the use of a process map.

References

- [1] Leitnaker, M.G. & Cooper, A. (2005). Using statistical thinking and designed experiments to understand process operation, *Quality Engineering* **17**, 279–289.
- [2] Scholtes, P.R., Joiner, B.L. & Streibel, B.J. (1996). *The Team Handbook*, 2nd Edition, Joiner/Oriel, June 1.

Further Reading

Sanders, D., Ross, B. & Coleman, J. (1999). The process map, *Quality Engineering* **114**(4), 555–561.

MARY LEITNAKER AND TONY COOPER

Value Stream Analysis in Quality Management

Introduction

Value stream mapping first gained popularity with the book *Learning to See* [1], published in 1998. The book highlighted a technique that had become commonly used within Toyota, although Toyota does not call it value stream mapping. The technique has since been closely associated with **lean** manufacturing and the Toyota Production System. It has been used widely by those directing lean activities, but value stream mapping can also be utilized to aid in quality improvement efforts.

A value stream can be defined as the sum of all process activities needed to convert the primary product (or service) into a form ready to deliver to the customer. It includes both value-add and nonvalue-add activities. The technique of value stream mapping is a simple tool that seeks to map both material and information flow.

The technique begins (see the section titled “Creating a Current State Map”) with mapping the “current state” of a value stream. The current state map details how the value stream is operating currently. It helps visualize the overall, bigger picture of the process and thereby encourages improvements to the whole value stream, not individual processes. In addition, the technique is useful in generating discussions about how the process, or value stream, is performing.

Once the current state of the value stream is defined, plans can be developed to achieve an improved, or “future” state. This transformation from current state to future state is established by developing the improvement actions that are necessary to achieve this improved state.

In short, value stream mapping has several advantages that have contributed to its popularity and widespread use. Some of these include the following:

- **Helps identify waste** Value-add and nonvalue-add activities, or waste, within the value stream becomes apparent.
- **Helps establish a common language** Those associated with the process can use common terms and common map icons to discuss the value stream. Decisions regarding the value stream become apparent.

- **Results in an implementation plan to guide improvement activities** Activities necessary to achieve the future state become the implementation plan for the value stream improvement efforts. Specific activities can be tied directly to intended results.

- **Shows linkages** Value stream mapping uniquely links material flow to information flow. It also links the symptom of waste (e.g. excessive inventory) to the root causes (e.g., long changeover, poor process capabilities, etc.).

- **Provides data** Unlike other mapping methods, value stream maps possess considerable process data. For example, cycle times, **process yield** rates, and uptime ratios may be captured on the maps.

- **Easy to learn** The technique and an understanding of how to read maps can be taught at all levels of the organization in just a few hours.

- **A useful tool for manufacturing and non-manufacturing processes** Value stream maps can be created for many different types of processes, including manufacturing, healthcare, and transactional processes.

The technique of value stream mapping can also be used to enhance our understanding of a process and be used within **quality management** efforts. Some of the value stream wastes (e.g., excessive inventories, long lead times, etc.) may be due to process quality issues.

One initial step in process improvement is to map the process! In fact, in **Six Sigma** this is typically a recommended step in the “define”, or initial, stage of the project. Simply mapping the process, using the value stream mapping approach, can be used as one of the initial steps to improve the overall quality of the value stream. Sometimes the process quality inefficiencies become painfully obvious when just simply mapped.

In addition to helping identify improvement opportunities and projects, the value stream map aids in identifying the overall business or financial impact the improvements may have on the value stream. For example, decreasing the amount of excessive inventory may have an obvious financial impact of reduced carrying costs. Using the value stream map as an analysis tool may also note that this same project will also decrease the overall lead time to the customer

2 Value Stream Analysis in Quality Management

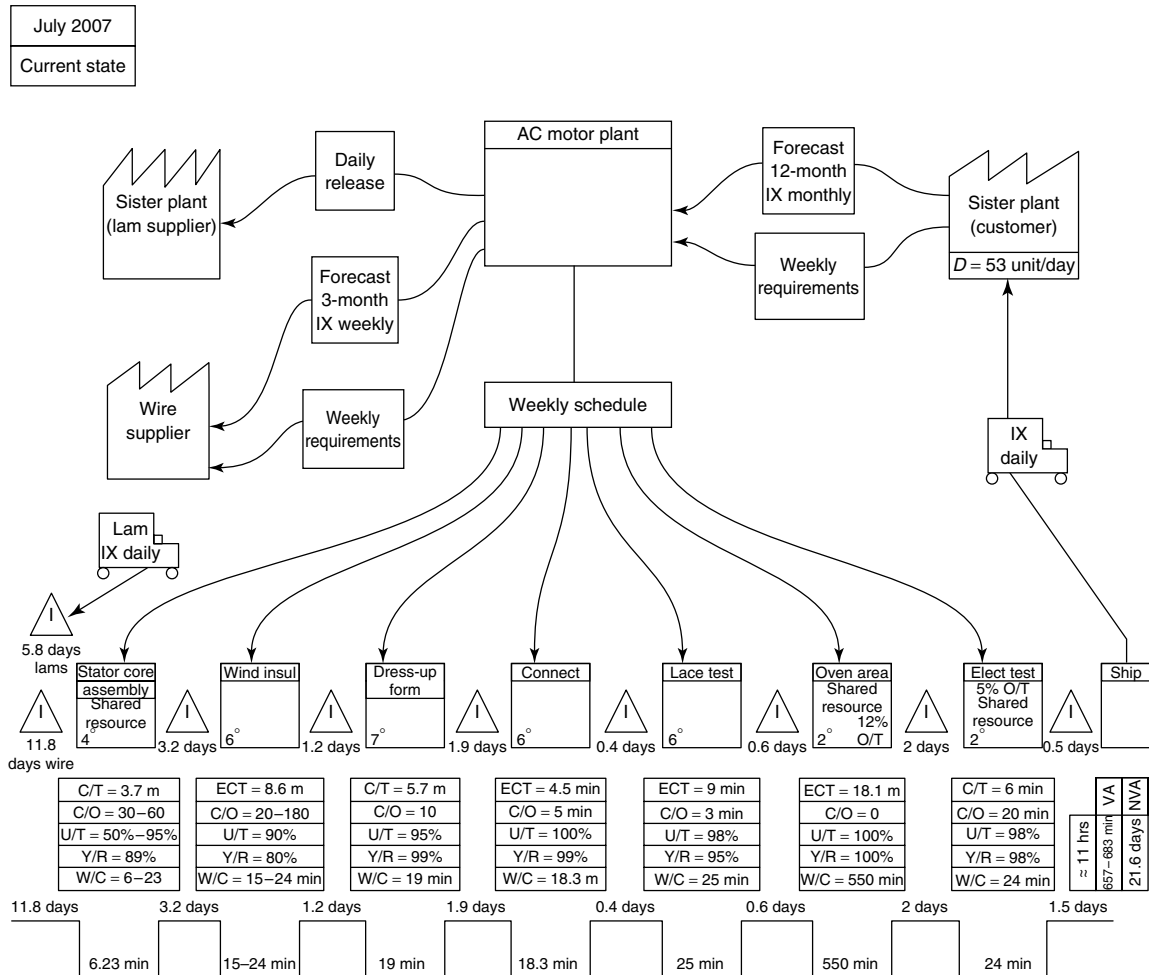


Figure 1 Current state map—manufacturing. This map details a manufacturing process for electrical motors. Notice the total value-add content for the value stream ranges between 657 to 683 min (approximately 10–11 h) and the nonvalue add is 21.6 days. Information flow is mapped across the top of the map

and thus could drastically influence sales. Thus, the impact from the increased sales may significantly exceed the gains from the reduced inventories. This analysis and realization may not be readily achievable without using the value stream mapping tool.

Creating a Current State Map

1. Define the scope of the map (i.e. the start and stop points for the value stream).
2. Determine the demand on the value stream (in terms of customer needs, e.g. patients seen per day).
3. Determine the operating characteristics of the value stream (e.g., how many shifts, the number of days worked each month, etc.).
4. Assemble a team of employees who know the value stream and/or know how to get to information and data concerning the value stream.
5. Walk the value stream, observing and collecting data such as cycle or process times, changeover times, yield rates, and number of operators in the process. Start the mapping process. Use paper and pencil. Stopwatch cycle times where possible. See the whole flow. Also look for waste (e.g., excess motion, transportation, scrap, etc.).

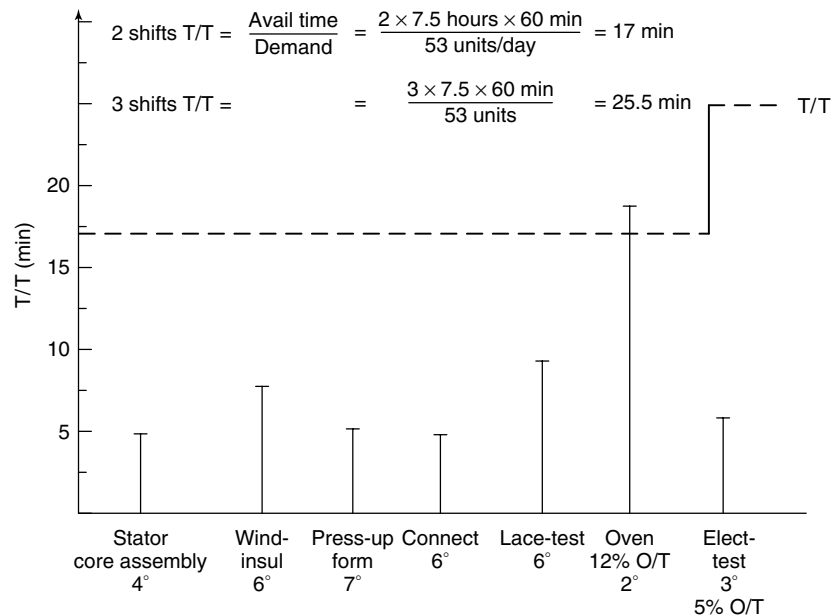


Figure 2 Operation balance chart. This diagram compares the process cycle time to *takt* time

Waste can be defined as activities the customer is not willing to pay for.

6. Gather information on inventory levels within the process. Often this is done by counting the inventory at different stages of the process (counting the inventory – if possible – is preferred to using computerized reports).
7. Determine and review the types of information obtained from the customer and how they drive activities within the value stream. Walk the information flow.
8. Reassemble the team and use the data collected to complete the current state map. Use icons to standardize and simplify the map.
9. Perform a “reality check” using observed cycle times against the calculated *takt* time for the value stream.

In most instances, the current state map can be created by the team in 4 hours or less.

Example Provided

In Figure 1 the team has created a current state map (see step 8 above). Notice the value added time (approximately 11 hours) is very small in comparison to the total non-value added time (21.6

days). This is very typical in both manufacturing and transactional processes. The value stream map should always be created in conjunction with the operation balance chart (see Figure 2). This chart compares each operation to the *takt* time.

Value Stream Mapping is also a tool that can be used in non-manufacturing too. Notice in Figure 3 a hospital has mapped an outpatient surgery process. Notice in this process too, the value-add time is low compared to the non-value add time.

Key Definitions and Calculations

Takt time.

The pace the value stream must achieve to meet customer demand. $Takt = \text{Available work time (e.g. seconds)}/\text{customer demand (units)}$.

Daily Demand Rate.

The number of units the value stream must produce each day to meet the customer demand. $\text{Daily demand rate} = \text{Monthly customer demand (e.g. units)}/\text{number of days worked per month}$.

Cycle time.

How often a part is produced (ideally timed by observation).

4 Value Stream Analysis in Quality Management

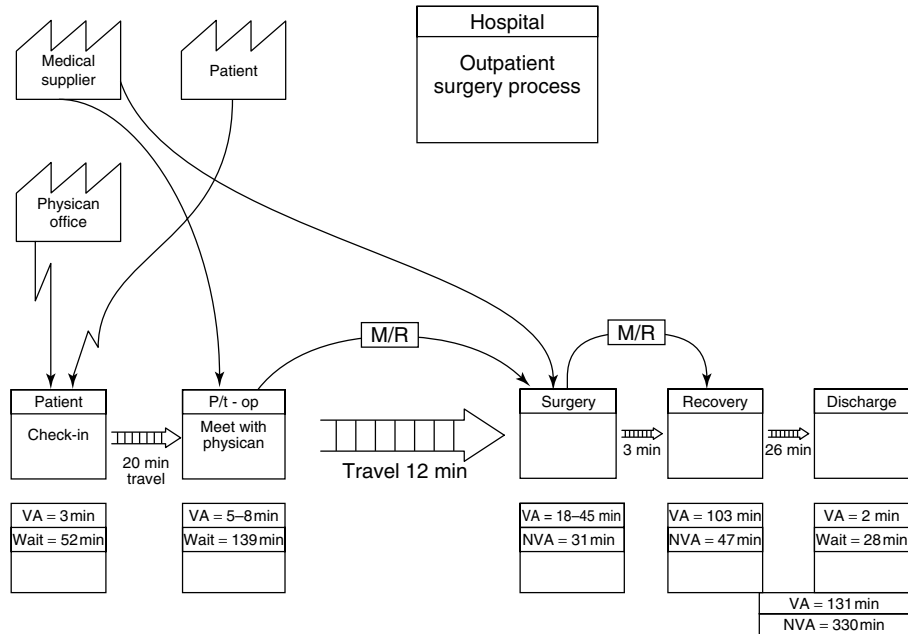


Figure 3 Current state map – healthcare process. VSM is often used to map processes other than manufacturing. This map is for an outpatient surgery process. Note again the VA time is small compared to the total time

Lead time.

The time required for one unit to move through the entire value stream. Lead time includes both value-add time and wait time.

Value-add time.

The time required to transform the product in a manner the customer is willing to pay for.

Changeover time.

The time required to change an element of the value stream (e.g. a machine tool) after one model has finished to the next model begins, “from last good piece to first good piece”.

First-pass yield rate.

The percentage of products that pass through each element of the value stream, the first time, without any type of quality problem (e.g., rework, sorting, etc.).

Uptime.

The amount of time the value stream element (e.g. machine) is available on average. Often in manufacturing this is a measure of machine reliability.

Working time.

Total planned time for work. Total seconds in a shift – any time not planned for work (e.g., breaks, shift meetings, etc.).

Pull.

Producing an item only when the next process consumes an item. A method of work that links production to the customer demand.

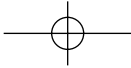
Push.

Producing an item without consideration of consumption from the next process. This often leads to the very common waste of overproduction.

Reference

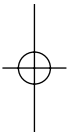
- [1] Rother, M. & Shook, J. (1998). *Learning to See*, Lean Enterprise Institute.

KEVIN GRAYSON



Creativity Tools

Introduction



How do people invent? Famous scientists and engineers sharing their memories, as well as psychologists studying the creativity process describe similar situations: an individual facing a difficult problem mentally explores various approaches, persistently trying and rejecting ideas until the right one comes. Psychologists call this process the trial-and-error method (T&EM).

T&EM has a long history. It was used to create the first stone axes, bows, guns, windmills, buildings, ships, and almost everything we can see around. Some of its results are astonishing, for example, Polynesian catamarans and old Chinese, Norwegian and Russian boats. However, archaeological research has shown that even 500 years ago these vessels were rather far from perfect since the builder tried only small changes in design. Some of these changes were unsuccessful causing fatal accidents and have been forgotten while others were successful and became standard. It was a long evolutionary process similar to the evolution of life.

With the acceleration of technological evolution, T&EM became less and less acceptable as a method of innovation and design. It is unreasonable today to build thousands of samples to select the best design of a modern aircraft or a steam machine. Engineering science has stepped in offering various methods for identifying the best design. These methods include scientific research, calculations, modeling, and computer simulation. As a result, engineering design today is a rather systematic, structured, and well-controlled process. However, the creative process of searching for new ideas still lacks these features.

In the typical creative process, people start by exploring apparent conventional solutions, slowly moving to the area of “wild” ideas under the weight of *psychological inertia*.^a

This T&EM for the creative process can be effective depending on how difficult the problem is. The effectiveness of this T&EM can be measured by the number of trials that have to be made to guarantee successful results. This number can vary within a wide range – from dozens for simple problems to hundreds of thousand for difficult ones. T&EM is sufficient for problems that do not require more than

10–20 trials; however, for difficult problems that require out-of-the-box thinking, T&EM leads to an unacceptable waste of time and effort.

In addition to low efficiency, T&EM contributes to poor problem statements. Often a problem is stated in incorrect format with a lot of unnecessary information while needed information is absent.

Until recently, T&EM deficiencies has been partially compensated by increasing the number of people working on a difficult problem. Nevertheless, since the mid-1950s it had become obvious that even the best use of human resources cannot satisfy the required pace of creative invention production. Accelerating technological evolution has demanded simple and affordable creative methods. This demand has lead to hundreds of them with different efficiencies, areas of application, and practical importance.

According to Higgins [1], creativity is the process of generating something new that has value. Creativity is an absolutely necessary ingredient for innovation, problem solving, and innovative design [2].

Creativity has always attracted scientists and practitioners, who have attempted to formalize it to the extent that its output becomes predictable and controllable, like in conventional engineering tasks. The history of the development of various creative problem-solving techniques extends back to the fourth century A.D. with Papp in Alexandria, Egypt. Papp was searching for a science of invention; he even found a name for this science: *heuristics*. For centuries, however, there was no real demand for such a science – until the nineteenth century when the industrial revolution started. Moreover, it seemed fairly obvious that creativity was a product of the human brain, and thus the main approach to creativity was focused on attempts to enhance the creative process by facilitating an individual’s mental processes, that is psychology-based approaches to creativity. In summary, these efforts were aimed at the following:

- unleashing natural creativity, eliminating mental blocks
- stimulation and mobilization of resources helpful for generating ideas by a group or an individual.

Later, knowledge-based approaches have been introduced including various analytical steps to organize, restructure, and involve available knowledge

2 Creativity Tools

and experience and eventually use of specially developed and structured external knowledge (innovation knowledge bases).

Classification of Creative Techniques

Depending on the methods and means used, the following categories for creative techniques have been identified [3]:

1. Conditioning/motivating/organizing techniques

The techniques, procedures and special conditions, and means belonging to this group help create an environment that facilitates the removal of various mental blocks and the unleashing of natural creativity.

Examples: Napoleon technique, listening to music.

Other techniques from this group merely suggest the use of various helpful tools including notebooks, stickers, boards, and flip charts.

2. Randomization

Since psychological inertia usually keeps an individual “inside the box” of their paradigms/perceptions/assumptions, helping an individual make more random attempts has been found to be very beneficial in solving difficult problems. Randomization makes the search more chaotic.

Example: brainstorming.

3. Partially focusing techniques

Many people have difficulty with random idea generation when no guidelines or focusing steps or subjects are offered. Special focusing techniques are used to help an individual focus on one issue at a time and avoid frustration. Focusing elements (steps) may be presented with or without any particular order (random focusing).

Examples: attribute listing or feature transfer; various sets of control questions.

4. Process-oriented (systematic) techniques

A process containing a set of steps to be followed in a specific order.

Example: morphological analysis.

5. Targeted techniques

These techniques offer single or multistep recommendations pursuing a predetermined, promising direction. This direction may be identified

as promising based on intuition, experience, or documented knowledge.

Example: problem reversal.

6. Evolutionary directed techniques

These targeting techniques offer directions according to statistically proven fundamental patterns of evolution.

Example: use of the pattern/lines of technological evolution for prediction of new generations of technological systems.

7. Innovation knowledge-based techniques

These techniques use structured knowledge derived from past human innovation experiences.

Example: contradiction table and 40 innovation principles.

Understandably, some techniques represent combinations of the approaches described above.

Over 90 creative techniques have been classified according to these seven groups [3].

Comparative Analysis of Selected Creativity Techniques

Significant work in the direction of unleashing and mobilizing natural creativity (psychological enhancement) was done by Osborn [4] (“Osborn’s direction”). Other important techniques that followed this direction are

- Synectics (Gordon [5])
- Fundamental design method (Matchett [6])
- Complex of techniques (De Bono [7]).

Results provided by methods of psychological enhancement strongly depend on personal issues and are not repeatable. Psychological means can stimulate a person to make the best use of his or her personal knowledge or that of a group of people. Unfortunately, this knowledge represents a negligibly small part of human innovative experience. Besides, psychological means cannot help organize the process of solving inventive problems as a logical step-by-step procedure (the way that conventional technological problems are usually solved).

The most successful technique for operating with available knowledge (*systematic techniques*) was offered by Miles (value analysis and engineering [8]). Other important techniques following “Miles’ direction” are

- Morphological analysis (Zwicky [9])
- **Quality function deployment** (Akoa [10])
- System thinking (Senge [11])
- **Failure Mode** and effect analysis (FMEA) [12].

The main weakness of this approach is lack of creative methods embedded in the working process. Often, known creative techniques are hardly compatible with the process.

Theory of Inventive Problem Solving (TRIZ)^b

New knowledge-based approach to creativity

Beginning in the mid-1940s, a new knowledge-based approach to creativity was introduced. First, attempts had been made to elucidate the intuition of successful inventors in a general way (Osborn's control questions, for example). The next step was made by Genrich Altshuller who embarked on a direct analysis of inventions documented in patents and other sources of technical information with the purpose of revealing so-called patterns of invention and of technological evolution (Altshuller's direction).

The basic advantages of the innovation knowledge-based techniques are the following:

- Accumulation of the best practices in creative problem solving is possible.
- Proven knowledge can be assessed.
- Results are repeatable and do not depend on personal (psychological) issues.

The most significant result of "Altshuller's direction" is the theory of inventive problem solving (TRIZ) [13].

Classical TRIZ

TRIZ is based on *patterns of technological evolution*, which represent its theoretical foundation. Among the basic discoveries of TRIZ the most important are:

- Any technical system develops according to certain patterns.
- The patterns of evolution for different systems have much in common.
- The patterns of evolution can be unveiled by researching the evolutionary history of a system (for the area of technology, this evolutionary

history is contained in the patent library and other sources of technical information).

- *Via* application of these patterns, one could accelerate the evolution of that system to its next generation.
- On the basis of these discovered patterns of evolution, universal methods for searching for new ideas can be developed.

TRIZ tools for systematic innovation include analytical tools that help understand and if necessary to reformulate the problem, and knowledge-based tools that represent best practices in solving problems extracted from patents and other sources of information. Analytical tools include ARIZ (Russian acronym for the *algorithm of inventive problem solving*) and *substance-field (S-F) Analysis*. Knowledge-based tools include patterns/lines of technological evolution, a number of *principles* (pathways) for eliminating (resolving) contradictions, *standard solutions*, and *innovation guides*.

Ideation TRIZ

Ideation TRIZ is a natural extension of Altshuller's TRIZ and has been in development since the mid-1980s targeting and addressing the main weaknesses of classical TRIZ which are as follows:

- TRIZ development was focused on studying patents without consideration of the feasibility and usefulness of solutions. Its practical application for the first 40 years of development was limited because of specific conditions in the former Soviet Union.
- Insufficient rigorousness of TRIZ tools (i.e., many analytical skills required for successful application of TRIZ tools to practical cases had not been transformed into documented rules, algorithms, and recommendations).
- Only a limited amount of the TRIZ knowledge base had been documented and was available for study and use.
- The tools did not support all stages of the problem-solving process. For example, problems had to somehow be preformulated in TRIZ terms before the tools could be applied.

Because of the above limitations, classical TRIZ was characterized by the following:

4 Creativity Tools

- Considerable education (from 100 to 250 h) was required to effectively use TRIZ.
- Extensive practice (from 1 to 5 years) was required to become self-sufficient in applying the methodology.
- Making TRIZ available for mass use posed an insurmountable challenge.

The main objective of Ideation research was to identify the most effective techniques covering all necessary components and issues associated with creativity and successful inventive problem solving, such as mobilization of personal capabilities, problem and system analysis, and the innovation knowledge-based approach, and integrate these techniques into a single, powerful methodology capable of addressing any problem or situation. Because of this integration, the following components have been selected:

From “Osborn’s direction”:

- methods of reducing psychological inertia
- team work.

From “Miles’ direction”:

- methods of collecting and organizing knowledge about a problem and the system in which it resides (enhanced and implemented in the Ideation Innovation Situation Questionnaire®)
- functional analysis (enhanced and implemented in the technique of *problem formulation*)
- morphological approach (used to ensure the exhaustiveness of the ideas developed).

From “Altshuller’s direction”:

- evolutionary approach (*patterns and lines of technological evolution*)
- innovation knowledge-based approach (various knowledge-based tools)
- TRIZ analytical tools.

Ideation TRIZ development has resulted in the following enhancements:

- Development of a streamlined, unified, and integrated process by which all types of problems can be treated in the same way, including the following steps [14]:
 - problem definition and documentation

- problem formulation
- identification and categorization of directions for innovations (solutions)
- development of solution concepts
- evaluation and planning of implementation.
- Development of new analytical tools to support relevant process steps, as follows:
 - The Innovation Situation Questionnaire® (ISQ) is used to facilitate understanding and documentation of the problem at hand.
 - Problem Formulator® is used to analyze a problem and develop an exhaustive set of potential directions for innovation.
- Development of a new integrated knowledge-based tool called the *system of operators*.
- Development of new TRIZ-based applications within inventive problem solving and beyond as follows:
 - Anticipatory Failure Determination (AFD®) including Ideation failure analysis and Ideation failure prediction [15]
 - **Directed Evolution**® – a methodology for predicting the next generations of various systems by inventing them [16]
 - evaluation and enhancement of intellectual property.

While being originated in technology, over the decades TRIZ principles and tools have proved their applicability to any area of human activity that requires creative solutions [17].

Other versions of TRIZ

For about 40 years, TRIZ development was confined within the Soviet Union borders. Russian “perestroika” in the mid-1980s opened the borders starting the process of TRIZ penetration in the Western part of the world. The new powerful approach to creativity and inventive problem solving started attracting various individuals, companies, and other institutions; however, the complexity of classical TRIZ became a serious hindering factor in its dissemination, causing frustrations and disappointments. To accommodate the new customer base, a number of simplified versions of TRIZ were developed (with various degrees of simplification) of which the best known is

structured inventive thinking (SIT) [18] and its modifications. Other attempts at simplification involved use of selected TRIZ tools like contradiction matrix and 40 innovation principles (original [19] and enhanced [20]). In all these cases, the simplicity was achieved at the expense of a serious limitation of TRIZ's problem solving power, eventually leading to frustration and disappointments.

A different approach was applied by Ideation International involving restructuring and enhancing classical TRIZ tools, developing of new ones and embedding them into software capable of supporting all steps of the innovation process [21].

TRIZ-based software

The first attempt to automate TRIZ was made by G. Altshuller in the mid-1960s when he built an electromechanical version of the contradiction matrix with the 40 innovation principles. Since the late 1980s, various TRIZ-based software products have been developed. The most comprehensive software products are the *TRIZSoft*[®] suite [22] developed by *Ideation International* and software developed by *Invention Machine Corporation* [23] (USA). Several other software packages are also available.

Summary

1. Over centuries, the T&EM was the main creativity tool. Moreover, creativity itself was associated with the T&EM process making the words *systematic innovation* sound like an oxymoron. While resulting in an enormous amount of incredibly valuable inventions enabling human survival, T&EM became less and less acceptable as a method of design and innovation with the acceleration of technological evolution.
2. As a response to the new requirements, hundreds of creativity techniques mostly based on psychological stimulation and thus lacking a systematic approach and structure have been developed in the last century.
3. In the mid-1940s, TRIZ, a new knowledge-base approach to creativity was introduced derived from the analysis of documented technological information and unveiled patterns of inventions and of technological evolution. TRIZ is the only innovation knowledge-based and evolutionary

directed technique that can provide the user with the accumulated power of the world's best inventors and innovations and the proven processes to invent in the most systematic and expeditious way. While being originated in technology, over the decades TRIZ underlined concept and tools have proved their applicability to any area of human activity that requires creative solutions.

4. To date, an integrated approach combining the most valuable achievements of the twentieth century, including Osborns, 'Miles', and Altshuller's directions has resulted in the Ideation TRIZ methodology that offers a unified and integrated process and tools capable of guiding and supporting individual and team creativity and innovation activity from problem definition to solution implementation. The process is supported by software that embeds all necessary analytical and knowledge-based innovation tools.

Acknowledgments

The authors would like to extend our gratitude to Rod Kornienko for the comprehensive collection of various creative techniques that enabled this research and to Zion Bar-El for encouraging us to write this paper.

End Notes

- ^a Psychological inertia could be defined as a complex of perceptions based on previous experiences, including one's background, education, and expertise. While being quite useful in performing routine operations, psychological inertia is the main reason for "in-the-box thinking" and is a critical hindrance in creative thinking.
- ^b TRIZ (pronounced "trees") is a Russian acronym for the theory of inventive problem solving.

References

- [1] Higgins, J.M. (1994). *101 Creative Problem Solving Techniques*, New Management Publishing Company.
- [2] Jones, J.C. (1997). *Design Methods*, John Wiley & Sons.
- [3] Zusman, A., (1999). Transactions of the Ideation Research Group. Overview of creative methods, *TRIZ in Progress*, Ideation International Inc.
- [4] Osborn, A.F. (1963). *Applied Innovation*, Scribener's Sons, New York.

6 Creativity Tools

- [5] Gordon, W.J.J. (1961). *Synectics: The Development of Creative Capacity*, Harper and Row, New York.
- [6] Matchett, E. & Briggs A.H. (1966). *Practical Design Based on Method (Fundamental Design Method)*, In the Design method (Edited by S. Gregory). Butterworths, London.
- [7] De Bono, E. (1970). *Lateral Thinking: A textbook of Creativity*, Penguin Books.
- [8] Miles, L.D. (1961). *Techniques of Value Analysis and Engineering*, McGraw-Hill.
- [9] Zwicky, F. (1969). *Discovery, Invention, Research—Through the Morphological Approach*, The Macmillan Company, Toronto.
- [10] Akao, Y. (1991). *Quality Function Deployment: Integrating Customer Requirements Into Product Design*, Productivity Press.
- [11] Senge, P.M. (1990). *The Fifth Discipline: The Art & Practice of The Learning Organization*, Currency Doubleday.
- [12] Stamatis, D.H. (2003). *Failure Mode and Effect Analysis: FMEA from Theory to Execution*, American Society for Quality, Quality Press, Milwaukee.
- [13] Altshuller, G. (1984). *Creativity as an Exact Science*, Gordon and Breach, Science Publishers.
- [14] Terninko, J., Zusman, A. & Zlotin, B. (1998). *Systematic Innovation: An Introduction to TRIZ (Theory of Inventing Problem Solving)*, St. Lucie Press.
- [15] Kaplan, S., Visnepolschi, S., Zlotin, B. & Zusman, A. (1999). *New tools for Failure and Risk Analysis*, Ideation International Inc.
- [16] Zlotin, B. & Zusman, A. (2001). *Directed Evolution: Philosophy, Theory and Practice*, Ideation International Inc.
- [17] Zlotin, B., Zusman, A., Kaplan, L., Visnepolschi, S., Proseanic, V. & Malkin, S. (2000). TRIZ beyond technology. The theory and practice of applying TRIZ to non-technical areas, in *Proceedings of TRIZCON*, Nashua, pp. 135–176.
- [18] Sichafus, E. (1996). *Structured Inventive Thinking*. The Industrial Physicist, 2(1).
- [19] Altshuller, G. (2002). *40 Principles: TRIZ Keys to Technical Innovation*, Technical Innovation Center, Worcester, <http://www.triz.org>.
- [20] Mann, D.L., Dewulf, S., Zlotin, B. & Zusman, A. (2003). *Matrix 2003: Updating the TRIZ Contradiction Matrix*, Creax Press.
- [21] Zlotin, B. & Zusman, A. (2005). Theoretical and practical aspects of development of TRIZ-based software systems, in Presented at the Conference TRIZFuture 2005, Graz, November 2005.
- [22] <http://www.ideationtriz.com>.
- [23] <http://www.invention-machine.com>.

BORIS ZLOTIN AND ALLA ZUSMAN

Directed Evolution®

Introduction

Ability to understand, predict, and control evolution of technology is becoming the centerpiece need of the informational era [1]. A proactive approach to all aspects of innovation requires the development of an integrated multistep process for “visualizing the future”. Directed Evolution,®^a (DE) is a technology involving systematic processes for building a sustained competitive advantage through the effective management of the evolution of various artificial (man-made) systems by utilizing evolutionary patterns for technologies, markets, business, social systems, and so on.

DE has two main roots: *technological forecasting* and the *theory of inventive problem solving*.

The first root of DE extends to the mid-1950s when technological forecasting [2, 3] – an approach for reckoning the future – was under development. From the late 1960s to the mid-1970s this resulted in the establishment of nonrelated techniques such as trend extrapolation, morphological modeling, the Delphi process, and others, all of which were based on probabilistic modeling of future characteristics of various systems. Although in the beginning much hope had been placed on the new method, most of the long-term forecasts that resulted have not proved correct, primarily due to the tools that were utilized to develop the forecasts.

The theory of inventive problem solving (TRIZ – Russian based acronym) was originated in the mid-1940s by Genrich Altshuller and represents structured systematic process for solving problem *via* unveiling and resolving contradictions (conflicts, dilemmas, creative challenges) supported by numerous tools [4]. These tools have been derived from the thorough analysis of worldwide innovative experience documented in patents and other sources of technical information. Unlike numerous methods for enhancing creativity based on trial and error and/or psychological stimulation, TRIZ is based on the assumption that inventions in technological systems appear not randomly, but rather in compliance with certain statistically proven *patterns of evolution* that could be revealed and utilized for organized and structured innovation. For example, one of the patterns states that every technological system in its

evolution becomes more dynamic, that is changeable, with higher degree of flexibility and variety. Unlike *trends* that could be short-lived (for example, the trend of increasing utilization of light metals could be in time replaced with the trend of increasing utilization of plastics), patterns of evolution reflect practically eternal customer requirements (like high performance, low cost, convenience, etc.). Typically, each pattern of evolution includes multiple *lines of evolution* – more detailed descriptions as to how this pattern could be realized step by step (for example, a mechanical system can become more dynamic *via* sequential introduction of a hinge, hinge mechanism, flexible elements, microlevel transformations, etc.). Over the years, numerous patterns and lines of technological evolution have been discovered through the analysis of hundreds of thousands of innovations spanning different areas of technology [5].

Since the mid-1970s, an entirely new approach based on utilization of predetermined patterns of evolution called *TRIZ forecasting* has been in development. Unlike traditional technological forecasting, TRIZ forecasting (as guided by patterns of evolution) offered directions together with proven standard ways on how they could be realized.

The main feature of DE is its proactive approach to the evolution of technology. Instead of making a prediction and waiting for it to be confirmed, the DE process uses numerous patterns of evolution for the purpose of identifying possible scenarios, analyzing them, selecting the most promising ones, and then planning the process of implementation. In other words, DE is a method to predict a future generation of a system by inventing it.

The DE technology was introduced in the early 1990s. Since then it has been under constant development by Ideation International’s research group and applied to various aspects of human life, including product and process development, evolution of technologies, markets, organizational development, and more [6, 7]. Later, significant progress has been made with the introduction of DE software, which incorporated powerful analytical tools and a substantial knowledge base.

Besides technological forecasting and TRIZ, other approaches and technologies were involved in the development of DE, in particular:

- methods of searching for new ideas from various modifications of the trial-and-error method to

the most recent works in the area of artificial intelligence;

- approaches and methods arising from studies in the evolution of biological and other systems (technological, social, arts, etc.) over the last two centuries;
- ideas developed within the theory of nonlinear systems and synergetics (autopoiesis); and
- ideas associated with the informational wave of modern civilization.

Directed Evolution Components

Directed evolution methodology includes the following components:

- DE forecasting – a systematic procedure for identifying all possibilities for system improvement and possible undesired events that might be associated with the system evolution and developing a comprehensive set of logically sequenced scenarios of evolution of the given system using patterns and lines of evolution.
- DE root cause analysis – a systematic procedure for identifying the root causes of any phenomenon (useful or harmful) and inventing ways to prevent harm and/or possibilities for its useful application.
- DE inventive problem solving – a systematic procedure for resolving tough technological problems, enhancing system parameters, improving quality, reducing cost, and so on, for products, processes, technologies, businesses, management, and so on.
- Evaluation and enhancement of intellectual property – a systematic procedure for unveiling and increasing IP value and strengthening IP protection against infringement and circumvention.

Directed Evolution Steps

The DE process comprises the following steps:

- analyzing the past evolution of the system of interest (whether technological, social, business, etc.) and revealing the main patterns in this evolution;
- selecting the goals for further evolution using the DE knowledge base, which incorporates information about typical needs of individuals, society,

and related environments together with basic patterns of market evolution;

- connecting the evolution of the given system with anticipated evolution of other related systems, the environment, and general technological evolution using the Ideation Bank of Evolutionary Alternatives™,b;
- defining how to pursue the selected directions and revealing evolutionary resources – i.e., the technologies, materials, products, processes, skills, knowledge, and so on, that could be utilized in the development using the DE knowledge base, which includes patterns describing the evolution of technological and social systems;
- identifying potential dangers and other undesired events, possible future problems, and so on, both in the given system and its environment, using Anticipatory Failure Determination (AFD)® and a special knowledge base of information about typical evolutionary dangers;
- revealing and solving problems, resolving contradictions, overcoming limitations, and so on, that can hinder the achievement of set goals using various instruments for problem solving including the most recent Ideation TRIZ methodology;
- developing a logically sequenced evolutionary scenarios for the given system, including a strategy for beating the competition;
- providing reliable and effective protection of intellectual property and (if necessary) eliminating possible legal obstacles, including invalidation and designing around blocking patents;
- monitoring and continuing development of the process of implementing selected evolutionary scenario(s) of the given technological or business system, utilizing feedback for prompt adjustment.

Directed Evolution Instruments

The DE process involves the utilization of a number of tools embedded in the software^c, in particular:

- DE questionnaire – a comprehensive set of questions and instructions for conducting a thorough diagnostic evaluation of the given system.
- Problem Formulator® – a software module utilizing elements of artificial intelligence for building and analyzing cause–effect diagrams describing the structure and functioning of the given system.

This tool generates lists of specific control questions for system analysis and evaluation, and then creates a set of directions for evolving the system, for revealing and solving problems, revealing the root causes of undesired events, and so on.

- A set of evolutionary patterns for artificial (man-made) systems, including over 400 patterns and lines of technological evolution and over 100 patterns and lines of marketing and social evolution.
- A comprehensive multilayer system of approximately 1000 patterns of innovation (“operators”) for:
 - solving technological problems;
 - solving nontechnological (business, marketing, logistics, etc.) problems;
 - conducting failure analysis; and
 - conducting failure prediction.

Typical Results of a Directed Evolution Project

Typical results of a DE project include:

1. A comprehensive diagnostic analysis of the DE subject, in particular:
 - (a) documenting existing problems and contradictions – and revealing hidden ones – that hinder the effective evolution of the given system,
 - (b) revealing the evolutionary potential and resources that allow for an increase in system ideality and provide new opportunities, and
 - (c) objectively evaluating the applicable intellectual property and developing suggestions for enhancing it and providing better protection.
2. Solving selected problems, generating new ideas and building alternative concepts for evolving the system, based on suggestions for selecting the most promising directions for effective strategic planning for the short, mid, and long term.
3. Predicting possible mistakes and undesired events associated with further evolving the system (for short-, mid-, and long-term planning) and developing recommendations for:
 - (a) timely prevention of potential mistakes and other undesired events;
 - (b) early diagnostics of potential undesired events;
 - (c) prompt actions capable of reducing harmful consequences from unexpected undesired events; and
 - (d) possibilities for capitalizing on undesired events such as changes in government regulations, social changes, and so on, even economical crises.
4. Providing recommendations for the effective growth of intellectual property and structuring an IP portfolio, including:
 - (a) analyzing and “sprucing up” existing ideas and inventions, including valid and expired patents; evaluating and enhancing these ideas/inventions to ensure their best utilization;
 - (b) adding new inventions to the IP portfolio; and
 - (c) ensuring effective protection of the IP portfolio, or of its most valuable elements.
5. Increasing the company’s creative potential.

To date, over 100 DE projects have been completed. The list of selected DE projects includes:

- chemical industry
 - new generation of existing chemical processing plant;
 - processing of new chemical;
 - catalysts, etc.
- heavy equipment: lifting cranes
- automotive industry
 - parking brakes;
 - doors;
 - cable applications, etc.
- medical industry: surgical instrumentation
- oil industry: control of oil production
- materials development
 - superabsorbent fibers;
 - plastics;
 - catalysts;
 - packaging, etc.
- fluid systems
 - water pumps;

4 Directed Evolution[®]

- hazardous materials pumps
- commercial products
 - sanitary products;
 - foot massagers;
 - electric shavers;
 - hair care products: combs, hair dryers, curling irons, and hair clippers, etc.
- stock/commodity exchange.

Acknowledgments

The authors would like to thank Vladimir Petrov for inspiring our special interest in patterns of Evolution, Dr Gafur Zainiev for suggesting the name of *directed evolution* and, for very useful discussions, and Zion Bar-El for encouraging us to write this article.

End Notes

^a. Directed Evolution is a registered Trademark of Ideation International Inc.

^b. The bank of evolutionary alternatives includes the results of completed and partially proven predictions

for numerous important areas such as energy, food, informational technology, and so on.

^c. Directed Evolution software is protected by US patent US 5.581.663: Automated Problem Formulator and Solver.

References

- [1] Toffler, A. (1981). *The Third Wave*, Bantam Books.
- [2] Jantsch, E. (1967). *Technological Forecasting in Perspective*, Organisation for Economic Co-operation and Development, Paris.
- [3] Martino, J.P. (1983). *Technological Forecasting for Decision Making*, 2nd Edition, North Holland.
- [4] Altshuller, G. (1984). *Creativity as an Exact Science*, Gordon & Breach, Science Publishers.
- [5] Zlotin, B. & Zusman, A. (2006). Patterns of evolution: recent findings on structure and origin, in *Altshuller's TRIZ Institute Conference TRIZCON 2006, Milwaukee, April 29–May 2, 2006*.
- [6] Zlotin, B. & Zusman, A. *TRIZ in Progress*, Transactions of the Ideation Research Group, Ideation International, 1999.
- [7] Zlotin, B. & Zusman, A. (2001). *Directed Evolution: Philosophy, Theory and Practice*, Ideation International.

BORIS ZLOTIN AND ALLA ZUSMAN

Quality Management, Overview

Introduction

Quality management (QM) is an essential feature in any organization that desires to be considered reliable by its customers – an organization that consistently keeps the promises that it makes. Reliable organizations design compelling promises for customers and consistently perform according to these promises at a high level of service thereby achieving a competitive edge over their rivals. This approach to quality requires two key elements – innovation to develop new products (whether goods or services) and customer care to effectively and efficiently deliver these products – and obtains predictable financial results for business owners. This means that the value proposition offered to consumers must deliver value that is beyond the offerings of competitors, so that consumers will consistently choose their own organization's products/services instead of the alternative offerings. QM is the approach to develop and maintain a “value edge” in competitive differentiation at a level that assures both market leadership and profitable growth. This outcome is achieved through an integrated system of business practices called *quality management* [1].

Quality Defined

To understand QM requires defining *quality* in a way that eliminates any subjective interpretation. In order to provide an operational definition of quality and illustrate how it is applied, three questions will be addressed: what is quality, how is quality created, and how is quality delivered.

First, We Must Ask What is Quality?

Quality is a comparison of expectations for performance outcomes with the perception of their achievement. Technically quality can have two meanings: it can describe either the characteristics of a product or service that give it an ability to satisfy stated or implied customer needs or it can refer to a product or service that is free of deficiencies. But, the most important aspect of quality is the one that

creates business competitiveness for an organization. To describe this perspective a theory of **attractive quality** was developed by Dr. Noriaki Kano of Tokyo Science University [2]. The Kano model of attractive quality (Figure 1) illustrates three functions that define the dimensions of competitive quality.

Kano observed that there are three functions that define how quality is perceived by customers as meeting their requirements. When customers describe their quality needs, they define their understood needs; they do not talk about unknown requirements for unanticipated needs or applications nor do they describe their assumptions regarding basic performance that is embedded within their concept of the good or service. The three curves result in different types of customer performance perceptions.

- The lowest curve represents “must be” or standard quality which is characteristic of a commodity product (for instance, there is no competitive value in the standard grade of octane in gasoline – people are likely to change brands for very marginal differences in price). These requirements are not typically described by customers as they are “expected” or “understood” to be part of the package – cars are expected to go and to stop, and most people would feel that these capabilities must not be specified, thus these features deliver no satisfaction to customers but only offer an opportunity to dissatisfy the customer through poor performance.
- The middle curve describes what Kano called *one-dimensional quality* or a product quality feature where companies compete “head to head”

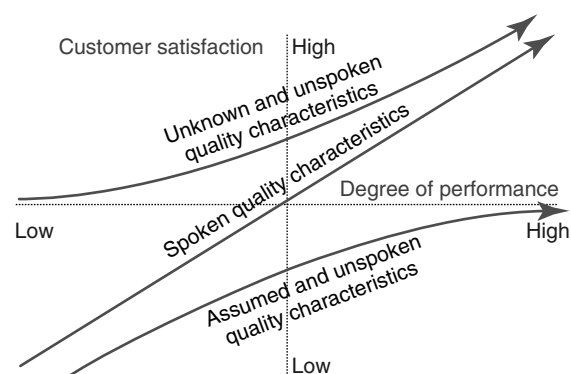


Figure 1 The Kano theory of attractive quality

2 Quality Management, Overview

to win customer approval (e.g., gasoline consumption of an automobile) and value as perceived by customers is directly proportional to the performance of this feature design. In this category of product feature, quality is directly proportional to the effectiveness of the design to the customer's experience in accomplishing the job they need to accomplish [3].

- The third curve describes the type of quality that always satisfies customers or what Kano refers to as *attractive quality* – quality that anticipates customer needs before the customer knows or understands what capability or feature they are missing. This type of quality comes from innovation – knowing the customer environment and application so intimately that new technology can be applied to create new competence or capabilities in the job that the customers need to get done. This type of quality provides a disrupting influence on markets and can totally change the balance of its competitive structure.

Thus, according to Kano, quality is directly linked to the ability of a product or service to deliver perceivable value to its targeted customers.

Second, We Must Ask How is Quality Created?

Given this understanding of the importance of the customer's experience, how is this type of quality created? Consider the Watson model for value delivery for a definition of the way quality is created and delivered in the customer experience (Figure 2) [4].

Quality creation is a two-phased process: designing value into products and services and delivering

the quality of value that was designed. The design process interprets the customer experience of a Kano analysis to make a commercial promise to the market. This first stage in quality creation establishes the expectations of customers for the value they will receive and defines the offer made to customers. Not understanding or misinterpreting the best value proposition for customers thus results in suboptimization of the organization's potential competitive position. The second stage of quality delivery is in the delivery of the promise that has been defined for customers – a customer care function in which the consistency of promise delivery is evaluated by customers. Managing the daily dynamics of the customer experience must satisfy a quality entitlement – customers are entitled to consistently perceive satisfaction in the delivery of the promised performance. By aggressively executing the delivery of this quality model through an organization's business model, persistent perceptions creating long-term customer loyalty are achieved.

Third, We Must Ask How is Quality Delivered?

Sustained success and profitable growth are achieved when an organization coordinates its activities to deliver quality above its competitors, costs below its competitors, and technology ahead of its competitors. These imperatives must be constructed into a system for quality creation, which is delivered through the business processes of the organization (Figure 3 describes such a basic business model). This model characterizes the three core processes of an organization: product creation, product realization, and the process of management.

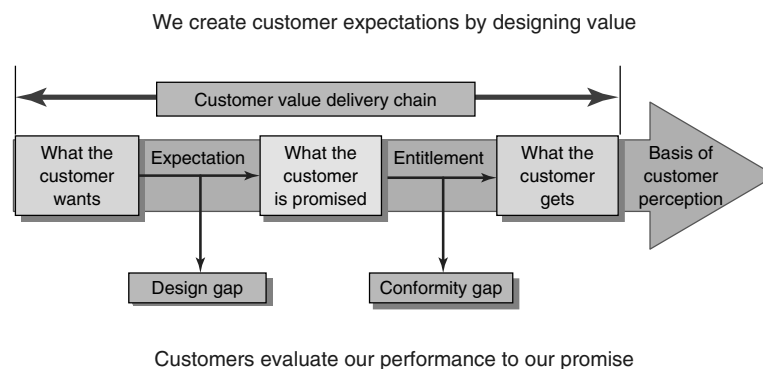


Figure 2 The Watson model for value delivery

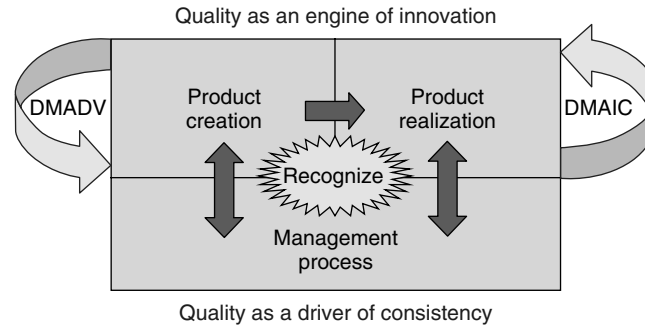


Figure 3 Quality-driven business model

Figure 3 overlays the linkage of **Six Sigma** methodology and policy deployment on top of the business model to illustrate the linkage between change management projects (business improvement seeking *kaizen* (continuous improvement) or *hoshin kanri* (breakthrough improvement)) to the daily management system (*nichijo kanri*) represented by these core business processes. Note that two types of innovation are required to implement this change management process and that these are expressed by the two different approaches that are taken to resolve issues using Six Sigma business improvement methods:

- **DMADV** (define–measure–analyze–design–verify) is the Design for Six Sigma (DFSS) process that delivers innovation through business development processes by focusing on value delivery through product creation, service creation, or development of new value delivery work processes. This Six Sigma approach resolves issues where a current performance capability is inadequate to meet its future requirements (the designed capability is unable to meet the customer’s need).
- **DMAIC** (define–measure–analyze–improve–control) is the **lean** Six Sigma process that delivers innovation in the product realization process by improving or optimizing the performance of current products, services, and the processes that deliver products or services by applying the methods of statistical problem solving and lean production management. This Six Sigma approach resolves issues where the current level of performance is not achieving its designed capability (achieved process capability is less than the designed process capability).

Quality must be managed over the long term but it is delivered in the short-term actions of the daily management system through a disciplined work process. What is the set of specific components that are contained in a quality management system (QMS) and how do they operate?

The Elements of Quality Management

QM differs from the management of quality. QM delivers excellence through a set of good management practices while the management of quality describes the profession of managing the quality function within an organization – the first is applicable to managers in all organizational functions and at all organizational levels, while the second applies more narrowly to a specific profession and its associated disciplines.

QM is an “unnatural act” of establishing and preserving excellence – QM defies the natural law of entropy whereby everything degenerates into a state of mediocrity before it obsolesces.

Management of quality combines several specific disciplines to assure the consistent delivery of the level of service or product quality intended at the point of the customer experience. In managing quality there are five quality disciplines that have evolved over the past century. These five disciplines for managing for quality relate to practices of engineering, assurance, control, and improvement. Taken together they follow a PDCA cycle (plan–do–check–act) that is often described as the *Deming cycle*.^a

- **Quality engineering (QE)**
QE accomplishes Deming’s “plan” step for designing an enduring level of product or service

quality. QE analyzes the steps of an operational process and gives a statistical basis for measuring and managing process performance so it consistently meets outcomes for customer standards or requirements. QE delivers predictable quality outcomes by maximizing process quality to produce the required quality results in the product or service.

- **Quality control (QC)**
QC describes the methods used to address the Deming “do” step for process management. Control actions include processes for actively testing or monitoring performance, evaluating current performance results against standards for desired performance, and guiding process activities to achieve a state of statistical control by taking corrective action when performance is observed to deviate from this standard. Thus, QC consists of all operational techniques and activities used to fulfill the requirements for quality outcomes by focusing on improving process performance.
- **Quality assurance (QA)**
QA delivers the “check” step of the quality discipline. QA gives confidence that the customer requirements will be achieved by demonstrating the product or service capability to fulfill performance results as delivered to the final customer. QA is distinguished from QC in that QA focuses on the final test results or compliance with the external quality requirement while QC focuses on the interim test results or compliance with internal quality requirements, which produce final results.
- **Quality improvement (QI)**
QI executes the “act” step of the PDCA cycle. The QI process is one of continuously improving or adjusting performance to reduce the cost of poor quality, improve the cycle time of process activities, and eliminate waste from all process steps. QI seeks to obtain and sustain the “ideal” process performance that was initially designed into the process capability. QI activities include the work of both Six Sigma and *kaizen* blitz teams who seek to eliminate waste or streamline process activities.
- **Quality audit and review (QAR)**
QAR is a systematic, independent assessment of the entire quality system to determine the effectiveness of both the quality plan and its execution through application of the four disciplines of quality in the work processes of the organization.

Thus, QAR “checks” on the value of governance in the PDCA activities of an organization’s formal QMS.

When these five quality disciplines are combined into a systematic approach for the delivery of quality results, they become the cornerstones of a sound QMS. Thus, the “professional” dimension of QM may be defined as follows:

- **Quality management (QM)**
The process of designing and deploying a system for managing processes to achieve maximum customer satisfaction with product or service quality at the lowest overall cost to the organization while continuing to improve the performance outcomes of the process.

This professional dimension of QM seeks to instill a systematic way of working in the entire organization for delivering quality outcomes – achieving organization-wide practices that deliver managing for quality. What are the ingredients of this system?

Management System for Delivering Quality

The formal QMS of an organization is a documented process for managing these five dimensions of QM (described above) and QMS describes the context for the design, development, deployment, and delivery of quality activities throughout the organization. The QMS documents and executes the organization’s quality policy (strategic direction about how the organization will work to deliver quality) and it can be summarized as the PDCA of the management system: quality planning, quality standards (defines the system of work activities, controls and tests, as well as the control plan to assure that the system operates effectively despite any potential contingencies), QA, QC, and QAR, which must occur supported by the activities of an iterative process for QI. The QMS is a formal system to document the organization structure, management responsibilities, and organization-wide procedures that are required to achieve effective QM as a routine way of working. The ingredients of the QMS documentation include the following:

- *Quality policy*
An organization's general statement of its beliefs about quality, its vision of how quality will come about and what should be the expected results.
- *Quality plan*
A document or set of documents that describe the standards, quality practices, resources, and processes pertinent to the assurance of quality for a specific product, service, or project through the application of the four quality disciplines.
- *Quality standards*
The documented description of accepted way to perform work and quality operating practices so that "best practice" is applied as a discipline in the daily working environment. This documentation includes standard operating procedures as well as control plans to assure proper measurement and interpretation of the critical quality process measures and the judgments made to interpret them regarding quality performance. Standards should also exist to describe how each of the core disciplines of QM will be managed.
- *Quality systems*
The information systems that manage data for product and service performance, customer and market information about complaints, failures, warranties, and the record of responses to such claims.

The typical system for delivering quality in a modern company contains a number of proven ingredients for success: ISO 9000, business excellence criteria, and Six Sigma which can be forged into the process management architecture that is shown in Figure 4.

This integrated QMS seeks to maintain a business control system based on the ISO 9000 Quality Management Standard while encouraging the organization to aspire to apply the "best practice" principles of QM, which have been distilled into the award criteria for business excellence prizes such as the Malcolm Baldrige National Quality Award or the European Quality Award. These award criteria address a system of areas for concern in the development of a business management system that were developed as a consensus from many of the world's best companies. No company has ever scored above 800 points using these criteria as a standard, so they provide a way to challenge organizations to initiate self-improvement projects that close the gap in opportunities for improvement that are discovered when the organization conducts a self-assessment or conducts formal benchmarking with an organization that is recognized as superior by their management team. These improvement projects can be conducted using the Six Sigma DMADV or DMAIC methods.

Achievement of exceptional performance using such a management system requires that the organization develop "leadership at all levels" of its structure. What does this mean in terms of the quality actions of the organization's business leaders? We have finished discussing the key dimensions of management of quality; now let us turn back to understanding how QM can be infused into an entire organization as it integrates its QMS into its core business to manage performance in all dimensions of business excellence while ensuring successful results from the perspectives of all the organization's stakeholders (e.g., customers, employees,

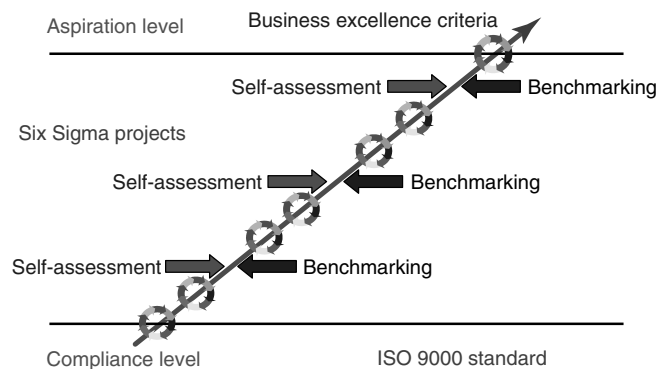


Figure 4 An integrated quality system for total performance management

suppliers, investors, legal and regulatory authorities, etc.).

Pervasive Quality Leadership

It is a tall order to deliver this kind of management system. It requires, as Warren Bennis observed, the kind of situation where: “a leader is a follower is a leader” [5]. Bennis noted that we are facing a “chronic crisis of governance – that is, a pervasive organizational incapacity to cope with the expectations of their constituents – [that] is now an overwhelming factor worldwide. If there was ever a moment in history when a comprehensive strategic view of leadership was needed, not just by a few leaders in high office but by large numbers of leaders in every job, from the factory floor to the executive suite, this is certainly it” [6]. In order to have a system of sound business management there must be good role models of management quality – not just at the top of the organization, but throughout its entire structure so that the people can see this behavior is not just applicable for the “super-leaders” but that is also part of the expectation for every ordinary worker – including supervisors – and they are also capable of exhibiting the shared values and behaviors that operationally define quality actions throughout the entire organization. This consistency assures that each person understands that the vision of leadership is not about wishing, hoping, or praying: it is an act of courageous leadership that must be duplicated at all levels of the organization. What does leadership at all levels look like?

Executive Leadership

Senior managers must establish the vision and values that will guide their organization into its future state. One way executives demonstrate leadership is by establishing a framework for action through a management process for cascading values and objectives that define the way the organization will work together and the goals that it will seek to accomplish in pursuit of its long-term mission. A critical aspect of the process of management that remains the active responsibility of the top management team is conducting regular leadership reviews of the business. This activity is both a due diligence responsibility

of the business owners and a fundamental approach for exercising leadership. Leadership reviews have at least two emphases: review of the governance structure and “strategic direction of the organization” as well as review of strategic problem areas that lead to business vulnerability through challenging technologies, violation of critical assumptions, or changing “rules of competition” in the market. Leadership reviews offer the top management an opportunity for

- demonstrating personal commitment to their “philosophy of management”;
- providing visibility for their “defining moments of quality encouragement”;
- mentoring the organization to achieve desired behavioral changes;
- guiding cross-organizational efforts to achieve desired systemic change; and
- encouraging people to “build a desire to win and a will to act”.

What else can executives do personally to demonstrate “positive leadership behavior”? A few examples include the following:

- *Executive touch program*
Program for top executives to get close to the leading targeted external customers. Requires quarterly visits that are intended solely for relationship building to establish a foundation for future business discussions.
- *Executive customer advocate program*
Program for the senior management team to facilitate significant problems encountered by major customers in each of the primary lines of business.
- *Complaint listening program*
Program for all levels of management to spend a few hours monthly listening to “real” customer complaints directly on call center lines.
- *Executive escalation program*
Program for the top level of senior managers to rotate through a monthly “duty day” where they are represent formal escalation “points of last resort” for resolving all customer complaints.
- *Executive compensation program*
Change the compensation so a significant element of the “reward component” is granted for a “statistically valid” increase in customer satisfaction as measured by a valid external method.

Business leaders must exhibit consistency in all that they do. Lack of consistency is considered by employees to countermand the organizational culture and may cause deterioration in a shared system of values. It is particularly important to demonstrate consistent performance when one of the tenants of the organization is “empowerment” – the ability to make individual choices within a set of boundary conditions. An organization’s vision provides direction for empowerment, but its values provide the boundaries for making choices.

Management Leadership

Managers may also become leaders. The definitions of manager and leader are not mutually exclusive. Each person appointed to a management function should at least aspire to becoming a local leader. What can be done to demonstrate leadership at the local level? Several actions can be suggested:

- taking an active role in leading the local implementation of a major business or QI initiative that has strategic value to the organization as a whole;
- showing a personal interest in developing the next generation of leaders through the personal mentoring of high-potential employees;
- “managing by wandering around” and taking time to talk with employees about any issues or concerns they may have, providing brief words of encouragement, and taking action to apply their insights for better management of work or business processes;
- sponsoring and reviewing improvement projects that deliver on the annual continuous improvement objectives of their own management area of responsibility;
- recognizing the improvement efforts of frontline teams and individuals;
- developing the core competence of their organization through team-based on-the-job education and training programs; and
- exercising information transparency by communicating to all employees about business results and briefing them on both the strategic direction and any news that directly concerns them or their livelihood.

Frontline Leadership

Leadership is often required at the frontline in order to coordinate action and encourage common behaviors that reinforce organizational values. Each person can exhibit leadership as a means to encourage fellow employees and provide an example that reinforces the behaviors of the value system. Some of the actions that can be taken at the personal level include

- participate actively in their work group or team process management activities and continuous improvement projects;
- develop their personal problem-solving and statistical analysis skills;
- mentor new employees in the cultural values and historical accomplishments of the organization;
- pursue certification in the core skills involved in the profession and demonstrate mastery of the tools and methods of their own trade;
- provide improvement ideas and suggestions to managers whenever any opportunities for improvement are observed in the work process;
- participate on teams for conducting self-assessments, audits, and cross-functional process improvement;
- take responsibility for their personal development and pursue a combination of both internal and external courses that deliver their career objectives; and
- recognize the contributions of colleagues and team members and encourage their achievements by expressing appreciation for their positive involvement in team improvement projects or personal suggestions made for process improvement.

Personal Leadership

True leadership is not in words, but in deeds: “The authentic test for mastery of learning is not in what a manager [person] says, but in what a manager [person] does” [7]. This requires a consistent practice that is developed from the inside out and exists on both personal and interpersonal levels both within the local work experience and across the whole organization. This effort requires two key focus areas: managers must empower the workforce to exhibit the principles of total quality leadership and frontline workers must become aligned with the strategic direction, improvement objectives, and

cultural values of the leadership team. Each employee should aspire to represent a role model of the behavior desired of colleagues in the organization. This principal is also a cornerstone of the lean methods of the Toyota production system where employees are held accountable for the quality of their work and are expected to manage the quality of their work using a self-regulated QC approach. This raises an important question about the responsibility for quality and the role of quality professionals in managing to achieve quality outcomes in contrast to the responsibility of both process owners and workers for the quality of their own work output. What is the role of a quality professional in managing for quality?

Professional Leadership: The Role of a Quality Manager

Responsibility for Quality Performance

Who is responsible for the quality outcome of work? Is it the individual doing the work, the supervisor of the work effort, the owner of the work or business process owner, or the leader of the business area? There is a quandary: if everyone is responsible for the quality of work, then nobody is accountable! The resolution of this quandary lies in the multiple levels of responsibility – which fosters different types of accountability! Responsibility accumulates from bottom up; however, accountability is delegated from higher levels of authority to lower levels in the organizational structure. Thus management delegates responsibility for quality performance of work to the workers, but maintains overall responsibility and accountability. Workers can only be held responsible and accountable in the area of work that has been properly delegated to them. You cannot have accountability without responsibility! What does it take to hold someone properly accountable for executing their responsibility and then delivering the desired level of quality in their work? According to the sage advice of Peter Drucker, there are three conditions that must be satisfied to hold a worker personally accountable for the quality of their work:

1. Employees must have specific knowledge of the job they are asked to perform in terms of specifications, measurements, systems for work performance, and contingency actions that should be performed when work quality is not acceptable according to standards and they must be trained to perform to these work requirements.
2. Employees must know the standard of quality they are expected to deliver (the targets or goal of their performance expectation).
3. Employees must have the ability to monitor their work progress and be delegated the authority to make decisions that allow them to self-regulate their own performance to consistently meet the standards of performance [8, 9].

When these three conditions are not met, then the business managers retain the responsibility and accountability for all aspects of work quality. The job of a professional quality manager is to develop the QMS that assures that responsibility for execution of work practices leading to consistently sustainable quality results is achieved. What behavioral characteristics does it take to be such a “local leader” of the organization’s quality practices?

Behavioral Qualities of an Exceptional Quality Manager

On the basis of a behavioral analysis of successful and unsuccessful quality managers, there are some common traits that have been identified as distinguishing between the most successful quality managers and those who were not perceived by their bosses as being successful. Sixteen traits were identified as leading to success [10]. Interestingly, only one trait was observed to detract from managerial performance. The positive behavioral traits that were discovered in this study included customer oriented, customer advocate, organizationally astute, influencing, goal oriented, interpersonally diagnostic, persistent, organizational planning, initiating, mentoring subordinates, collaborative, professional, conceptual, innovative, communicative, and self-confident. The one negative trait that leads to loss of personal influence and ineffective personal interactions was “making fast decisions.” On the basis of these characteristics, it is clear that a “local quality leader” must behave like an internal consultant acting without the direct management authority and responsibility for line management (e.g., management of the profit and loss centers of the organization) but must convince others, through persuasive use of data, as to what is the right course to take. In an interview with Peter Drucker he was asked: how can a quality professional

convince their business leader to “do” quality? His answer was most educational. Drucker said, “It is not your job to train your CEOs. They are bright people and can understand quickly what needs to be accomplished. However, it is the job of staff to clearly report necessary information that CEOs can easily assimilate and understand, so they may draw their own conclusions” [11]. The “secret weapon” of quality professionals is the use of data as a means to convince management to act. Communicating concerns with clarity and confidence is perhaps the greatest competence that a quality professional can have.

Summary

What is quality? Quality is the never-ending pursuit of excellence in product or service performance as judged according to the standards of customers relative to alternative choices they have for selecting similar goods or services from competitors. These customer standards for quality change over time and are influenced by a wide variety of external forces such as technology and regulatory legislation. As quality improves in one area in life, it often has a positive influence by increasing the requirement for quality in other areas. Thus, managing to improve quality is an essential ingredient to continued success and QM is a business requirement that will continue to be part of the “essential work” of organizations.

End Note

^a Although PDCA is often called the *Deming cycle* it was first described by Walter Shewhart and the PDCA version was actually developed by the Union of

Japanese Scientists and Engineers (JUSE) translators of Dr Deming’s lectures under the supervision of Kaoru Ishikawa.

References

- [1] Watson, G.H. (2003). Persistent leadership: a key to sustainable quality, in *Quality into the 21st Century: Perspectives on Quality and Competitiveness for Sustained Performance*, T. Conti, Y. Kondo & G.H. Watson, eds, ASQ Quality Press, Milwaukee, Chapter 5.
- [2] Kano, N., Seraku, N., Takahashi, F. & Tsuji, S. (1984). Attractive quality and must-be quality, translated by Glenn Mazur, *Quality, The Journal of the Japanese Society for Quality Control* **14**(2), 147–156.
- [3] Christensen, C.M. & Raynor, M.E. (2003). *The Innovator’s Solution*, Harvard Business School Press, Boston.
- [4] Watson, G.H. (2003). Customers, competitors and consistent quality, in *Quality into the 21st Century*, T. Conti, Y. Kondo & G.H. Watson, eds, ASQ Quality Press, 2003, *Ibid.*, Chapter 2.
- [5] Bennis, W. (1989, 1994). *On Becoming a Leader*, Perseus Books, Cambridge, p. 39.
- [6] Bennis, W. & Nanus, B. (1997). *Leaders: Strategies for Taking Charge*, 2nd Edition, Harper Business, New York, p. 2.
- [7] Watson, G.H. (1994). *Business Systems Engineering*, John Wiley & Sons, New York, p. 116.
- [8] Drucker, P.F. (1954). *The Practice of Management*, Harper Collins, New York.
- [9] Drucker, P.F. (1973). *Management: Tasks, Responsibilities and Practices*, Harper Trade, New York.
- [10] Watson, G.H. (1999). Building quality competence: successful management behavior, in *Proceedings of the 53rd Annual Quality Congress*, May 25, 1999.
- [11] Watson, G.H. (2002). Selling Six Sigma to CEO’s, *Six Sigma Forum Magazine* **2**, 26.

GREGORY H. WATSON

Benchmarking in Project Definition

Introduction

In 1990, Roger Milliken called benchmarking *the art of stealing shamelessly* and since then many have used benchmarking to generate “quick fixes” for their performance problems. However, benchmarking is not just for quick fixes, it can also provide a rigorous process that requires both “sweat equity” – learning about one’s own processes and performing the logistics of coordinating study missions to other organizations – and “analytical thoroughness” – measurement and analysis of sustained work process performance as well as the detailed mapping of processes and side-by-side evaluation of process differences. Benchmarking uses the analytical information contained in a benchmark, or a comparative measure of process or results performance to establish which organization is the candidate for a “best practice” in a specific business process. Then the business process must be specified in order to understand how the benchmark was achieved and to identify enablers of successful overall performance. Finally, a cultural adaptation of the learning must be made in order to apply this new knowledge to another organization. For benchmarking to be successful, it must heed the warning of Dr. W. Edwards Deming who said, “It is hazard to copy. One must understand the theory of what one wishes to do.” Cultural adaptation and business model adaptation are necessary to assure that the lesson observed in one organization can be successfully transferred to another organization. As Deming also cautioned, “Adapt, don’t adopt. It is error to copy.”

Benchmarking Definitions

What is benchmarking? In order to understand this methodology, it is essential that definitions be agreed upon for the meaning of some key terms. The first definitions that must be agreed upon are terms that identify the types of benchmarking studies and their operational approaches:

Process Benchmarking

The method for comparative analysis of work process performance between two unique or distinct implementations of the same fundamental process is called *process benchmarking*. It includes both the internal inspection of an organization’s own performance as well as the external study of an organization’s recognition for achieving superior performance as evidenced by objective standards of comparison (benchmarks). The objective of process benchmarking is not to calculate the quantitative gaps in performance, but to identify best practices that may be adapted for improvement of organizational performance. There are four types of process benchmarking studies: strategic, operational, performance, and perceptual benchmarking.

Strategic Benchmarking

The process benchmarking of organizational strategy or key business process performance in order to determine breakthrough opportunities for profitability and productivity improvement is called *strategic benchmarking*. This type of benchmarking focuses on those critical business areas that must change to attain or maintain the competitive advantage of a business. Strategic benchmarking studies focus on critical business assumptions, primary competence areas, core business processes, technology inflection points, or business fundamentals that define organizational purpose. The purpose of strategic benchmarking studies is to challenge the management to move from a current state to a desired state of the whole business. Examples of strategic benchmarking studies include evaluation of options for the design of an organization’s governance structure, assessment of approaches used to implement advanced technology (e.g., enterprise management software or paperless document handling), or strategic business issues that are faced by the organization (e.g., creating a web-based business capability; managing the technology transition across generations of advancement; or managing the routine work of the organization through management methods such as balanced scorecard, performance management, and business excellence assessments).

Operational Benchmarking

The process benchmarking of work processes or practices in order to discover opportunities that will

2 Benchmarking in Project Definition

provide productivity improvement in the areas of effectiveness, efficiency, or economy of the routine business operations is called *operational benchmarking*. This type of benchmarking focuses on specific work activities that need to be improved and seeks to identify the work procedures, production equipment, skills or competence training, or analytical methods that result in sustained performance improvement as indicated by objective measures of process productivity (process throughput, cost per unit, defect opportunities, cycle time, etc.). Examples of operational benchmarking studies include analysis of invoicing procedures to determine the most productive process, evaluation of production methods to determine the highest throughput methods that deliver lowest cost and least defects, and study of logistics distribution methods that result in both high-delivery service performance and low levels of finished goods inventory.

Performance Benchmarking

The process benchmarking of product or service results using a standard comparison or test under known operating conditions is called *performance benchmarking*. This type of benchmarking seeks to answer the question which product or service is better based upon rigorous assessment using objective performance criteria. Examples of performance benchmarking studies include consumer product analysis that evaluates products on a “head-to-head” basis using a fixed set of criteria for performance, evaluation of product performance using a standard test, such as operating time to run a specific application, or endurance tests that identify the ability of a product to perform over a fixed period of time under comparable operating conditions.

Perceptual Benchmarking

The process benchmarking feelings or attitudes about process, product, or service performance by the recipient of the process output is called *perceptual benchmarking*. This type of benchmarking seeks to answer the question how do you perceive the delivery of service, performance of product, or execution of process by the people who are recipients of these outputs. Perceptual benchmarking uses attribute or categorical data to quantify subjective feelings and establish relative ranking of performance based on such criteria as timeliness of performance, goodness

of knowledge transfer, soundness of information, courtesy of delivery agents, and so on. Examples of perceptual benchmarking include surveys of training satisfaction at the completion of a training event, employee satisfaction surveys to determine either the work climate or structural issues regarding compensation and benefits, or customer satisfaction with the product or service delivery to the marketplace.

The next definitions that must be agreed on are the terminologies that identify the sources of data used in conducting a specific benchmarking study. This is an older and somewhat less helpful way to identify benchmarking studies. This categorization of benchmarking practice according to the source of data leads to conclusions about the relative utility of information from these sources.

Competitive Benchmarking

An approach to benchmarking that targets specific product designs, process capabilities, or administrative methods used by one’s direct competitors (e.g., the study of performance in the laptop computer industry that features only those companies that produce these products).

Industry Benchmarking

An approach to benchmarking that seeks information from the same functional area in a particular application or industry (e.g., benchmarking the purchasing function to determine the most successful approach for managing a supplier base).

Internal Benchmarking

An approach to benchmarking where organizations learn from “sister” companies, divisions, or operating units that are part of the same operating group or company (e.g., the study of internal research and development groups to determine best practices that reduce the time to market for the new product introduction process).

Generic Benchmarking

An approach to benchmarking that seeks process performance information that is from outside one’s own industry. Enablers are translated from one organization to another through the interpretation of their

Table 1 Benefits analysis of benchmarking data sources

Source of data	Advantages	Disadvantages
Competitive benchmarking	Provides a strategic insight into marketplace competitiveness and a “wake-up” call to action	Legal issues regarding data sharing among competitors
Industry benchmarking	Takes advantage of functional and professional networks to gain study participants	Study detail may not be good enough for process diagnosis Functional concentration tends to support operational rather than strategic studies
Internal benchmarking	Provides highest degree of process detail and simplified access to process information	Does not challenge paradigm of functional thinking The internal focus tends to be operational, rather than strategic, and reinforce the organization’s cultural norms
Generic benchmarking	Has the greatest opportunity for process breakthroughs Because organizations do not compete, reliable detailed information is usually available Provides incentive for strategic change initiatives	Difficulty in developing an analogy between dissimilar businesses Difficulty in identifying the companies to benchmark Difficulty in establishing the appropriate contact for a study

analogous relationship (e.g., learning about reducing cycle time in production operations by the study of inventory management methods used in stocking fresh vegetable in grocery stores).

Comparative advantages and disadvantages of these benchmarking data sources are presented in Table 1.

The final definitions that must be agreed on are the terminologies that are used to describe the different aspects of a benchmarking study:

Benchmark

A benchmark compares performance between products, services, or processes of analogous organizations in order to establish which one is superior in its sustained performance. Note that many of the benchmarks that are publicly promoted indicate only “spot” performance at a specific point in time and do not meet the criteria of “enduring success” by failing to establish the difference in performance between a “special cause event” and a “common cause” management process. A lack of statistical discipline in the use of benchmarks threatens to diminish the perceived value of the process of benchmarking (see the section titled “Presenting Benchmarking Results”).

Best Practice

Best practice defines the set of activities, tasks, resources, training, and methods that created the observed benchmark level of performance in an observed work process. In a process benchmarking study, in order to qualify as a “best practice” the performance must be observed and mapped to assure that the work performed is properly identified and that process experts have validated and verified the distinctions between observed best practices and merely good practice. Without the objective assessment by work process experts, “best practice” becomes a subjective claim that is not verifiable.

Critical Success Factors

These are quantifiable, measurable, and auditable indicators of process performance and process capability in key business processes. They indicate in basic business terms the performance level obtained in a comparative manner using such basic building blocks of processes to describe the performance of business effectiveness (quality), efficiency (cycle time), and economy (cost). Key critical success factors are universal and may be used for cross-organizational comparisons for the same process.

4 Benchmarking in Project Definition

Enabler

The specific activity, action, method, or technique that stimulated progress in one process over the comparative processes and led to identification of a best practice (e.g., the way **quality function deployment** or **failure mode** and effects analysis was used in a product design process; a process for data presentation that more clearly indicated the action to be taken by frontline operators; or an employee training and development system that delivers the appropriate skills and competence to process workers as they require these methods to perform their work in a changing technological environment).

Entitlement

The set of work process lessons learned that are derived by examination of one's own processes and discovery of wasted activities, duplicated steps or nonvalue-added work that can be eliminated or modified solely on the basis of the self-analysis phase of benchmarking. An organization is "entitled" to make such process changes without relying on the lessons learned from external discovery. Such improvements permit the process to operate as intended and represent gap closure between original process design and current process performance. *Entitlement* also refers to the gap that may exist between the design process capability (C_p ratio) and the achieved process capability (C_{pk} ratio) as management is entitled to the performance that they purchased with their capital investment.

Gap Analysis

Gap analysis evaluates the performance difference between current internal performance and benchmark performance at the best practice organization. To be effective, a gap analysis should include both the use of statistical **confidence intervals** and the tests of difference (for both means and variance) to demonstrate that a real performance gap has been observed and not a gap due to chance observations (see **Hypothesis Testing**).

Radar Diagram

A radar diagram provides a multivariate display of comparative performance for several dimensions

of interest (e.g., cost, cycle time, quality, and productivity). These dimensions are displayed on a single chart as spokes from the center where each measure has its own unique scale; however, all indicators are shown on the same graphic to illustrate a performance profile for a specific process. The radar diagram provides a more complete benchmarking assessment than a single-point measure of performance comparison.

Baseline Analysis

At the beginning of a study a baseline analysis is conducted to compare performance data across all benchmarked processes. A common scale is used for each comparison based on the variation observed in process performance. A best process is one that has the highest average sustained performance and the lowest variation in the daily results. The performance baseline comprehends both of these factors using a standardized metric for process comparisons (e.g., process standard deviation as calculated using the defects per million opportunities as evaluated against a common customer requirement for targeted performance). The baseline analysis may be presented as an analysis of variance showing sampled performance as a function of the different process locations.

World Class

Although it is intuitively clear that there is no one world best performance that exists at a particular point in time (the enormity of analysis to support such a claim would be unmanageable), it is possible to define a category of performance as "world class" by the fact that using a standardized measurement process (e.g., the performance baseline analysis), the process was observed in the top 5% of all performance noted in the study. This indicates that there is a high confidence level that the process is in a leadership position and worthy of investigation for potential best practice areas.

Benchmarking as a Discipline of Total Quality Management

In the development of total **quality management** (TQM), benchmarking has a unique place as both a tool to stimulate improvement and a management

Table 2 Defining the scope of benchmarking

Benchmarking is	Benchmarking is not
A discovery process	A cookbook process
An improvement methodology	A panacea for problem solutions
A source of breakthrough ideas	About “business as usual”
A learning opportunity	A management fad
An objective analysis of work	A subjective “gut feeling” or opinion
A process-based learning approach	Just measurement of process performance
A means to generate ideas	Just quantitative comparisons

technique that aids in strategic positioning of an organization. Benchmarking provides opportunities for full organizational participation in business process improvement by engaging the management team in the architecture of change and choice of focus areas for study; involving the middle managers in self-assessment of the work processes that they own and in adapting the lessons learned from other organizations; and relying on the study of related processes by the organization’s frontline process experts who are charged with discovery of the significant differences that lead to performance gaps.

What is benchmarking and how is it used? Benchmarking is a structured approach for learning about process operations from other organizations and applying the knowledge gained in your own organization. It consists of dedicated work in measuring, comparing, and analyzing work processes among different organizations in order to identify causes for superior performance. Benchmarking is not complete with just the analysis, however, it must be adapted and implemented in order to have a complete cycle of learning.

The objective of benchmarking is to accelerate the process of strategic change that leads to breakthrough or continuous improvements in products, services, or processes, resulting in enhanced customer satisfaction, lower operating costs, and improved competitive advantage by adapting best practices and business process improvements of those organizations that are recognized for superior performance. Benchmarking is a method that forces organizations to look outside themselves in order to avoid myopic illusions of grandeur that come from reflecting on internal experience without external validation.

Benchmarking is not just a checklist or set of numbers that are used to make management feel better about their current performance. Benchmarking should make management uncomfortable due to

the identification of gaps in business performance. Benchmarking should challenge management due to the discovery of performance enablers that could help them to improve. Perhaps the juxtapositions shown in Table 2 can help describe this situation.

Benchmarking Logic

It is important to observe that the logic of the benchmarking process does not fail the test that was issued by Dr. Deming in the early 1980s, when he cautioned executives against deadly diseases in the management of business that were derived from setting arbitrary goals based solely on visible performance measures, without understanding the depth of profound (process-related) knowledge that lay underneath most high-level performance measures.

Dr. Deming would call “arbitrary” the use of benchmarking with the following logic:

- “Our competitor’s price is 15% lower than ours; therefore, we must lower our cost by 15%.”

The logic of benchmarking is much more process oriented and requires the development of the type of profound knowledge that Deming advocated:

- “Leading companies have operations that are 20% more effective than our operations.”
- “The reasons that their operations are more effective is because . . . ”
- “The practices that they used to improve these operations include . . . ”
- “The following enhancements to our processes are appropriate for our business model and our culture and will improve our performance: . . . ”
- “The estimate of performance gain due to implementation of these enhancements is . . . ”

The ability to apply this logic comes from developing an understanding of the root cause of process improvement at the benchmark organization and translation of their lessons learned into appropriate change for your own organization. By a process of conscientious learning and cautious adaptation, a company can learn the lessons needed to move it to the level of world-class performance.

Operational Definition of World-Class Performance

Benchmarking seeks to deliver performance that is “best of the best” (the Japanese word for this is *dantotsu*) or world-class performance. World class is an elusive performance level. To be the “best of the best” it is necessary that you have both a high level of performance (top 5%), and that this level of performance be sustainable across changes in product life cycle, underlying technology in both the product and the process, as well as successive generations of executive leadership. That is a tall order for any organization.

What does it mean to be world class? One definition [1] considers a world-class company as one that is able to achieve and sustain a leadership performance level, while at the same time exhibiting the following competitive considerations that are significant learning areas for a TQM-oriented organization. According to this definition, a *world-class organization*

- knows its processes better than its competitors know their processes,
- knows its industrial competitors business better than their competitors know them,
- knows its customers better than their competitors know their own customers,
- responds more rapidly to change in customer behavior or needs than its competitors,
- engages employees more effectively than its competitors, and
- competes for market share on a customer-by-customer basis.

Clearly this definition of world class requires both an objective performance standard as well as the profound knowledge of its business and commercial

environment in order to enjoy sustained performance at this level.

The Process of Benchmarking

Not surprisingly, there is a process for process benchmarking. Many different organizations have illustrated benchmarking processes from 4 to 42 steps [1]; however, the process that I favor has seven steps which highlight the work that must be done in a benchmarking study and which also follow a four-phase process that is generic to all benchmarking models:

1. Identify subject – choose what to benchmark
2. Plan study – identify your partners and plan your **data collection**
3. Collect information – actively collect the data and visit partners
4. Analyze data – analyze the data for performance trends and consistency over time
5. Compare performance – compare results and test differences for statistical significance
6. Adapt applications – prepare the lessons learned for transition to your own culture
7. Improve performance – implement projects to improve your processes

The generic four-phases that these steps cover roughly follow Deming’s plan–do–check–act (PDCA) process management and improvement cycle. These four process benchmarking phases are summarized as follows:

- Plan: design the study and evaluate the baseline performance.
- Collect: collect internal and external data about the process.
- Analyze: conduct a gap analysis and determine enablers of success.
- Improve: adapt the recommendations and implement the process improvements.

Each phase of this benchmarking process is described in the following sections.

Benchmarking: Plan

The steps involved in planning a benchmarking study include

- Select a process.
- Gain the owner's participation.
- Select a leader and a team.
- Identify customer expectations.
- Analyze process flow and measures.
- Define process inputs and outputs.
- Document the process.
- Identify process critical success factors.
- Determine data collection elements.
- Develop a preliminary questionnaire.

Some of the questions that must be answered during this phase of a benchmarking study include

- What process should we benchmark?
- What is our process and how does it work?
- How do we measure it?
- How well is it performing today?
- Who are the customers of our process?
- What products and services do we deliver to our customers?
- What do our customers expect from our process?
- What are the critical success factors for this process?
- What is our process performance goal?
- How did we establish that goal?
- What data should we collect for comparisons?

Notice that many of these questions are similar to the basic questions that one confronts in any TQM improvement project.

Benchmarking: Collect

The steps involved in collecting data for a benchmarking study include:

- Collect internal data.
- Perform secondary research.
- Develop partnership criteria.
- Identify benchmark partners.
- Plan data collection.
- Develop survey or interview guide.
- Solicit participation of partners.
- Collect preliminary data.
- Conduct site visits.

Some of the questions that must be answered during this phase of a benchmarking study include

- What companies perform this process better?

- Which company is best at performing this process?
- What can we learn from that company?
- Who should we contact to participate as our partners?
- What is their process?
- How representative is the process across different areas of their organization?
- How do they measure process performance?
- What is their performance goal and how was it set?
- How well does their process perform over time?
- Is there any difference in performance at different locations or based on seasonal change?
- What business practices, methods, or tasks contribute to the process performance?
- What factors could inhibit the adaptation of their process into our company?

Notice that many of these questions are the same as questions that would be addressed in any TQM-based data collection activity.

Benchmarking: Analyze

The steps involved in analyzing benchmarking study data include

- Aggregate the data across units but preserve subgroup relationships.
- Normalize performance to a common performance base (e.g., σ scale).
- Compare current performance to historical data.
- Test performance for difference in both trend of performance average and variation.
- Test differences for statistical significance.
- Identify gaps in performance and the root cause of all significant differences.
- Identify entitlement opportunities.
- Forecast performance observations to the business planning horizon.
- Develop case studies of best practice.
- Isolate process enablers.
- Assess adaptability of process enablers.

Some of the questions that must be answered during this phase of a benchmarking study include

- What is the basis for comparing our process measurements?

8 Benchmarking in Project Definition

- How does their process performance compare with our process performance?
- What is the magnitude of the performance gap?
- What is the nature or root cause of the performance gap?
- How much will their process continue to improve?
- What characteristics distinguish their process as superior?
- What activities within our process are candidates for improvement?

Note that many of the questions addressed above are the same as would be addressed in any TQM or Six Sigma improvement project.

Benchmarking: Improve

The steps involved in implementing the lessons learned from a benchmarking study include

- Set goals to close, meet, exceed the gap.
- Modify enablers for implementation.
- Gain support for change among all involved parties.
- Develop the action plan.
- Communicate the plan.
- Commit resources to achieve the plan.
- Implement the improvement plan and document the changes.
- Monitor and report progress on targeted schedule.
- Identify opportunities for further process improvement.
- Recalibrate the benchmark measure after implementation.

Some of the questions that must be answered during this phase of a benchmarking study include

- How does our knowledge of their process help us to improve our process?
- How should we forecast the future effectiveness of their process performance?
- Should we redesign our process or reset our performance goal based on this benchmark?
- What activities in their process need to be modified to adapt it into our business model?
- What have we learned during this study that will allow us to improve on “best” practice?
- What goals should we set for our own process improvement?

- How can we implement the changes in our process?
- How will other companies continue to improve this process?

Note that many of the questions addressed above are the same as would be addressed in managing implementation in any project improvement process.

Methods of Data Collection

In conducting a benchmarking study, there are several different approaches to data collection that can be pursued by a benchmarking team. Tables 3 and 4 illustrate several schemes for data collection and identify when to use each approach as well as the advantages and disadvantages associated with each of the methods.

Presenting Benchmarking Results

Some final points should be made about the process of benchmarking relative to the analysis and presentation of benchmarking data. Care must be taken in the data analysis efforts to assure that benchmarks are representative of real-world performance. Specific cautions include the following statistical problems in benchmarking:

- single data point measurements that are passed off as “benchmarks”,
- **measurement systems** not validated for sensitivity of observation or **calibration**,
- averages used to represent performance benchmarks,
- missing variation data in process characterization,
- components of variance not identified according to their source,
- comparative charts not indicating *both* mean and variance,
- process changes not correlated with performance shifts, and
- interactions not identified among the different process variables.

Clearly, there can be many issues that create problems in the measurement and analysis of results from benchmarking studies. Careful planning and solid data collection and analysis efforts can achieve the elimination of these opportunities

Table 3 Data collection methods–1

Method	Existing data review	Questionnaire/survey	Telephone survey
Definition	Analysis and interpretation of data that already exists in-house or in the public domain	Written questions sent directly to benchmarking partner – can contain any type of question: multiple choice, open-ended, forced choice, or scaled answer	A written script of questions that are used to solicit data or guide a conversation during a telephone call to gather data and information
When to use	Before conducting original research to establish what is the historical baseline	When you need to gather the same information from a large number of sources	When information is needed quickly or to conduct a fast screen of potential sources
Advantages	A large number of sources of information are available and most are accessible from your own organization	Permits extensive data gathering over time, can be analyzed by computer (if OCR form is used), and data is easy to compile	Can cover a wide range of respondents quickly, people are more candid over the phone and can also amplify answers in follow-up questions
Disadvantages	Gathering this information can be very time consuming	Response rates are low; and interpretation of questions can be subjective, creative ideas rarely surface, probing for “how to” answers is difficult	Locating the right person to answer, no exchange of process information occurs, and often requires multiple phone calls

Table 4 Data collection methods–2

Method	Interview	Focus group	Site visit
Definition	A face-to-face meeting with a benchmarking partner using questions that are prepared and distributed in advance	A panel discussion between benchmarking partners with a third-party facilitator at a neutral location	An on-premises meeting at a benchmarking partner facility that combines an interview with work process observation
When to use	When you need one-on-one interaction to probe and drive data collection to a specific objective or level of detail	When you want to gather information from more than one source or perspective at the same time – when there are diverse opinions or ways to approach an objective	When you need to observe specific work practices
Advantages	Encourages interaction, in-depth discussion and open-ended questions – a flexible style can provide unexpected information	Direct sharing of information on best practice – brings the partners together to discuss a mutually established agenda	When interpersonal or face-to-face interaction is needed to evaluate the “human aspect” of a process Can observe actual practice and verify the process capability, enablers, and assess the measurement system
Disadvantages	Takes time – the people interviewed may be reluctant to talk about sensitive issues	Logistics must be managed – result may be the “lowest common denominator” among the participating partners	Requires careful planning and preparation – who asks what of whom?

for error introduction into a benchmarking study. Whenever possible, analysts conducting benchmarking projects should have the same education as Six Sigma Black Belts in statistical analysis to assure the analytical soundness of study results.

Benefits and Pitfalls of Benchmarking

Benchmarking is a business change process that encourages managed change. It encourages an organization to take an external and objective perspective in evaluating its performance. The benefit of benchmarking comes from three specific actions:

- The gap between internal and external practices creates the need for change.
- Understanding the benchmarked best practices identifies what must change.
- Externally benchmarked practices provide a picture of the potential result from change.

However, lest benchmarking appear to be a panacea for problem resolution, the following set of potential pitfalls in conducting benchmarking studies must also be highlighted:

- selecting benchmarking partners that do not convince management (not respected);
- choosing benchmarking partners to meet popularity tests with no performance substance;
- accepting public relations claims as process performance benchmarks;
- assuming that measurements are the same in different organizations (without checking);
- identifying process measures that are not traceable from strategic to operational levels;
- conducting statistical analyses that represent surface observations – not root **causality**;
- failure to validate performance with on-site inspection to verify benchmark claims;
- enforcing implementation of a benchmarking lesson across a cultural barrier; and
- use of “benchmarks” for management decisions without recalibration over time.

Comparative Analysis and Competitive Advantage

What does an organization gain in the way of competitive advantage from benchmarking? In the long

run, competitive advantage comes from outthinking and outperforming competition. When an organization uses benchmarking effectively, they are able to think ahead of their industry and to act efficiently by adapting lessons learned from cross-industry studies to permit them to creatively imitate the best performing processes in the world. Over the long haul this can establish them as the thought-leader within their own industry. In the final analysis, it is not outthinking or prior knowledge that results in competitive advantage, it is in the excellence of execution of such new knowledge and the creative application of breakthrough insights that wins in the long term. To achieve a dominant position in a market, a company must both know and do better than its most aggressive competitors. Benchmarking can help develop the competence to achieve this position, but it must be supplemented by management will in order to make success happen.

Integrating Benchmarking with Strategic Planning

The leading organizations in the world plan to win and win by planning. Benchmarking improves performance for these organizations by providing them with a methodology to learn and challenge their critical business assumptions by thinking differently about their strategic direction. Applying benchmarking as a tool of business strategy is an effective way to evaluate options and perform an assessment of alternatives by considering the strategic implications that may be observed in other analogous situations. Such lessons reduce the likelihood of “repeating the mistakes of others”.

Conclusions

Process benchmarking is an important tool of TQM and can help to improve both the strategic direction and the operational process performance. It is a discovery methodology that is used to stimulate learning and help organizations to think about creative options for the design and implementation of its business processes. Coupled with solid statistical data analysis, best practice identification and cultural adaptation can help organizations to both “learn” and “do” their business process improvement more effectively.

Reference

- [1] Watson, G.H. (1993). *Strategic Benchmarking*, John Wiley & Sons, New York.

Further Reading

- Camp, R.C. (1995). *Business Process Benchmarking*, ASQ Quality Press, Milwaukee.
 Camp, R.C. (1989). *Benchmarking*, ASQ Quality Press, Milwaukee.
 Deming, W.E. (1982). *Out of the Crisis*, Massachusetts Institute of Technology, Center for Advanced Engineering Study, Cambridge.

- Hammer, M. & James, C. (1993). *Reengineering the Corporation*, Harper, New York.
 Porter, M.E. (1980). *Competitive Strategy*, The Free Press, New York.
 Porter, M.E. (1985). *Competitive Advantage*, The Free Press, New York.
 Watson, G.H. (1992). *The Benchmarking Workbook*, Productivity Press, Portland.
 Watson, G.H. (2007). *Strategic Benchmarking Reloaded with Six Sigma*, John Wiley & Sons, Hoboken.

GREGORY H. WATSON

Six Sigma

Introduction

There are many ways in which people try to define Six Sigma; some emphasize the philosophy, others the statistical tools, some talk about the measurements and the 3.4 parts per million defect level targets, and still others focus on the similarities and differences with **total quality management** (TQM). Perhaps the simplest and most correct way to think of Six Sigma is the current state of the art of quality management. Six Sigma is more an evolution of best practices in quality management than being something entirely new.

More than 20 years ago, when Motorola coined the term *Six Sigma*, it was after years of adding new tools, methods, and concepts to what they had learned first from the Japanese in **total quality control** and then from many leading companies in the United States, and throughout the world. Many companies follow similar paths starting with specific tools and proceeding to develop comprehensive quality systems. The normal path was as follows:

- quality circles
- statistical quality control
- ISO 9000:1987 system of quality standards
- continuous quality improvement/*Kaisen*
- reengineering
- theory of constraints
- Malcolm Baldrige National Quality Award criteria or European Foundation for Quality Management criteria
- total productive maintenance
- supply-chain management
- ISO 9000:2000 system of standards
- lean manufacturing – Toyota Production System
- Six Sigma
- Design for Six Sigma
- lean Six Sigma
- ? – we are not sure what will be next, but are sure something will be.

It is not surprising to see such a list. Companies throughout the world are in an intense competitive race. The easy communication we now have through conferences, **benchmarking** visits, and the Internet makes the spread of new ideas and best practices

faster than ever. If a method, tool, or management practice improves the performance of a company or even part of a company, it is now implemented quickly. Senior managers are not interested in the name of a program, they are simply interested in what works.

We will not try to trace the entire evolution of Six Sigma here, but just cover the most important events. Some years ago George Box presented a wonderful summary of the evolution of science. For thousands of years progress was slow. New scientific breakthroughs were often by chance, an “intelligent observer” being in the right place at the right time. A significant event occurred, someone asked why, and after some experimentation, began to understand the cause and effect. W. C. Roentgen’s discovery of X rays is a classic example.

By the late 1800s, this was beginning to change. Thomas Edison built his laboratory in 1887 to make things happen, not wait for them to occur accidentally. He ran thousands of experiments in a well-planned fashion. Throughout the twentieth century industrial, academic, and government laboratories, often employing tens of thousands of researchers, became common. The art and science of experimental design became widely used to drive improvements in products and processes and in developing entirely new products and even services. In almost every leading organization, teams of researchers well trained in statistical methods, industrial engineering, and operations research supported these researchers and operations people throughout the organization in running experiments, analyzing data, and making significant improvements.

In the latter half of the twentieth century, another phenomenon took place, first in Japan and then quickly in other parts of the world. Large numbers of employees in organizations were taught the basics of the scientific method and given a set of tools to make improvements in their part of the company, changing the process and even the product to achieve continual improved production and product performance. This basic teaching of the scientific method for quality improvement was first documented in English by Joseph M. Juran in his classic 1964 text, *Managerial Breakthrough*. The steps he laid out were straightforward:

1. breakthrough in attitudes
2. the **Pareto** principle

2 Six Sigma

3. mobilizing for breakthrough in knowledge
4. the steering arm
5. the diagnostic arm
6. breakthrough in knowledge – diagnosis
7. resistance to change – cultural patterns
8. breakthrough in performance – action
9. transition to the new level.

He then stressed the importance of control – holding the gains. The second part of the book he devoted to quality control. Later in *Quality Planning* and *Quality by Design* he added the importance of design outlining many of the tools and techniques we now widely use in the Design for Six Sigma.

Many Japanese companies had made these methods a standard practice and taught them to all employees. An example is Komatsu, the heavy equipment manufacturer. They had a small pocket book printed to easily remind each employee of the steps for improvement. An English version of this small pocket book was created by GOAL/QPC. More than 10 million copies of this small guide have been purchased by leading companies for use in training their people in the scientific method and quality improvement. The steps in the guide are the following:

1. Decide which problem will be addressed first (or next).
2. Arrive at a statement that describes the problem specifically.
3. Develop a complete picture of potential causes.
4. Agree on the basic cause(s) of the problem.
5. Develop an effective solution.
6. Implement solution and control.

Some years later these steps have evolved into the basic quality improvement methodology of Six Sigma.

- *Define* – understand what is critical to quality, the process, and define the problem and the action plan.
- *Measure* – measure all critical aspects of the process, perform the **measurement system** analyses, focus on what is important.
- *Analyze* – identify the factors that truly drive performance – the Xs.
- *Improve* – establish the optimum levels for each factor, define the functional relationship,

$$Y = f(X).$$

- *Control* – define the control plan, hold the gains.

The improvement structure of Six Sigma is not different, it is just a refinement of continuous quality improvement methods used for years by leading organizations – the scientific method widely taught with useful tools for use at each step. The question then is what is different about Six Sigma? We shall focus on seven major differences.

Philosophy

The most startling factor in Six Sigma for many is the focus on perfection. In 1994 Robert Galvin, the then chairman and chief executive officer of the Motorola Company, startled a group of executives during a dinner talk by stating, “We now believe perfection is possible.” Many writers about Six Sigma have confused Motorola’s calculation of “Six Sigma quality” as 3.4 parts per million defects with Galvin’s philosophy of perfection. The 3.4 parts per million defects are just Motorola’s practical realization that for some products perfection was quite difficult and therefore they needed, at that time, a more realizable goal. They had defined perfection as a product that was on target with respect to its specifications and had six standard deviations (Six Sigma) from the center target to both the upper specification limit and the lower specification limit. With the basic assumption of an underlying normal (Gaussian) distribution, this process is close to perfect – only two parts per billion would be out of specifications.

Someone in the company soon realized that not all processes are exactly centered; some may shift for a variety of reasons and not be exactly in the true center of the specification range. A study in one of Motorola’s plants found that their process shifted as much as 1.5 standard deviations over time thus becoming 4.5 standard deviations from one limit and 7.5 standard deviations from the other limit. A quick calculation showed that this worst-case analysis created a process with 3.4 parts per million out of specification. A Whole lot of companies have adopted this shift assumption and now consider Six Sigma quality as 3.4 parts per million defects. The true philosophy underlying Six Sigma is the relentless pursuit of perfection.

Financial Focus

A weakness in many quality initiatives in the past including the way total quality management was practiced in some organizations was a lack of hard numbers on the bottom-line results of various activities. The leading organizations practicing Six Sigma have an almost fanatical focus on quantifying actual results for every project using techniques such as “hard numbers” and “soft numbers”. The hard numbers are ones that can be proved to directly result from completed projects – reduced structural costs, reduced labor, reduced materials, decreased warranty payments, reduced machine downtimes, increased number of customers – actual savings. The soft numbers are not usually counted in terms of the actual savings but calculated as potential benefits, hard to quantify, but definitely positive results. These soft numbers often include increased customer satisfaction, reduced cycle times, improved employee morale, faster development times, and so on. Most companies now include people from the finance and/or accounting departments on project teams. Project values are estimated before projects are launched, and the teams work hard to document actual savings at each step. Return-of-investment calculations are critical for each project.

Training

For many companies the training methods now used in Six Sigma are completely different from how employees were trained in the past. People are selected for training after the leadership of the organization has selected the projects or project areas. Teams are trained together. No one comes to training without a project. During the training cycle they work on the project. The goal is to finish the project shortly after the training is complete. Tools, methods, and concepts are applied to the project as soon as they are learned. Teams present their projects during every week of training receiving valuable feedback, criticism, and ideas from the instruction team and the other participants. The goal is not to complete the training; the goal is to solve a problem with significant financial results. The emphasis is not on understanding statistical methods but to know to use sophisticated tools, statistical and nonstatistical, and to solve complex problems.

A new structure has evolved for the training. The team leaders, the fundamental building blocks of Six Sigma, are the black belts. These people receive 4–5 weeks of intense training in problem solving, statistical methods, project management, and team leadership. The companies producing significant results in Six Sigma select these people from the pool of future company leaders and make their assignments as black belts full time for several years. Some of these black belts decide to make this their company performance path and become permanent problem solvers or task force leaders, and some go on to become master black belts gaining other advanced skills and becoming trainers and consultants within the organization.

Team members are usually trained as green belts in a 2-week course again focusing on problem solving and working during training on selected projects. Green belts will often lead small project teams within the organization and work as team members under the leadership of a black belt on large projects. Many companies now have “white belt”, “yellow belt”, or even “red belt” training programs to give a basic understanding of Six Sigma methods to every member of the company. These are usually 2- or 3-day sessions focused on the basic problem-solving methodology (define–measure–analyze–improve–control (DMAIC)) and an introduction to the basic tools.

Software

One of the striking differences in the application of Six Sigma is the widespread use of sophisticated software. Several leading software companies now sell outstanding software packages directly supporting Six Sigma projects. This software makes sophisticated problem solving possible by all members of the organization. The software is not taught as a software package, it is totally integrated into the green belt, black belt, and master black belt training. As problem steps are taught, the appropriate tools are also taught. The class uses the software to solve example problems, work on complex case studies, and to work on their own problems.

Most leading companies have standardized on selected software packages to enable rapid sharing of best practices, easy understanding by all project reviewers, and enhanced support for all employees. These software packages have continued to

evolve over the past decades becoming more and more focused on the tools needed by large numbers of employees engaged in breakthrough quality projects and developing new products, services, and processes.

New Tools

For many people the most obvious new addition is the plethora of new tools—both statistical and non-statistical—introduced in Six Sigma classes and practice. Most of these tools are not truly new. The sophisticated, easy-to-use software has just made it practical to make these tools available to every level of employees in all parts of the organization. Simple “hand tools” that were taught in earlier days of continuous quality improvement, statistical quality control, or total quality management have now become practical for large datasets of thousands or even millions of data points. Other tools, for example multiple regression, that were thought to be in the purview of only advanced-degree specialists are now commonly used by people from all parts of a company. New tools are continually being added to both training classes and software as they prove useful. The problem-solving “toolbox” has been greatly expanded. A whole lot of books have been written describing these tools, giving examples of their use, and step-by-step instructions on how to use.

Structured Improvement Process and Design Process

Companies have created very structured improvement and design processes that allow much innovation and creativity but are also easy to teach and understand. Most use some form of the DMAIC process described above for improvement projects (breakthrough quality), but there seem to be slight variations in the Design for Six Sigma project steps. Some of these variations come from companies trying to match the new ideas of Six Sigma with their current design processes.

The advantages of the structured improvement and design processes are obvious. All levels of employees throughout the organization know how to lead a project from start to finish. Communication across the organization is greatly enhanced, people can

quickly learn from others working on similar issues. Managers easily review large numbers of projects, thus becoming involved in the critical implementation steps. Replication of results across many areas of the organization is facilitated.

Leadership

The driving force in successful Six Sigma implementation has not been different from the driving force in successful total quality management or other initiatives. Leadership is critical. What has been different is the amount of attention paid to senior management, executive leadership, and specific training programs developed just for the senior management team. Many companies now have 2- or 3-day special courses just for project “champions”, leaders specifically selected to oversee black belt and green belt projects. This training usually covers the basic methods (DMAIC) and introduces the statistical tools, but focuses on the actual roles of senior managers in getting the results.

This focus includes the role of senior managers in selecting the project areas and specific projects and selecting the black belts and green belts who will be involved. The senior managers understand their role in providing resources to assure that the projects are successful. They understand their roles in the implementation of the company’s organizational or process changes that will be necessary to achieve the desired results. They know to deal with the various sources of resistance to change the status quo in any organization. They become truly involved in all aspects of the Six Sigma initiative becoming major drivers in the success of projects.

Summary

In many ways Six Sigma shares most of the ideas, concepts, tools, and methods developed over the past hundred or so years in quality management. It truly is far more an evolutionary change than a revolutionary one. However, the differences discussed above have made it far more effective for a large number of organizations throughout the world than any approach to quality management or organization excellence than ever before.

Further Reading

- Breyfogle, F.W. (2003). *Implementing Six Sigma: Smarter Solutions Using Statistical Methods*, 2nd Edition, John Wiley & Sons, Hoboken.
- Chowdhury, S. (2002). *Design for Six Sigma*, Dearborn Trade Publishing, Chicago.
- Feigenbaum, A.V. (1991). *Total Quality Control*, McGraw-Hill, New York.
- George, M.L. (2003). *Lean Six Sigma for Service*, McGraw-Hill, New York.
- Gitlow, H.S., Levine, D.M. & Popovich, E.A. (2006). *Design for Six Sigma for Green Belts and Champions*, Prentice Hall, Upper Saddle River.
- Godfrey, A.B., Stephens, K.S. & Wadsworth, H.M. (2002). *Modern Methods for Quality Control and Improvement*, 2nd Edition, John Wiley & Sons.
- Goel, P.S., Gupta, P., Tyagi, R.K. & Jain, R. (2005). *Six Sigma for Transactions and Service*, McGraw-Hill, New York.
- Gupta, P. (2005). *The Six Sigma Performance Handbook: A Statistical Guide to Optimizing Results*, McGraw-Hill, New York.
- Hayler, R. & Nichols, M.D. (2005). *What Is Six Sigma Process Management?* McGraw-Hill, New York.
- Hayler, R. & Nichols, M.D. (2007). *Six Sigma for Financial Services*, McGraw-Hill, New York.
- Hoerl, R. & Snee, R.D. (2001). *Statistical Thinking: Improving Business Performance*, Duxbury, Pacific Grove.
- Juran, J.M. (1964). *Managerial Breakthrough*, McGraw-Hill, New York.
- Juran, J.M. & Godfrey, A.B. co-editors (1999). *Juran's Quality Handbook*, 5th Edition, McGraw-Hill, New York.
- Levine, D.M. (2006). *Statistics for Six Sigma Green Belts*, Prentice Hall, Upper Saddle River.
- McCarty, T., Daniels, L., Bremer, M. & Gupta, P. (2005). *The Six Sigma Black Belt Handbook*, McGraw-Hill, New York.
- Mizuno, S. (1988). *Company-Wide Total Quality Control*, Asian Productivity Organization, Tokyo.
- Pande, P.S., Neuman, R.P. & Cavanagh, R.R. (2000). *The Six Sigma Way*, McGraw-Hill, New York.
- Pande, P.S., Neuman, R.P. & Cavanagh, R.R. (2002). *The Six Sigma Way Team Fieldbook: An Implementation Guide for Process Improvement Teams*, McGraw-Hill, New York.
- Sleeper, A. (2006). *Design for Six Sigma Statistics: 59 Tools for Diagnosing and Solving Problems in DFSS Initiatives*, McGraw-Hill, New York.
- Smith, D., Blakeslee, J. & Koonce, R. (2002). *Strategic Six Sigma: Best Practices from the Executive Suite*, John Wiley & Sons, Hoboken.
- Tang, L.C., Goh, T.N., Yam, H.S. & Yoap, T. (2006). *Six Sigma: Advanced Tools for Black Belts and Master Black Belts*, John Wiley & Sons, Hoboken.
- Tennant, G. (2002). *Design for Six Sigma: Launching New Products and Services without Failure*, Gower Publishing, Burlington.
- Wang, J.X. (2005). *Engineering Robust Designs with Six Sigma*, Prentice Hall, Upper Saddle River.

A. BLANTON GODFREY

Total Quality Management (TQM)

Introduction

Many associate the philosophies of **quality management** with a number of our better-known quality gurus. And while that is certainly true to an extent, there have been many contributions to quality management philosophies, beliefs, and tested truths. These philosophies have by now also undergone evolution over long time periods, with the greatest number of additions formulated and placed in practice over the last century. In fact, the quality discipline has matured so well that there is now a large body of knowledge pertaining to associated philosophies.

The quality discipline itself has been like a “learning organization”, learning by mistakes (e.g., assigning inspection alone as a control on production, and separate from the production department, essentially removing responsibility for quality from production and applying it too late in the process stream); learning from application experiences and the sharing of these experiences in professional societies, seminars, publications, and many forms of media; formal university learning on the discipline; learning from **benchmarking** best practices; learning from other related disciplines such as leadership theories, motivation theories, organizational development, and so on; learning from the application of standards and awards such as the ISO 9000 Quality Management System, and the US Malcolm Baldrige National Quality Award for business excellence, as well as many others in Europe, Japan, and other countries. Numerous innovations of the philosophies and techniques of total quality management (TQM) have come from all parts of the world, with notable additions by the Japanese in their postwar rise to prominence.

Total Quality

Central to the term *total quality management* and to its associated philosophies is the meaning of “quality”. What is quality as it applies to the TQM philosophies? An important TQM philosophy is how we define and apply the concepts of quality. Generally, most dictionary definitions do not include sufficient

breadth of the concept as applied to the quality disciplines. Many dictionary definitions of quality do, of course, address the general subject of quality – but they lack strategic and/or operational meanings that can be used to advantage by an entrepreneur/manager in structuring products or services and processes that contribute to successful businesses. And while there are numerous contributions to defining quality in the context of the quality disciplines, this brief *exposé* cannot explore them all. One rendition of an attempt at this is by the author [1], and is expressed diagrammatically in Figure 1. The gist of this diagram is to emphasize that in the quality disciplines quality has many aspects.

The Fundamental Aspects of Quality

There is the quality of design – usually defined by way of specifications, standards (internal and/or external, company, national, regional, and/or international), grades, drawings, and so on. We refer to this as the fundamental aspects of quality and picture it as a core element in the overall definition of quality. It should be understood in this dialog that quality, in and of itself, doesn’t necessarily mean high quality to the degree of unnecessaryness; it means uniformity, consistency, and conformity to the standard or specification – a statement of what the user wants and can afford *and* what the producer can provide. Note in the fundamental aspects of quality (design quality), the reference to product “grade”. The quality statement or specification should be based on what the existing process can produce economically with reasonable controls and improvements (otherwise recognition that a new process must be designed and implemented). Hence, the producer and user must cooperate in defining a practical, reasonable, and economical specification of quality. This is, in fact, an important part of quality planning – to match the needs of the customers with the capability of the process (existing or new) and product design, taking into consideration the intended use of the product.

Two Contrasting Examples

For example, if a producer of plastic measuring and drawing rulers were to specify a requirement of ± 0.5 cm in a length of 30 cm, then even with good conformity to this requirement, the rulers would be unfit for reasonable use. Conversely, if consumers

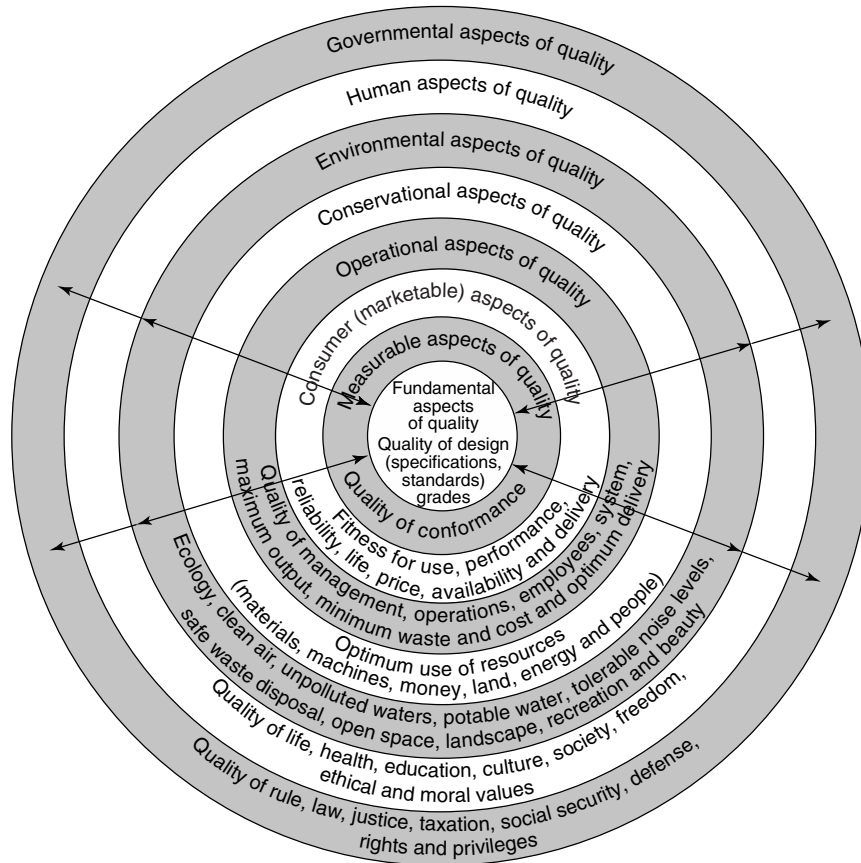


Figure 1 Concepts and expansion of the influences and the areas of applicability of quality

were to request a 0.001-mm tolerance on the diameter of the body of ballpoint pens, their manufacture would become extremely impractical and/or uneconomical, while such a tight specification would fail to provide any extra fitness for use.

The Measurable Aspects of Quality

Next are the measurable aspects of quality, which include the concept of conformance to the standards, specifications, and requirements that are part of the fundamental aspects. Here we define important quality characteristics, often referred to as the *building blocks* of quality. In this section we address different types of characteristics such as variables, attributes, visual, sensory, and the like. And these apply to both products and processes, including management

processes. Hence, embedded in this aspect are important principles such as “fact-based” decision making as opposed to opinion-based decisions. It enhances product and process control as well as continuous improvement (CI). This aspect of quality also reaches out to bring into the TQM arena, the whole field of metrology – the science of measurement, including, across broad ranges of products and services, scientific metrology, process metrology, and legal metrology. If we are to make good and effective use of collected facts and data, they must be accurate, precise, reproducible, repeatable, and meaningful. The conformance aspect of quality was emphasized by Phil Crosby and it could represent a fairly “complete” definition of quality under the assumption that the “conformance” is to requirements that reflect sound design; customer wants, needs, and desires; and interests of other stakeholders. As we expand the diagram

of Figure 1, it needs to be pointed out that inner aspects influence outer aspects, and considerations of properties of outer aspects, in turn, influence inner aspects. This is reflected by the double directional arrows that are shown in the diagram to emphasize that each aspect is influenced and modified by the other. For example, while measuring conformance we may discover areas for which the specifications (quality of design) must be modified or improved. Measurement also touches on the broader management aspects, discussed subsequently, and includes management information systems (MIS), uses of such techniques as the balanced scorecard, the dashboard technique, the measurement of costs of poor quality (COQ), and various measures of customer and employee satisfaction.

The first two aspects of quality are particularly important for a proper understanding of the “as conceived” and “as made” quality of products and services. Concern for quality is often not exercised adequately, particularly by production and management personnel, because these two aspects of quality are not known or understood properly. In such a setting, these aspects of quality are known as “design quality” (as labeled above) and “manufactured quality”, respectively. As discussed subsequently, both should be aimed at providing quality satisfactory to the consumer – the “as delivered” and “as perceived” quality. Therefore, it can be seen that design quality is not the only aspect of quality. After a product has been designed, even to meet customer expectations and placed in manufacture, we find that the process does not always produce each unit of product (or each instance of service) in conformity to the design – defects in materials, parts, subassemblies, assemblies, and in final product arise. These defects, of course, may be related to the design of the product. The correction, reduction, improvement, and/or prevention, may call for a redesign of the product. But all too often these nonconformances are the result of incomplete, inadequate, poorly planned, and poorly controlled processes. The products or services that fail to perform as intended, most often result in customer dissatisfaction. Any defective or nonconforming materials, parts, assemblies, and finished products that are discarded or reworked during the manufacturing process, result in increased cost. And incidentally, these costs are often not identified or recognized as quality costs by owners/managers who otherwise consider “quality” too costly. The lost time and effort in

producing these defective or nonconforming parts and products, the loss of business resulting from delays in delivery, and other associated ills of poor quality products are the consequence of “manufactured quality” (i.e., the quality of conformance to the design). We will continue this discussion in still another aspect of quality subsequently.

The Consumer Aspect of Quality

We then move to the “consumer” or marketable aspects of quality. In this aspect the “customer” is king! This was emphasized by Karou Ishikawa and others. Related characteristics of this aspect are such things as “fitness for use” emphasized by Joe Juran; performance (as opposed to conformance); reliability; life; price (yes, price is a “quality”); availability; and delivery. This alludes to the philosophies of customer satisfaction, customer retention, building a base of loyal customers, making them a part of your “external sales force”, customer satisfaction/retention measurements, and tools and techniques such as focus groups, surveys, “voice of the customer” and the use of **quality function deployment (QFD)** (which also impacts other aspects of quality). Also useful here is Noriaki **Kano**’s “customer quality model” in which he emphasizes dissatisfiers, satisfiers, and excitors/delighters as well as the dynamics of these characteristics over time, requiring diligence in these efforts.

There are numerous examples of products (and services) that have been designed, specifications prepared (usually by a design engineering department), placed in production, conformance to specification checked, and confirmed – only to find that the products (or services) did not meet with customer satisfaction. Is it possible for specifications to be prepared without reflection on the needs, requirements, or desires of present, or potential customers? Consider the case of the US automobile manufacturers in the 1960s–1980s, and even beyond. Specifications do not always reflect customer needs and desires. We may have *conformance*, but do we have satisfactory *performance*?

Example

When wooden furniture is made in a tropical, humid climate for distribution to a frigid or temperate dry climate, it is the producer’s responsibility to learn

something about the effects of this climate, or for that matter, any artificially produced environment such as forced hot and dry air, on the long-term performance of the product. It may be necessary to include additional operations (e.g., kiln drying of raw materials) or controls (e.g., selection of aged raw materials) in the manufacturing process. If cracks develop in the product after some period of use and the producer is irresponsible in this regard, a decline of business after the initial orders should be expected, deservedly so. In similar situations, costly recalls, **warranty claims**, or liability judgments may be incurred, depending on the nature of the product or service.

It is one aspect of quality to establish a specification based on a quality of design. It is another aspect of quality to determine if the specification is being met – that there is conformance to the specification. It is still another aspect of quality that the specifications – as written and implemented – reflect the needs, requirements, and desires of the customers (or at least a majority of present and potential customers).

The Operational Aspects of Quality

Next we move to the operational aspects of quality. It is here that we expand well beyond the usual dictionary definitions of quality, beyond products, or entities, to the quality of management, operations, employees, systems, and the like. And it is here that our concepts of quality begin to take on a more “total”, even strategic, dimension. Characteristics of this aspect of quality include such things as maximum output (yes, productive output, throughput, cycle time, and the like are qualities!). Quantity is, in fact, a characteristic of the operational aspects of quality! A characteristic that should be monitored, controlled, and improved – an integral part of quality, as some might say, with a big “Q”, rather than with a little “q”. Other characteristics of the operational aspects of quality are minimum waste, minimum scrap, minimum repair, or rework, minimum downtime, minimum absenteeism, and certainly, minimum cost, as well as optimum delivery, and the like. Here an important philosophy is that optimum production, quality, and cost can be planned, established, improved, and achieved simultaneously. As emphasized by Deming this strikes at the core of some now (hopefully) obsolete theories of “trade-offs” of these characteristics. If quality, that includes

this aspect, is implemented properly with planning, control, and improvement, it is entirely feasible to achieve these characteristics simultaneously. And it is with respect to this aspect of quality that a great many philosophies of TQM have come to light.

If we examine these characteristics, it is not difficult to understand how quantity (e.g., maximum output) can be achieved with a correct approach to quality. Recall our earlier discussion regarding the characteristics of “manufactured quality”. When our quality orientation includes a constant effort (and associated techniques) to reduce waste, scrap, rework, downtime, absenteeism, and others, more output is a natural result at lower costs. For every defective unit prevented, another good unit is produced and available for sale. For every day of absenteeism prevented, another person-day of output is produced. Yes, our overall definition of quality must address these matters as well as those of the actual final products and/or services.

Included in this is the whole realm of variability reduction, control and improvement *via* statistical methods; statistical **quality control**, including use of **control charts**, designed experiments, and numerous such tools; process control, including concepts of **process mapping** and flowcharts, **Pareto** analysis, Ishikawa diagrams, Poka-yoke (mistake proofing), the “five-whys” technique, the “triple-role” principle, cause-and-effect analysis, and the like. The recognition that defects and/or nonconformities are caused and we can seek and find their causes is essential. And here a related quotation by Shakespeare brings this concept to the forefront from so many years ago, “*Find out the cause of this effect, or rather say, the cause of this defect, for this effect defective comes by cause*”, from Hamlet, Act ii, Sc. 2. An important philosophical understanding in this connection is the type of cause, so essential to the approach taken to locate and eliminate unwanted effects, namely, common (or chronic) causes and special (sporadic) causes, with the first set of terms used by Deming and the matching ones in parentheses used by Juran. The following quotation from Deming [2] is applicable here, “**Special causes and common causes.** A fault in the interpretation of observations, seen everywhere, is to suppose that every event (defect, mistake, and accident) is attributable to someone (usually the one nearest at hand), or is related to some special event. The fact is that most troubles with service and production lie in the system. Sometimes the fault is

indeed local, attributable to someone on the job, or not on the job when he should be. We shall speak of faults of the system as common causes of trouble, and faults from fleeting events as special causes." Further principles and disciplines involve sampling inspection, reliability, and many others. Here, we owe a great deal to pioneers and gurus such as Shewhart, Dodge, Fisher, Ishikawa, Taguchi, and others. While sampling inspection may be considered "passé" by certain elements of our profession, it is still being applied effectively in a great many enterprises to serve as a measure of control, improvement, and assurance, especially influenced by many of the regulated industries. The philosophy of CI is important here, achieved both by incremental step-by-step improvements (of the "**Kaizen**" type) as well as breakthroughs of the innovation type. **Six Sigma** programs and efforts are applicable here, including six sigma for improvement and cost reduction as well as Design for Six Sigma (DFSS). **Failure mode** and effects analysis (FMEA) to help reduce the necessity for extensive improvement efforts downstream is important, essentially moving much of process quality control and improvement upstream to the design process. Just-in-time (the JIT concept) is applicable as a production system, not just as an inventory reduction system, as many misunderstand. Juran's "Trilogy" alluded to earlier, involving planning, control, and improvement, is applicable. Elimination of the "hidden-plant" (existing to produce defective products) as emphasized by Feigenbaum is useful here. Cross-functional approaches, abetted by a "teams" approach are useful here. Deming's entire philosophy of his "profound knowledge" applies here, with emphasis on (1) appreciation for a system, (2) understanding of variation, (3) theory of knowledge, and (4) psychology, including his 14 points for management, and his seven diseases and obstacles.

As this concept of quality touches on the quality of management, employees, and systems, it includes the whole realm of managerial principles, tasks, practices, and skills as well as a "systems" approach. This includes strategic planning, including the recognition that "quality" as expounded here is a strategic element. It also includes leadership, managerial audits, system audits, and the like. It embraces the Malcolm Baldrige National Quality Award for business excellence, as well as the ISO 9000 Quality Management System procedures and requirements. It

includes the entire concept of "supply-chain" management. This involves relationships with suppliers, including the "partnership" concepts of Deming and "single-sourcing", where feasible. It includes organizational structure with a philosophy of flattening the "silo" effect of a high managerial hierarchy, in order to achieve more involvement from all levels of the organizational chain and reduce communications and errors related thereto. It also includes the philosophy of defining, communicating, and achieving for all stakeholders, the firm's vision, mission, guiding principles, and values. And while we will consider subsequently the human aspects of quality, from a management of employees point of view, this aspect of quality includes the whole realm of human resources management, recruitment of appropriate personnel, adequate training, employee involvement, suggestion systems, empowerment, fair wages, cross-functional and team deployment, recognition, awards, internal promotion, and measurement and enhancement of employee satisfaction with the additional philosophy that "a satisfied employee helps to achieve satisfied customers". It embraces the concept that quality is everyone's responsibility, not relegated to a specific function. And it emphasizes the necessity to manage quality as a shared responsibility, that every function will have various quality matters that are unique to that function.

The Conservational Aspects of Quality

The next aspect of quality from Figure 1 is that of the conservational aspects of quality. Here we are considering the optimum use of resources, including materials, machines, money, land, energy, and people. Not all waste in a business or other setting is the result of actual defective products. Materials, machine run-time and length of downtime, money, space, energy and labor of personnel can be wasted at any step in a process of production or service. How we use our resources is crucial to optimizing those resources. It is not difficult to realize that there is a strong link between this aspect of quality and the operational aspects discussed above. The same tools and philosophies of management apply. Much of this is attitudinal, but once the proper attitude toward quality has been established, actions to make optimum use of resources follow.

The Environmental Aspects of Quality

Still another part of our overall definition of quality is the environmental aspects of quality. This pertains to such things as ecology, clean air, unpolluted waters, potable water, tolerable noise levels, safe waste, open space, landscape (natural and developed), recreation, and beauty. Safety at work, home, and play is also included as an important characteristic. Therefore, our overall definition of quality includes the environmental aspects of quality as shown on Figure 1. Significant activities associated with this aspect of quality are being undertaken at the international and national levels. The International Organization for Standardization (ISO) has promulgated a series of standards on Environmental Management Systems, namely, the ISO 14000 series. These, like the ISO 9000 series on Quality Management Systems, are already generating a lot of attention and form the basis for specifying requirements, assessing compliance, and certifying satisfactory implementation.

The Human Aspects of Quality

While the quality of employees was considered as part of the operational aspects of quality in recognition that “our people are our most important asset”, the broader scope of the human aspects of quality may be considered separately. Characteristics of this aspect of quality include the quality of life, including the quality of work life, the quality of health, education, culture, society, freedom, ethical, and moral values. Therefore, our overall definition of quality includes the human aspects of quality, as shown on Figure 1 and figures significantly in the philosophies of TQM. Many of these characteristics of the human aspects can be included by forward-looking enterprises in creating an overall work environment in which employees enjoy and contribute to their productive capability and capacity. Other characteristics are challenges to larger segments of our society and emphasize that quality, indeed, is a multifaceted discipline.

The Governmental Aspects of Quality

What about the governmental aspects of quality? This aspect is shown as the last concentric ring on the diagram of Figure 1. Is the quality of how we are governed, regulated, taxed, policed, defended,

cared for, etc. of interest and importance to us? Is the “quality” of our rights and privileges of any interest to us? And these questions apply not only to the individual, but also to the corporation. Yes, the quality of our government and all of its actions and deliberations are important to us individually and collectively. The Sarbane–Oxley (SOX) act, and its implementation, is a further example of government involvement.

Total Quality Management

The preceding is the use of an extended description of “Quality” as an aid to help expose many of the philosophies of the meaning and practices of TQM. And in conclusion, quality can be seen to be a complex, multifaceted subject – not limited to a single dictionary (or other) definition. Attention to quality – in its broadest sense – and recognition that quality – as a discipline – offers sound business strategies and philosophies are evolutionary and revolutionary forces that have made (and are making) major contributions to our current way of life and progress. Managing this concept and aspects of quality is essentially what we mean by TQM. The following is one attempt to define the term.

TQM

The process that integrates fundamental management art and techniques with the principles and methodologies of total quality to develop and implement successful business strategies throughout the organization with emphasis on:

- a clear focus on customer needs, wants, desires, and specified and unspecified requirements;
- continuous planning, control, improvement, and innovation of all processes, services, and products;
- effective empowerment and recognition of individuals under a team involvement approach, including essential programs of learning, education, and training;
- sound planning for quality with fact-based decision making, variability reduction, defect prevention, and fast response systems that include recognition of the internal customer and the triple role of supplier, processor, and customer;
- integration of and mutual cooperation with suppliers based on long-term partnership relations;

- sensitivity to competitive comparisons *via*, e.g., benchmarking, including “best in class” on non-competitive but significant processes, and effective implementation of learning gained;
- productivity, cost reduction, profitability, and quality enhancements;
- strong leadership by management at all levels, sound strategic planning and deployment, and faithful dedication to ethical and moral principles and practices.

Benefits expected to be realized from TQ and TQM include the following:

- improving the quality of products and services to meet customer needs;
- increasing the productivity of manufacturing processes and commercial businesses;
- reducing manufacturing and service costs;
- determining and improving the marketability of products and services;
- reducing consumer prices of products and services;
- ensuring on-time deliveries and availability;
- assisting in the management of an enterprise.

Total quality as a strategy can bring better quality, price, delivery, sales, growth and profitability, leading to economic prosperity.

References

- [1] Wadsworth, H.M., Stephens, K.S. & Godfrey, A.B. (2002). *Modern Methods for Quality Control and Improvement*, 2nd Edition, John Wiley & Sons, New York.
- [2] Deming, W.E. (1982). *Out of the Crisis*, Massachusetts Institute of Technology, Cambridge, pp. 314.

KEN STEPHENS

Scorecards

Introduction

Conventional wisdom in business management has long held that “what gets measured gets done”. But because conventional wisdom does not always win the competition for time and resources, business processes cycle year after year with hardly more than a wet windy finger’s worth of directional assessment. This relegates management decisions to gut instinct and to reactions to obvious stimuli. Disciplined **data collection** and scorecarding can add clarity to an organization’s direction and goal setting.

Visible performance measurement, when coupled with attention from top management and shareholders, drives a pattern of execution and improves the quality and speed of decision making. We link decision making about improvement efforts to customer expectations, which is the essence of quality. And although a rush to having good data for planning and execution is a heady stuff, the first step of robust dialog with customers about what exactly are their expectations is a moment to savor and in which to invest (Just coming to a consensus about who exactly is the customer, or about how different customers should be prioritized, is often contentious. This is particularly the case in business support organizations, e.g., human resources, information technology, and facilities management).

How to Build a Scorecard System

Build your measurement system by identifying your customers and their strategic objectives. Customer identification should start with a definition of “customer”, generally the recipient of the output of a process and/or the funder of an activity. Business support organizations are typically funded by corporate budgets. And the direct recipients of the output of their processes are the management or staff, who help in defining service expectations but do not pay from their own budgets. Business support organizations with charge-back models can drive a higher level of alignment in this regard; but the cost of charge-back models often drives firms to avoid them. This makes clear performance metrics tied to customer expectations more important, and usually leads to the

broadening of the definition of customer to include important internal stakeholder groups (we shall use “customer” and “stakeholder” interchangeably in the remaining portion of this paper).

Business support organizations never have external customers and the customer identification conversation is the first opportunity for team alignment. Once customer groups are identified through brainstorming techniques, each group’s expectations should be documented at a “financial statement” level: revenue generation, cost management, brand management, business continuity, and so on. It is important to start the conversation at this level to maintain line of sight to the real business imperatives. Next we turn to the requirements that these customer groups have for the support organization.

Customer Expectations Dialog

This dialog can be used to build a bridge between the customer group’s business expectations and an organization’s original programs. This can be a difficult task and the temptation is to go in the other direction, starting with the portfolio of the organization’s programs and justifying them by showing how they support the bigger objectives of your stakeholder groups. This bottom-up approach can often result in assumptions about the program’s impact that will not make sense to the stakeholders when presented. Common ground must be reached between stakeholders’ expectations and the organizational team’s program objectives.

Depending on the nature of the organization, customer expectations might include “supply skills to match future revenue growth demand”, “automate core tasks to reduce cost”, and “provide office locations with better adjacencies to customers for improved market coverage”. Once these objectives are brainstormed and consolidated to a working list, answering a few basic questions can provide a conceptual framework for performance metrics that will matter to stakeholders.

- How can we monitor and demonstrate that our programs, which are a priority to our customers, are achieving the desired scope?
- Are we achieving the customer-desired results on schedule?
- How do we know if these programs are being delivered cost effectively?

By now we have established a value chain showing the translation of customer business needs to expectations of our specific organization to performance metrics for those expectations. The quantified version of this chain creates an algebraic “**transfer function**” where the independent variables of our program performance drive the dependant variables of customer-required results. (In fact, we sometimes refer to this chain as “business algebra”.) This transfer function when first presented to customer groups in the form of a dashboard or scorecard must pass the customer’s test in assessing whether your cause and effect assumptions are logical. Data in this transfer function become the basis for future statistical analysis, process modeling, and graphical display, not to mention group and individual performance goal setting.

Without further development, the business transparency created by a good measurement system can align and influence behavior to focus on what customers expect. However, in people-based businesses, metrics alone are not sufficient. Management must work with their leadership teams and staff members to be explicit about how to prioritize work; inventory the programs and projects that comprise the bulk of the organization’s work and match them to the metrics defined against customer’s expectations for your group. Taking this step can – especially the first time – help an organization to see where to make investment and disinvestment decisions. Work that contributes most directly to customer expectations is resourced. Work that makes less of a contribution is eliminated. We now have a list of programs and projects that are prioritized by contribution to our customer expectations (note that a single program with a strong contribution to one customer expectation might be prioritized at the same level as a program with a lesser contribution to multiple customer expectations).

Program Prioritization

Most teams addressing the prioritization exercise for the first time start by reviewing each *program* one at a time and determining which metrics (or customer expectations) are impacted by those programs. This results in a generous assessment of program contribution because rarely will a program escape discussion without being assigned a contribution to a metric.

But these connections are less important from the customer’s point of view. The proper technique is to review each *metric* one at a time and ask, “If this metric is changing, which programs will our customer assume are performing better or worse”.

For example, if a customer expectation is to “grow profit in the Chicago region next year” and the support organization does no further translation of relevance to their group’s activities, the group might then decide that their performance should be measured by “dollars of profit change year over year”. If the support organization is the company’s real estate function and one of its programs is to “maintain the Chicago sales office”, the temptation would be to assign this program a contribution to the “dollars of profit” metric. Sales management, however, when seeing an increase in profit figures in the Chicago region, will not likely credit the result to maintenance of the office. Some interim analysis was missing in defining the transfer function.

We would drive the team to break down the customer expectation to the next level of detail. If profit is a function of revenue and cost, then the customer expectation for the real estate group could be to “reduce cost in the Chicago region next year” measured by “dollars of cost reduction year over year” or even “dollars of cost reduction per square foot year over year”. Since the program to “maintain the Chicago sales office” goes directly to cost per square foot, there should be no doubt that an increase in “dollars of cost reduction per square foot” should be credited to the performance of this program.

During the dialog to prioritize programs, it may be necessary to refine the list of metrics.

Process Mapping

Once a prioritized list of programs is defined, the focus shifts to the underlying processes. We are looking to instrument our business on the basis of a process framework. This approach allows for better continuous measurement and improvement compared to considering a business to be a portfolio of programs. Programs and projects should be positioned as process or business improvement efforts – they come and go. The customer-facing metrics and underlying processes create a stable, long-term management model.

In the above example, “maintain the Chicago sales office” is more accurately described as a program

comprising one or more processes scaled to a specific geography (Chicago) or asset type (sales office). For the sake of simplicity, let us assume that only one process is involved: facilities management. Although the facilities management process comprises a variety of subprocesses (e.g., janitorial service, heating and air conditioning, and cafeteria services), we consider this to be one process for the purposes of our scorecard. Cost per square foot is a consideration for process efficiency (less cost equals a more efficient process). But extreme cost elimination would also come with service elimination. Some measure of service quality should be assigned; in this example, we will assume that a defect-minimization metric is applied. In summary, we have at least two metrics that indicate the performance of the facilities management process: cost per square foot (\$) and service defects (#). Monitoring both metrics is required to determine if this process is meeting customer expectations.

Metric Quality

In continuing to develop the scorecard, it is important from a business standpoint to consider both the cost of data collection and the quality of component metrics in the overall deployment of a measurement system. Passive, automated systems that collect data on an entire population as work is done are easiest. Active systems like employee surveys can fatigue the audience. Managers should set expectations about how much time in a given week employees should be spending on manual reporting. This decision should balance the consequences for unforeseen process failure with the desire for complete information. Most groups overcollect data, which is costly. Sampling techniques can be used to define the data collection process. Any more than 15 min per week of data reporting for the average employee will likely be difficult to maintain.

For analysis and decision-making purposes, consider the quality of the metrics in which the team is investing. Stoplight metrics (a status of “red”, “yellow”, or “green”) provide focus on major problems at the time of review. They are common because they are easy to produce and to understand. But by themselves they provide no information about trends over time and have limited statistical analysis capability. Leading indicators allow managers and teams to take improvement action before customers experience dissatisfaction in process performance. For

example, a rolling estimate of the net present value of facilities maintenance cost, taking into account all known factors that will impact future cost, might be a more useful planning metric than simply knowing the previous month’s charge. Forward-looking metrics always have an element of uncertainty and this can require a mindset change among their users; but the value of looking forward is worth the effort. Once a business is reframed as a collection of processes, management can take advantage of some basic facts of process models. For example, every process has inputs and activity to change those inputs into an output valued by the customer. Input, activities (processes), and outputs can each be measured separately, with input and process metrics being leading indicators for the outputs that go to customers.

As statisticians, we find the use of continuous data more flexible for performance analysis. To illustrate, consider the need of a process to deliver an output to a customer on a certain day and that this process is repeated dozens or hundreds of times each month. The customer complains that the delivery is “always late”. A business might be tempted to quantify how often they are late (a laudable goal), so a “% late” metric is created, showing, on a monthly basis, what percentage of them has been late out of all the deliveries. Month after month the management can only say “Go faster!”; meanwhile, the delivery staff work to improve and feel they are making progress. But the metric shows barely a change. We would suggest that a better metric might be the “# days late” where each month the total number of days by which the delivery is late is summed and reported. When recalculated, the process showed vast improvement by driving down number of days late across all projects. Incremental progress and its causes are rewarded and investigated, leading to real improvements that are eventually felt by the customer.

Metric Management

At this point, we will likely have a dozen or so collectible, high-quality metrics for intended display on our dashboard. To put the dashboard into practice as a tool, it must be populated with data. Each metric needs a detailed set of instructions for population and management. This set of instructions should be owned by each metric owner and include the following:

- *Operational definition* of the metric. How it is to be calculated, and what system or report is the source for the data (e.g., if the data comes from a standard spreadsheet report, the operational definition must include the name of the spreadsheet and the exact cell reference for the number representing process performance).
- Clear definition for the *unit of measurement*, e.g., numbers, dollars, days, percent and so on.
- *Specification limits* allowing for calculation of a good or bad result, e.g., no more than \$35/ft² for facility maintenance cost. These limits are set with the customers' input and final approval. There are two criteria to set: the performance expectation of an individual event and the expectation for all events. For example, 90% of facilities should achieve the target square foot cost of \$35 for aggregate performance to be considered "green".
- *Frequency of data reporting*. How often will the data reasonably change based on sensitivity to process performance? We find that a monthly interval is about right for the management to review performance dashboards (unless the dashboard can be automated for real-time data on high-volume transaction processes). Thus any metric that is not meaningful on a monthly basis should be eliminated. Perhaps what is more important for performance management is that if the data are to be used to drive corrective

action then the nonmanagement review must be more frequent.

- *Performance owners*. Who will be tasked with improvement actions if the data show performance trending negatively? (This is usually a person who has a span of control over the resources doing the work and not usually the same person collecting the data.)

Conclusions

Team involvement through the dashboard development process is critical for the understanding of each individual's contribution to meeting the customer expectations. This development should not be done in a backroom by the quality manager. Now that the dashboard is live it should be used to inform customers of ongoing performance, to incent the right behavior among staff, and to identify improvement opportunities in the business. The dashboard itself becomes a focal point for management meetings where context is discussed and results are put in the spotlight.

Management's use of the dashboard in regular reviews ensures that what is getting measured is not only getting done but is adding value to its customers' operations.

MICHAEL E. JORDAN AND THOMAS MCCARTY

Theory of Constraints

Introduction

According to theory of constraints (TOC), every organization has at least one constraint (still not many – very few constraints) which prevents the system from achieving a higher performance relative to its goal (similar to Liebig’s law of the minimum – the growth of agricultural systems is controlled not by the total of resources available, but by the scarcest resource). The constraints can be broadly classified as “internal constraints” and “external constraints”. “Internal constraints” can be a physical one like a “resource constraint” or a nonphysical like a “policy constraint”. An “external constraint” can again be a physical one like “raw-material constraint” or a nonphysical one like a “market constraint”. In order to manage the performance of the system, these constraints must be identified and treated accordingly.

Background Information

The TOC was developed by Eliyahu M. Goldratt and other loosely coupled communities of practitioners around the world. Goldratt started to spread his initial ideas and thoughts in the late 1970s by first developing solutions to scheduling problems of manufacturing plants. These solutions were extended to an overall management philosophy during the years 1985–1990 immediately after the publication, in 1985, of the first edition of his best-seller book, *The Goal*. Since then additional applications of TOC were developed, and are still being developed, which enables this methodology to be applied to additional areas (business as well as others like education, health, and more). Currently, there are applications of TOC that cover most of the managerial applications. TOC is sometimes referred to as *constraint management*.

The TOCICO (Theory of Constraints International Certification Organization) is a global not-for-profit certification organization for TOC practitioners, consultants, and academics to develop and provide certification standards, and to facilitate exchange of latest developments.

A Process of Ongoing Improvement: The Five Focusing Steps

TOC is based on the premise that the rate of revenue generation is limited by at least one constraining process (i.e., a bottleneck in the case that the constraint is physical). Only by increasing throughput (flow) at the bottleneck process can overall organizational throughput be increased.

The key steps in implementing an effective TOC approach are as follows:

- Step 0: articulate the goal of the organization. Frequently, this is something like, “Make money now and in the future.”
- Step 1: *Identify* the constraint(s) – the thing that prevents the organization from obtaining more of the goal.
- Step 2: Decide how to *exploit* the constraint – e.g., make sure the constraint is doing things that the constraint uniquely does, and not doing things that it should not do.
- Step 3: *Subordinate* all other processes to the above exploitation decision – align all other processes/resources/departments to the decision made above.
- Step 4: *Elevate* the constraint – as long as the constraint is not a strategic constraint, permanently increase capacity of the constraint.
- Step 5: If, as a result of these steps, the constraint has moved, return to Step 1. Do not let *inertia* become the system’s constraint. The inertia is the tendency to continually do things the old way.

The TOC Thinking Tools

The thinking processes are a set of tools to help managers walk through the steps of initiating and implementing the TOC. When used in a logical flow, the thinking tools also help walk through the buy-in process:

1. Gain agreement on the problem.
2. Gain agreement on the direction for a solution.
3. Gain agreement that the solution actually solves the problem.
4. Agree to overcome any potential negative ramifications of the solution.
5. Agree to overcome any obstacles to implement the solution.

2 Theory of Constraints

TOC practitioners sometimes refer to these in the negative as *working through layers of resistance to a change*.

The thinking processes, as codified by Goldratt and others, are as follows:

1. Current reality tree (CRT) – evaluates the network of cause–effect relations between the undesirable effects (UDEs) in the defined system and helps to pinpoint the root cause(s) for most of them.
2. Evaporating cloud (conflict resolution diagram or CRD) – defines the conflicts that usually perpetuate the causes for an undesirable situation and helps to develop a win–win solution.
3. Core conflict cloud (CCC) – a combination of conflict clouds based on several UDEs. This is a deeper conflict that creates the undesirable effects.
4. Future reality tree (FRT) – once some actions (injections into reality) are chosen (not necessarily detailed) to solve the root cause(s) uncovered in the CRT and to resolve the conflict(s) in the CRD, the FRT shows the future states of the system as a result of implementing the changes and helps to identify possible negative outcomes of the changes (negative branches) and to prune them before implementing the changes.
5. Negative branch reservations (NBR) – identify potential negative ramifications of any action (such as an injection or a half-baked idea). The objective of the NBR is to understand the causal path between the actions and negative ramifications so that the negative effects can be “trimmed”.
6. Positive reinforcement loop (PRL) – desired effect (DE) presented in FRT amplifies intermediate objective (IO) that is earlier (lower) in the tree. Although IO is strengthened it positively affects this DE. Finding out PRLs makes FRT more sustaining.
7. Prerequisite tree (PRT) – states all of the IOs necessary to carry out an action chosen and the obstacles that will be overcome in the process.
8. Transition tree (TT) – describes in detail the action that will lead to the fulfillment of a plan to implement changes (outlined on a PRT or not).
9. Strategy and Tactics (S&T) – the overall project plan and metrics that will lead to a successful implementation and the ongoing loop through

(POOGI) (a process of ongoing improvement – the five focusing steps).

Throughput Accounting

Throughput accounting refers to a specific accounting methodology linked to the TOC. Throughput accounting suggests that one examine the impact of investments and operational changes in terms of the impact on the throughput of the business. It is an alternative accounting method to cost accounting and is very effective while making internal decisions.

Specific TOC Application Solutions

Operations

Within manufacturing operations and operations management, the solution seeks to pull materials through the system, rather than push them into the system.

Drum–buffer–rope (DBR) is a manufacturing execution methodology, named for its three components. The *drum* is either one of the two: (a) the schedule to run the physical constraint of the plant – the work center or machine or operation that limits the ability of the entire system to produce more or (b) the shipment schedule. The rest of the plant follows the beat of the drum. They make sure the drum has work and that anything the drum has processed does not get wasted.

The *buffer(s)* protects the drums, so that it always has the right work flowing to it. Buffers in DBR have time as their unit of measure, rather than quantity of material or capacity. This makes the priority system operate strictly on the basis of the time at which an order is expected to be at the buffered operation. Traditionally DBR usually calls for buffers at several points in the system: prior to the constraint, prior to the synchronization points, and at shipping. The method is known as *S-DBR*, which is a simplified form of the DBR assuming that the only real constraint of the system is the market. Consequently, it requires only a single buffer at shipping.

The *rope* is the work-release mechanism for the plant. Only a “buffer time” before an order is due (in an S-DBR environment) does it get released into the plant. Pulling work into the system earlier than a buffer time guarantees high work-in-process, reduces flexibility, and slows down the entire system.

Supply Chain/Logistics

The solution for supply chain is to move to a replenishment model, rather than a forecast model. This is based on fast replenishment capability. The two major elements in this solution are the following:

1. TOC – Distribution: enable high flexibility while using a color-based alerting system that indicates low probability to deliver on time.
2. TOC – (vendor managed inventory) VMI: Continuous and fast replenishments based on daily consumption.

Finance and Accounting

Use throughput accounting with T (throughput which is defined as revenues minus true variable costs), OE (operational expenses), and I (inventory) as the most important measures to support managers in the decision-making process.

Project Management

Critical-chain project management (CCPM) method is the derivative of TOC to project management. It is based on the realization that most of projects' tasks contain internal unseen buffers. Therefore, the "critical-chain" method suggests that these buffers be extracted from the tasks and to create well-managed buffers to the project itself. These buffers are to protect mainly (a) the completion time of the project from deviation in the critical-chain tasks (tasks with no time slack to the past and to the future) and (b) to protect the critical chain from the deviations in the noncritical tasks.

In addition, in an environment where there are more than a few projects sharing the same resources, the problem turns into one of managing the time of the most loaded resource(s), the one which is scheduled to participate in many projects. This schedule is also the basis for the synchronization between the projects and it provides the mechanism to avoid penetration of lateness from one project to other projects via the shared resource(s).

Marketing and Sales

Although originally focused on manufacturing and logistics, TOC has expanded lately into sales management. First data shows that the sales system is massively constrained and TOC offers significant opportunity to increase enterprise throughput.

The solution is based on the ability to turn the operation's superb performance (from the implementation of the above concepts in the manufacturing and project management environments) into value for the clients and, as a result, to be able to charge clients with premium prices. This will eventually increase the organizational throughput.

Buffer Management

Wherever buffers are applied, in manufacturing, in project management, in the supply chain, and so on, buffer management plays an important rule in getting significant results from the implementation of TOC. Buffer management contains three elements: (a) a well-defined color-based system to alert when probability to be on time is deteriorating; (b) dynamic modification of buffer size according to fluctuations in performance and demand; and (c) statistic analysis of reasons to penetrate into the "red" zone of the buffer providing effective information to the POOGI process.

Further Reading

Novels

- Goldratt, E.M. *It's Not Luck*, ISBN 0-88427-115-3.
- Goldratt, E.M. *Critical Chain*, ISBN 0-88427-153-6.
- Goldratt, E.M. & Cox, J. *The Goal: A Process of Ongoing Improvement*, ISBN 0-88427-061-0.
- Goldratt, E.M., Schragenheim, E. & Ptak, C.A. *Necessary but not Sufficient*, ISBN 0-88427-170-6.

Theory of Constraints

- Goldratt, E.M. *Essays on the Theory of Constraints*, ISBN 0-88427-159-5.
- Goldratt, E.M. *What is this thing called Theory of Constraints and how should it be Implemented?* ISBN 0-88427-166-8.
- Goldratt, E.M. *Beyond the Goal: Eliyahu Goldratt Speaks on the Theory of Constraints*, ISBN 1-59659-023-8.
- Goldratt, E.M. *Dr Lisa Lang. Achieving a Viable Vision: The Theory of Constraints Strategic Approach to Rapid Sustainable Growth*, ISBN 0-9777604-1-3.

4 Theory of Constraints

Manufacturing

- Goldratt, E.M. *The Haystack Syndrome – Sifting Information Out of the Data Ocean*, ISBN 0-88427-089-0.
- Goldratt, E.M. *Production the TOC Way*, ISBN 0-88427-175-7.
- Goldratt, E.M. & Fox, R.E. *The Race*, ISBN 0-88427-062-9.
- Woepfel, M.J. *The Manufacturer's Guide to Implementing the Theory of Constraints*, ISBN 1-57444-268-6.

Supply Chain

- Schragenheim, E. & Dettmer, H.W. *Manufacturing at Warp Speed*, ISBN 1-57444-293-7.

Strategy

- Dettmer, H.W. *Strategic Navigation – A Systems Approach to Business Strategy*, ISBN 0-87389-603-3.
- Lang, L. *Achieving a Viable Vision: The Theory of Constraints Strategic Approach to Rapid Sustainable Growth*, ISBN 0-9777604-1-3.
- Spector, R.E. (2006). *How Constraints Management Enhances Lean and Six Sigma Supply Chain Management Review*, Jan/Feb.

Accounting and Finance

- Bell, J., Swain, M., Bell, J. & Ansari, S. *The Theory of Constraints and Throughput Accounting*, ISBN 0-07-027589-0.
- Corbett, T. *Throughput Accounting*, ISBN 0-88427-158-7.
- Lang, L. *Maximizing Profitability: The Theory of Constraints Approach to Maximizing Profits*, ISBN 0-9777604-0-5.
- Noreen, E.W., Smith, D.A. & MacKey, J.T. *Theory of Constraints and Its Implications for Management Accounting*, ISBN 0-88427-116-1.

Project Management

- Hutchin, T. *Enterprise-Focused Management – Changing the Face of Project Management*, ISBN: 0-7277-2979-9.

- Leach, L.P. *Critical Chain Project Management*, ISBN 1-58053-903-3.
- Newbold, R.C. *Project Management in the Fast Lane: Applying the Theory of Constraints*, ISBN 1-57444-195-7.

Continuous Improvement and the Thinking Processes

- Cox, J.F. II. & Spencer, M.S. *The Constraints Management Handbook*, ISBN 1-57444-060-8.
- Dettmer, H.W. *Goldratt's Theory of Constraints – A Systems Approach to Continuous Improvement*, ISBN 0-87389-370-0.
- Dettmer, H.W. *Breaking the Constraints to World-Class Performance*, ISBN 0-87389-437-5.
- Scheinkopf, L.J. *Thinking For a Change: Putting the TOC Thinking Processes to Use*, ISBN 1-57444-101-9.
- Schragenheim, E. *Management Dilemmas: The Theory of Constraints Approach to Problem Identification and Solutions*, ISBN 1-57444-222-8.

Sales and Marketing

- Klapholz, R. & Klarman A. *The Cash Machine: Using Theory of Constraints for Sales Management*, ISBN 0-88427-177-3.
- Woehr, W.A. & Legat, D. *Unblock the Power of your Salesforce*, ISBN 3-7083-0082-3.

Healthcare

- Wright, J. & King, R. *We All Fall Down: Goldratt's Theory of Constraints for Healthcare System*, ISBN 0-88427-181-1.

Education

- Ean, K.C. *Thinking Smart: Applying the Theory of Constraints in Development Thinking Skills*, ISBN 967-978-918-7.

AVRAHAM MORDOCH

Target Costing

Introduction

Target costing is a market-driven financial strategy for product and service design. Its goal is to assure profitability in the design and introduction of new products and services. Target costing was first conceived in the 1960s in Japan. It expanded the concept of value engineering to a customer-focused, profit-driven cost system that focuses on price. Target costing has been successfully adopted by most Japanese manufacturing firms including, among others, Sony and Toyota. More recently, it has begun to gain acceptance among US companies, as well.

Traditional “cost plus” reasoning sets the price of a product or service by starting with anticipated costs, and adding the desired profit margin. In competitive markets, the “cost plus” approach often leads to producers setting an “asking price” that customers are unwilling to pay. The result is a failure of the new product or service to achieve market acceptance, with an accompanying substandard return on the producer’s investment, or, in the worst case, a partial or total loss of the capital invested in product development.

Target costing, on the other hand, starts with the price that market conditions warrant, and then factors in required profitability to yield a “target cost” that must be met in the design and delivery of a product or service. To assure sound financial decision-making, target costing can and should be applied from the very start of the concept phase. As design and development progress, updated information permits continuous reassessment of whether costs are on track to provide a feasible return on investment (ROI). Through this process, target costing can be integrated with – and provide critical information to – corporate investment decision tollgates, which exist to provide go/no go decisions on funding of the development project. Target costing thereby provides a realistic and rational discipline that places the focus squarely on profitability throughout. Moreover, it provides a financial management discipline that is consistent with, and a valuable addition to, other related best practices such as **Six Sigma** and Design for Six Sigma.

Methodology

The concept of target costing is simple; however, achieving the benefits of the approach requires business commitment and discipline throughout the organization. The following example illustrates the requirements and challenges that companies face along the path to success.

Consider the definition of target cost:

$$\text{Target cost} = \text{Market price} - \text{Profit} \quad (1)$$

Then the current shortfall relative to target costing requirements is

$$\text{Cost gap} = \text{Target cost} - \text{Current expected cost} \quad (2)$$

The cost gap can be further detailed by attribution to specific design components or activities.

Consider a company that determines, through competitive analysis and customer feedback, that the market-driven price that a product will sell at is \$100. Leadership decides that for a sustainable ROI, the contribution margin must be 30%, and hence the target cost is \$70.

On the surface, this sounds simple. However, implementation of the target costing plan requires considerable thought and planning. At the outset, successful target costing requires an accurate and comprehensive assessment of market price. While gathering competitive analysis is straightforward in theory, and often will be even in practice for a commodity, for most products this exercise can pose quite a challenge. Often the sale of one product is bundled with others, or there will be warranties and services attached, which impose additional costs on the producer. So the price of an individual product may not be obvious. The target costing team will be required to work closely with marketing, sales, and service functions, and to look as well for confirmation of pricing projections from industry trade data. It should be evident that considerable expertise is required to confidently project the \$100 starting point for our hypothetical target costing analysis.

Similarly, a cross-functional effort will be required for identifying the impact on profitability of the many aspects of the company’s performance, which will in turn impact the required \$70 target cost. These include not only variable costs of production, but also manufacturing yield losses, scheduling and utilization

2 Target Costing

efficiencies, distribution, warranty and service costs, and many others. For realistically projecting these costs, expertise in finance, operations and supply-chain management, product life and warranty cost estimation, procurement, supplier quality, and statistical quality control are required.

Cost management entails trading off wins in cost reduction in some of the above areas against challenges in others to assure that the “cost gap” is closed. Effective use of target costing promotes resource reallocation, as needed, to assure success. In instances where best practices are being applied for the first time, “low-hanging fruit” will often account for a significant component of cost reduction; but it will seldom close the gap by itself. The engineering, design, service, and finance functions – and perhaps others – will be asked to turn the *status quo* on its head to achieve the remaining gains.

In addition, while traditional methods typically ignore trends of key financial characteristics over the life of the product – including, for example, price erosion, productivity increases, share gain/loss, and so on – a long-term focus on target costing can estimate these items from the outset and update them on a continuing basis. This activity can continue even after product introduction, to update profitability projections or run “what if” scenarios to determine the potential impact of market developments.

Application Issues

Target costing offers a wide array of advantages and best practices for promoting profitable introduction of new products and services. It promotes cross-functional awareness and cooperation, supports an expansive long-term view of product/service profitability, and can be implemented in a manner consistent with a variety of corporate financial metrics (e.g., ROI, internal rate of return (IRR), and others).

At the same time, as evidenced by the foregoing example, a company’s target costing success requires strong management coordination. In addition, a comprehensive long-term approach to target costing may require estimates of factors that are difficult to predict – for example, price erosion, share gain/loss, and demographic shifts. However, target costing tech-

niques compensate for this uncertainty by offering management the opportunity to specify alternative scenarios, and to run “what if” analyses across an array of potential market developments.

Software for implementing target costing requires substantial financial modeling skill, and represents a nontrivial development effort. Partly for that reason, most companies which have adopted target costing have implemented a form which is built around the more commonly estimated aspects of new product development, such as variable costs, manufacturing quality, and perhaps warranty/service costs. Incorporating additional factors such as R&D investment, time to market, and sales and distribution costs entails significant additional complexity.

Conclusion

Given financial expertise, cross-functional cooperation, and effective management support, target costing can provide an unprecedented means for assuring new product/service profitability. Target costing provides extraordinary information quality in support of capital expenditure planning, investment decision tollgates, and related corporate financial practices. By starting from market realities and explicitly accounting for business needs, target costing makes business success a design requirement, and helps avoid costly financial disasters in new product or service introduction.

Further Reading

- Berliner, C. & Brimson, J.A. (eds) (1988). *Cost Management for Today's Advanced Manufacturing*, Harvard Business School Press, Boston.
- Cooper, R. & Kaplan, R.S. (1998). *Design of Cost Management Systems*, 2nd Edition, Pearson Education POD.

Related Article

Sampling Inspection of Products; Statistical Quality Control; Six Sigma; Warranty Claims and Costs: Statistical Analysis of; Warranty Modeling.

DONNA M. FALTIN

Lean Accounting

Introduction

The Association for Operations Management, also known as *APICS*, defines a **lean** enterprise as one with a continuous focus on (a) waste elimination and (b) the customers' needs in all parts of its operations, manufacturing, and administration. Emphasis is given to lean structures and processes, flexibility of response, and methods and techniques to continually seize new opportunities as they arise.

Achieving a lean enterprise requires a culture built on a foundation of lean thinking, which is a continuous process of eliminating wasteful activities while increasing the value of a process or function. By creating visibility into operations and processes – no matter what industry or function – management gains the ability to analyze the elements of a process, thus allowing the enterprise to eliminate unnecessary activities and create more time for additional value-added activities.

Lean accounting focuses on accounting for activities in such a way that it allows for an understanding and measuring of the value created for the customers, and uses this information to enhance customer relationships, product design, product pricing, and continuous lean improvement. A successful lean accounting implementation enables a company to take advantage of benefits such as providing accurate, timely, and understandable information to motivate the lean transformation throughout the organization, and for decision-making leading to increased customer value, growth, profitability, and cash flow. It also facilitates the use of lean tools to eliminate waste from the accounting processes while maintaining sound financial control. It can support a lean culture by motivating investment in people, providing information that is relevant and actionable, and empowering continuous improvement at every level of the organization.

The DMAIC Framework

An example of a DMAIC framework can be found in Figure 1.

Every process within an organization has inputs and outputs, suppliers and customers, and beginning and ending boundaries. Under DMAIC, we start

with defining the process we want to analyze, measure, improve, and control. Developing a SIPOC (supplier, inputs, process, outputs, and customer) diagram provides a high-level understanding of process stakeholders, contributors, and customers. A critical component to the success of the initiative is capturing the voice of the customer (VOC). Understanding the customers' needs and continually soliciting their feedback will drive the analysis and improvement components. Most businesses that we know exist for one sole reason: to make money. In order to maximize their profitability, businesses must assure that they are always working on delivering a product and/or service that continually meets their respective marketplace requirements. One key element to understanding a businesses marketplace is to capture the voice of the external customer. There are many tools to do so, with the most common method being to conduct a **Kano** analysis. By stratifying collected customer data, the business may get to know its customers better. By capturing the VOC, the business may identify current functional business to customer as well as learn where their customer's strategic direction is heading. This is important when considering a new product or service design. The Kano analysis chart in Figure 2 depicts the relationship between product performance and customer satisfaction.

A cross-functional **process map** for the “as is” state is then developed to provide visibility into the components and their respective relationships within the process to be analyzed and improved.

Knowing the VOC allows the business to understand what is important. This dictates what needs to be measured from the customer's viewpoint; after all, the customer is the final judge of product and service.

We then define the measurements by which to monitor our progress as we make improvements and implement controls. During this phase, we will develop a baseline **value stream map** (VSM), and collect some baseline data as to the performance of the process. Some noteworthy process performance indicators include but are not limited to the following:

- *Process cycle efficiency (PCE)*, which is calculated by value-added time (VAT) divided by the total elapsed time (TET) of the process.

$$PCE = VAT / TET \quad (1)$$

2 Lean Accounting

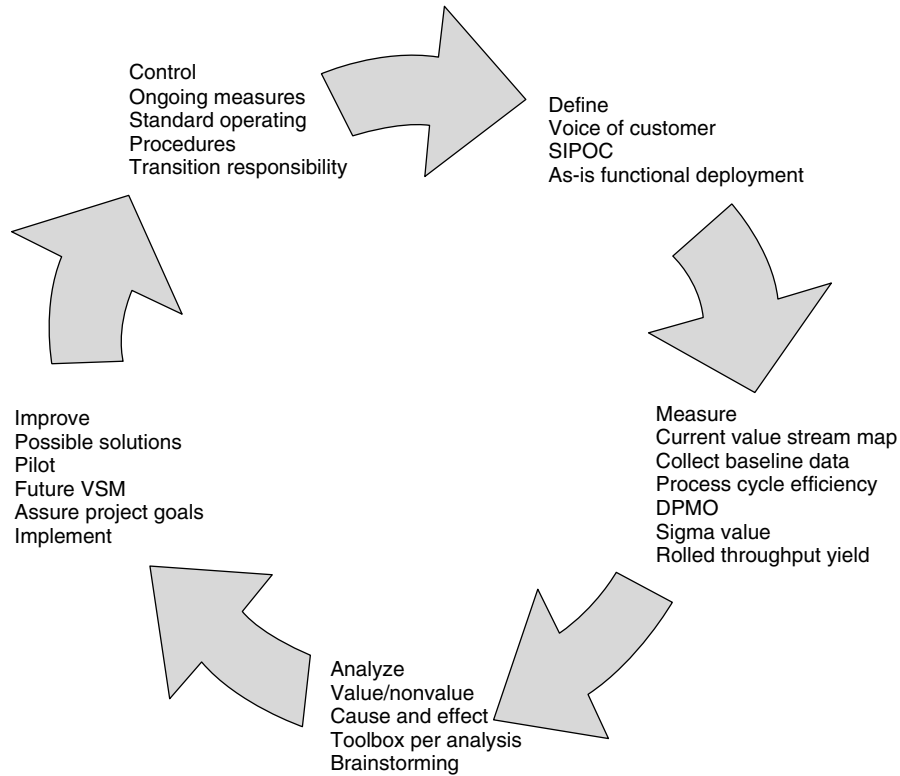


Figure 1 DMAIC framework for lean accounting

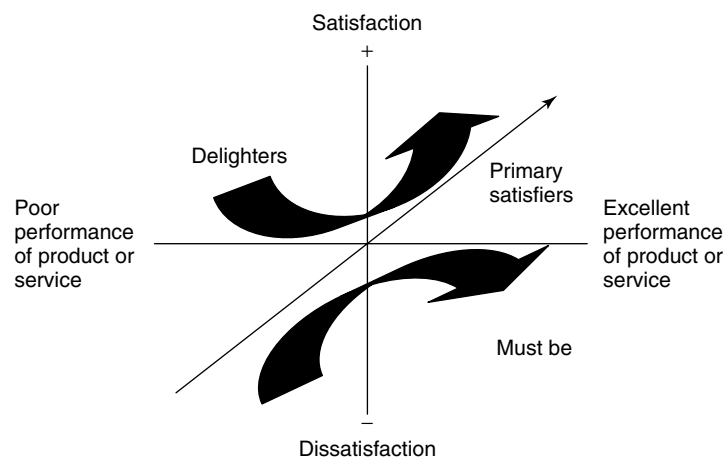


Figure 2 Classifying customer needs by Kano analysis

Processes that are benchmarked as being lean are processes whose $PCE \geq 25\%$. The typical Pareto principle applies where 20% of the investigated

process steps are responsible for 80% of the process bottlenecks. This finding typically serves both service and manufacturing environment processes.

- *Defect per million opportunities (DPMO)* displays the average number of defects per unit of measure observed during average process throughput

$$DPMO = M \times \frac{D}{N \times O} \quad (2)$$

D = total number of defects counted in the sample

N = number of units of product or service

O = number of opportunities per unit measured of process product or service to have the defect occur

M = one million (1 000 000)

- *Rolled throughput yield (RTY)* is the likelihood that a single unit can pass through the series of required process steps free of defects. This measure considers “rework” as often witnessed in accounting reconciliation. To calculate RTY, one must consider the defect rate at each subprocess step. If we consider a process to have four subprocess steps, P1, P2, P3, and P4 with each delivering 80% usable output per process step criteria, the RTY would equate to 0.4096

$$\begin{aligned} RTY &= P1 \times P2 \times P3 \times P4 \\ &= 0.8 \times 0.8 \times 0.8 \times 0.8 \\ &= 41\% \end{aligned} \quad (3)$$

Consequently, the defect rate for each subprocess is 20% and the defect rate on the final product or service from this process is 59%.

Essentially, a value stream encompasses all the activities performed to create value for a customer that can be associated with that product or service. Remember, value is defined as a product or service that a customer is willing to pay for, and must be balanced with what it takes a company to deliver that product or service to the customer’s satisfaction. A company cannot just decide that it believes engineering is a nonvalue adding function. Without the engineering expertise of the company, the product or service could not be developed to meet customer specifications. Some of the activities within a value stream might include marketing, selling, estimating, pricing, designing, purchasing, engineering, production, warehousing, delivery, customer service, and

collection. A value stream blends costing principles with value principles, since both are necessary to derive a true picture of customer profitability.

With this VSM and baseline data, we can then begin to analyze ways to eliminate waste and increase value within the process. Some examples of types of lean waste include overproduction, motion, inventory, transportation, waiting, underutilized people, defects, and overprocessing. The analysis of the process will involve techniques such as **cause and effect**, and the identification of activities that are value adding, nonvalue adding, unnecessary, or are considered to be rework that should be avoided.

Once this visibility into the process has been created, we can begin to identify ways to improve the process, which includes methods for gaining buy-in from the constituencies. Unless the affected parties are on board, a successful implementation is at risk.

Finally, we want to be sure that our improvements are lasting and that we have created an environment of continuous improvement. Most changes take place at a point in time and reflect marginal improvements. A continuous improvement program (Figure 3) ensures that processes will perform at peak efficiency into the future. Today, businesses draw benefits from continuous improvement programs by endorsing the resource development in both soft project skills (people and general **project management**) and hard skills (analytical and quantitative). There is usually a good blend of both types of skills within an organization, but it is critical for any continuous improvement effort to address optimization of team talent as well as to foster a culture for skill development.

Utilizing an overview of the elements of lean accounting as provided by the American Institute of Certified Public Accountants (AICPA), we can classify lean accounting components into four levels to demonstrate the progression of the use of financial and accounting data within an organization, from basic transaction recording to use of the information in strategy development. The levels are defined as follows:

- Level 1: lean financial accounting, which includes items such as procurement and payment, order fulfillment and receivables, and month-end closings.
- Level 2: lean operational accounting, which includes items such as materials cost, labor and

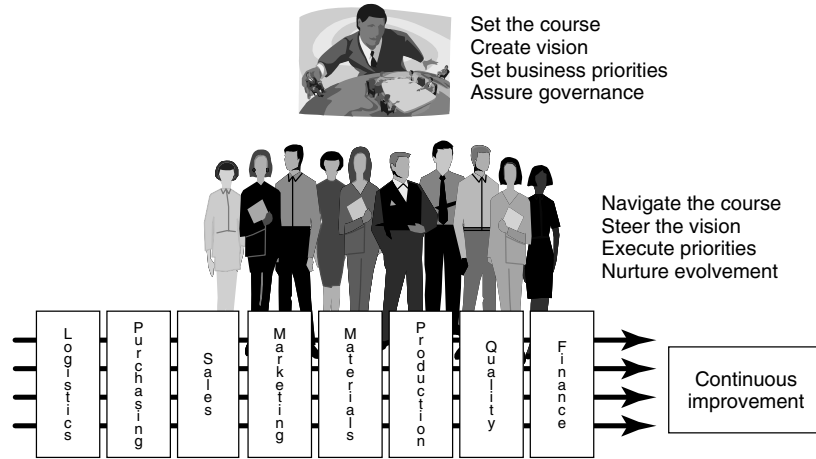


Figure 3 Hierarchy for Lean Accounting's Continuous Improvement

- overhead cost, direct cost, and inventory tracking.
- Level 3: lean business management which includes items such as organization of the value stream, customer value and **target costing**, and rewards and recognition.
- Level 4: lean management accounting, which includes items such as alignment of organizational strategy with lean goals, budgeting and planning, managing service line or product profitability, and performance management.

Example Applying the DMAIC Framework

Let us look at an example of how a multinational communications company implemented lean accounting within its organization utilizing the DMAIC framework. The company began the lean implementation initiative by defining the Level 1 financial accounting processes to be analyzed, where the accounting transactions were processed and recorded. Those processes included billing, revenue recognition, cash management, accounts payable, fixed asset accounting, and subsidiary accounting. We will use the billing process in this example.

Once SIPOC and functional deployment maps were created for the billing process, management had a baseline to begin analysis. One of the issues that became apparent was the amount of time and resources involved in pricing products and preparing

invoices. The first step was to identify wasteful activities in the process. The measurements they used included task time for the process, percentage of rework, number of price quotes, time to prepare a price quote, number of invoices processed, time to process an invoice, PCE, and number of customer inquiries on invoices. The process included some significant rework and unnecessary activities because the people performing the billing function thought it was their job to fix other people's errors. Management's challenge was to identify ways to prevent those errors upstream of the data.

Once nonvalue activities were identified and eliminated and a future VSM was created, the company realized significant benefits. One improvement was to change its billing methods from the traditional bundled service to a more customer-oriented breakdown of value components. This helped customer and product profitability to be measured more accurately. The timeline for pricing and billing activities decreased, and management had greater visibility into the value stream. Customer service was better equipped to anticipate and address customer value issues. They had laid the foundation for a continuous improvement program.

Discussion and Summary

There is a saying that we frequently hear from successful business leaders: you cannot manage what you cannot see or control. Implementing lean

accounting using the DMAIC framework provides both the intelligence and visibility necessary to see how the company is servicing the customer, managing operations, and meeting expectations. Once management has a view as to what is value adding and what is unnecessary, it can refocus efforts toward adding value to the customer experience, maintaining a lean operation, and improving profitability. The key to a successful lean accounting

implementation utilizing the DMAIC framework is commitment from management, clear direction, and involvement of the constituents. Once the team is on board and the culture is established, the company should have the fuel to maintain a continuous improvement program, thereby improving its profitability and value.

FRANK GRAYESKI AND LEN STAVISH

Organizational Assessment Models

The Quality Assurance Models

The seeds of organizational assessment models can be traced back to the period between the 1930s and 50s, when leading companies operating in strategic sectors started to elaborate the quality assurance (QA) concept. The concept was based on the definition of system requirements (QA standards) and on independent assessments (audits) aimed at checking compliance of the involved organization with such requirements. The purpose of QA was to give executives confidence that the audited organization was capable to meet the quality targets and all the relevant activities (planning, execution, measurement, review) had been put in place. The first structured QA standards appeared in the 1950s/1960s, in the wake of the progresses made in the definition of the “product development cycle”.

QA extended rapidly to all sectors where defect prevention was a strategic issue (defense, aerospace, nuclear), and then to commercial sectors. Geographic diffusion was also rapid. But companies using QA soon realized that the next logical step was to extend the concept to their suppliers and to assess them for compliance (second-party audits). *External* QA standards were created. Proprietary at the beginning – being derived by the customer internal standards – external standards could not obviously enter into the same details as internal standards. External QA standards soon took the lead over internal standards. In B2B and B2A (Business to Business and Business to Administration) relations, earlier system assessments were seen as an important means to save time and money with respect to the traditional product testing procedures [1].

The need to reduce the burden, for both the suppliers and customers, coming from multiple standards and audits, fostered joint definition of sector standards (e.g., automotive). Increased awareness of the commonalities among production systems – whatever the product – led some supply-intensive administrations (typically defense) to create a wide scope of general standards. The next step was the creation of national QA standards, and finally, in 1987, of the ISO 9000 quality system standards [2]. Such

international standards marked the decisive turn from second- to third-party audits, conducted by independent, recognized bodies (certifications).

Total Quality Management Models and Awards

The 1980s showed the limits of a standard-based quality vision. The new customer- focused quality approach inaugurated by the Japanese was winning out over the inward- focused approach induced by standards. The Western “quality prophets” (Deming, Juran, Feigenbaum), unheard at home in the past, were rediscovered. New quality models – called *total quality management (TQM) models* – proliferated, rotating around the concepts of customer focus, process-based management, employee involvement, and continuous improvement (see **Total Quality Management (TQM)**). Proliferation dwindled sharply when some highly respected quality organizations thought of creating national and international quality awards, as a means to diffuse the new quality culture. To that purpose, experts were gathered to review and consolidate the findings of the period and define, through a consensus process, the relevant award assessment model and process. The American Malcolm Baldrige (MB) National Quality Award [3] was issued in 1987, followed, in 1991, by the European Foundation for Quality Management’s (EFQM) European Quality Award [4]. Figures 1 and 2 represent the block diagram of the two models.

The two models brought forth significant changes with respect to QA standards. First, they do not look for compliance with minimum requirements but define methods to measure performance and capabilities, based on a number of factors (the blocks in the figures) and subfactors that are considered critical in relation to excellence. The award assessment is based on the analysis of an application report presented by the candidate and a site visit. Scores are assigned according to defined measurement rules. Second change: results appear on the assessment scene – *perceived results* (or outcome) in particular. That means recognition that *compliance* with a model is not sufficient for a quality judgment; *effectiveness*, measured by results, completes the assessment scenario. Finally, the subdivision between *results* and *enablers*, introduced by the EFQM model, starts to shape a real organizational model, where management can identify – and test – the links between

2 Organizational Assessment Models

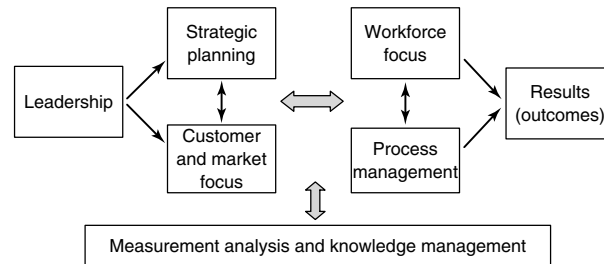


Figure 1 The Malcolm Baldrige model [Source: The National Institute for Standards and Technology – <http://www.quality.nist.gov>.]

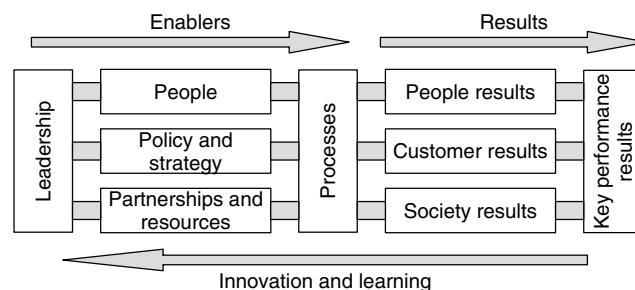


Figure 2 The EFQM excellence model, 2007 [Source: European Foundation for Quality Management – <http://www.efqm.org>. © EFQM. 1999–2003.]

organizational causes and their effects. That helps to understand the underlying dynamics of TQM models: continuous improvement.

Models for Continuous Improvement

The narrower the scope of a model, the better it works. Award models are conceived to measure, compare, and publicly recognize organizational excellence, and to that purpose they need rigorously defined standards. Is it correct to use the same standardized models and processes for performance improvement also? Arbitrarily enlarging the scope of a model is risky; better to reshape it. Unfortunately, the award models and processes have become, in the last 20 years, the normal choice for self-assessment. An acceptable choice when seeking for comparable quantitative measurements of performance; less acceptable when the aim is organizational diagnosis, to identify areas that need improvement (we use this term here both for incremental and quantum leap improvement). Experience shows that focus on scoring is often at odds with focus on diagnosis.

The dominance of award models in TQM – combined with the dominance of ISO 9000 in QA – created a kind of conformism that somehow hampered the development of free, improvement-oriented models. In fact, diagnostic effectiveness is enhanced by the use of *contingency* models (that is, models that fit the organization's specific characters). Standardized models and procedures, necessary for the awards, turn out to be a burden.

The Present Situation

In the area of QA and certification, ISO 9000 is the reference worldwide. In the area of TQM, the MB and EFQM are the most popular models in the West, the Deming Prize Model [5] in Japan. Recognition that the more the model fits the organizations' characters the better it performs, has led to the development of sector TQM models (e.g., education, public administration).

Several TQM/excellence models have been developed by companies, consultant organizations, and experts, often merging TQM concepts with other

successful approaches, such as the *balanced scorecards* [6]. They can be divided into two categories.

1. Models aimed at measuring organizational performance and capabilities

Similar in principle to the award models, they are however – usually – more inward than outward directed. The aim is to endow management with a set of high-level indicators to measure the state of the organization from a plurality of perspectives. Indicators are chosen from among those factors that are considered critical for excellence (critical success factors, CSF). The literature provides many examples of CSF [7–10]. Albeit experts agree on comprehensive lists of CSF, differences arise when a short list (typically, 7–10 elements) is pursued to get to a manageable model. CSFs are normally high-level, abstract constructs (e.g., leadership, organizational learning) that require careful declination into assessable elements. Separation of results and enablers helps, since it confines the more intangible evaluation to the enabler side. Inconsistencies between enablers and results turn out to be important symptoms for further investigation. For these models, tests and validation of the chosen CSFs, taking into account interdependence, is important, as well as measurement reliability. Structural equation modeling is often used for the purpose.

2. Models aimed at organizational diagnosis and improvement

These too imply performance and capability measurements. But the role of measurement is secondary, in relation to results. Measurement is mainly used to identify performance gaps (the starting point for diagnosis); in relation to enablers, it serves to quantify weaknesses and then define priorities of intervention. Possible errors are controllable, since the normally used heuristic cycle “plan/execute/test” evidences them. Diagnostic models tend to overcome the use of abstract constructs, like the critical success factors, as basic elements of the model. Given their purpose, the use of elements that are more directly linked to the reality of the organization is recommended. They can be called *critical systemic factors*, since they are normally subsystems or important elements in the systems perspective. The critical systemic

factor constructs are normally deployed to lower level elements and questions, to the point that the critical organizational causes of problems (very often interdependent) can be addressed. Diagnostic models, albeit focused on assessment, have shown the potential to be also used in planning and execution, that is, in all phases of the organization PDCA (plan/do/check/act) cycle [11].

Development Trends: The Systems View

The evolution from organizational assessment models to organizational improvement models seems the most promising trend. However, progress is slow. What is the reason? Why do not managers generally recognize TQM models as management tools, useful to keep their organizations fit. Probably one of the answers (possibly the most relevant) is lack of systems perspective. One could object that all TQM/excellence models claim the systems view. Declarations of interdependence between the components of the system are made, but little is said on how to enhance value generation capabilities through relations – and how to assess consistency of relation maps with strategies. In fact, by developing appropriate relations within a system, unique properties can emerge (the so-called emergent properties) – and organizational excellence is a unique property, that cannot be reached through technology or individual excellence alone.

Separation of enablers from results is a first step toward integration of the systems concepts into the model. Enablers in fact represent the components of the systems (elements, subsystems, and throughput processes) that are critical in relation to the purpose (Figure 3). The purpose is to be found mainly outside the system, in the environment, where customers and stakeholders are located. In the figure, internal stakeholders (like managers and employees) are placed across the border, since they are independent as persons but part of the system in relation to their role. The environment (that represents the suprasystems where the examined system lives and operates) is divided into two parts: the transactional environment (that the organization can influence) and the independent environment, that the organization is bound to understand as well as possible.

A challenge today seems to be developing organizational models that merge quality thinking and

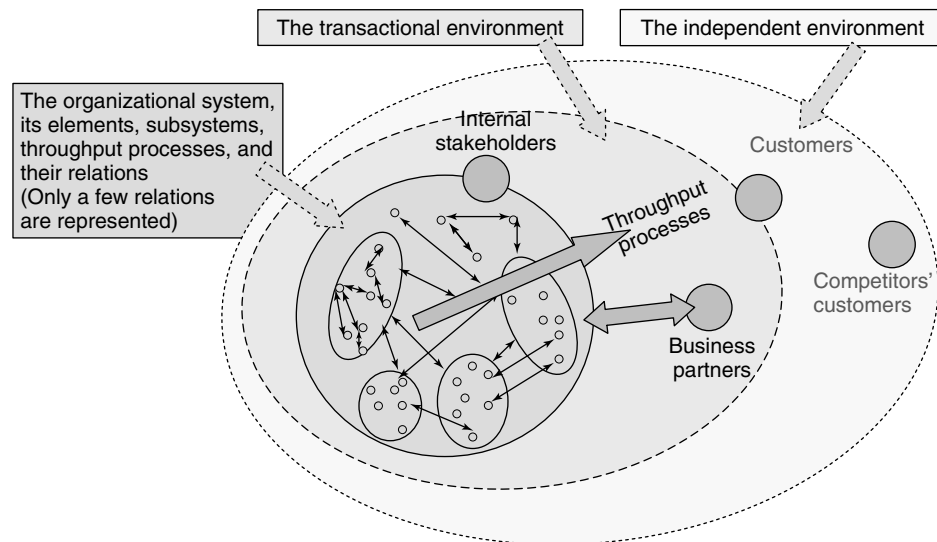


Figure 3 A systemic model of the organization highlighting the main relations within the system and between the system and its environment. It can be used as the basis for building TQM models [12]

systems thinking – and making them acceptable to managers [12].

References

- [1] Juran, J.M. & Godfrey, A.B. (1999). *Juran's Quality Handbook*, 5th Edition, McGraw-Hill, New York, pp. 2.13–2.14, 11.1–11.7.
- [2] Juran, J.M. & Godfrey, A.B. (1999). *Juran's Quality Handbook*, 5th Edition, McGraw-Hill, New York, Section 11.
- [3] Baldrige National Quality Program, <http://www.quality.nist.gov>.
- [4] EFQM, European Foundation for Quality Management, <http://www.efqm.org>.
- [5] The W. Edward Deming Institute, Deming Prize, <http://www.deming.org/demingprize>.
- [6] Kaplan, R.S. & Norton, D.P. (1996). *The Balanced Scorecards*, Harvard Business School Press, Boston.
- [7] Black, S. & Porter, L. (1996). Identification of the critical factors of TQM, *Decision Sciences* **27**(1), 1–20.
- [8] Bittici, U.S., Suwignjo, P. & Carrie, A.S. (2001). Strategic management through quantitative modeling of performance measurement systems, *International Journal of Production Economics* **69**, 15–22.
- [9] Anderson, J., Rungtusanatham, M. & Schroeder, R. (1994). A theory of quality management underlying the Deming management method, *The Academy of Management Review* **19**(3), 472–509.
- [10] Curcovic, S., Melnyk, S., Calantone, R. & Hardfield, R. (2000). Validating the Malcolm Baldrige National Quality Award framework through structural equation modeling, *International Journal of Production Research* **38**(4), 765–791.
- [11] Conti, T. (1999). *Organizational Self-Assessment*, Kluwer Academic Publishers, Dordrecht, Chapter 3.
- [12] Schnauber, H., ed. (2006). Conti T: integration of quality concepts into systems thinking, *Kreative un Konsequent*, Carl Hanser Verlag, Munchen, Chapter 4.

TITO A. CONTI

Assessment of Experimental Designs

Introduction

When we collect data with a designed experiment, we gain the advantage of being able to control the input factors and determine not just a **correlation** between the input factors, also called *explanatory factors*, and the response of interest, but also a causal connection between inputs and output. This more powerful conclusion comes with the additional advantages of being able to specify the region of interest for studying this relationship between factors and response, and to select which data should be gathered during the experiment. In this article, a number of different criteria are presented for consideration when deciding what combinations of input factors to select.

To help frame our discussion and to make some of the trade-offs to be weighed more concrete, we consider an experiment to understand what factors affect the strength of a part being produced during manufacturing.

Some of the preliminaries to consider before selecting a particular design, include how many and which factors to consider. Choose too few, and the possible interrelationships (*see* **Interactions**) between the explanatory variables can never be studied; choose too many, and budgetary or logistical constraints may not allow adequate information to be gathered about each of the factors. In many cases it makes sense to prioritize factors based on current knowledge, and then select designs with potential follow-up designs being possible under a sequential learning and experimenting framework. It is common during the exploratory phases of experimentation to start with a screening design (*see* **Screening Designs**) such as a factorial or fractional factorial design (*see* **Factorial Experiments; Fractional Factorial Designs**) which focuses on estimating major trends in the relationship between the inputs and the response (*see* **Main Effects; Main Effect Designs**). Later when we have gained an understanding of the key inputs on which to focus, a more detailed study of the form of the relationship can be achieved with a response surface design (*see* **Response Surface Methodology**). For our example, determining which factors to include in the experiment will depend

on what is already known about the manufacturing process. If we are just beginning our study, then including many factors (such as temperature, speed of the production line, proportion of various ingredients, type of machines, setup of the machine, suppliers, etc.) of potential interest using a screening design is recommended. If we already have a good understanding of the key factors influencing the process (say, temperature, setup of the machine, and speed of production line), then a response surface design will help with more precise **estimation** of the underlying relationship, and allow us to explore how to optimize the strength of the part within the allowable ranges of our key factors.

Once the factors of interest for the current experiment have been identified, it is important to specify an assumed model for the relationship between the inputs and response, and a range of values for each of the input factors. The assumed model most typically has a parametric form, although nonparametric models are also possible. The form of the model indicates our assumptions about the shape of the underlying relationship to be studied, with more complex relationships requiring bigger models involving more parameters to be estimated. For our part production example, during initial exploration for determining key factors with the screening design, a simple form of the model will allow us to more efficiently gather data on all possible explanatory variables; once we have narrowed our focus to a smaller number of factors, a higher-order model will allow for more flexible underlying relationships to be considered. The designs considered should be matched to the assumed model that we would like to be able to estimate.

The range of input values for the input factors can be independent of each other, which will yield a rectangular region of the appropriate dimension, or constraints on the values of the factors conditional on each other can lead to spherical or oddly shaped regions. For an interesting discussion about the relative merits of different shaped design regions, see Voelkel [1].

There may also be restrictions on how many observations can be collected while experimental conditions are thought to be similar. Hence, choosing a designed experiment which can be implemented in blocks of an appropriate size will yield better consistency of information and unbiased estimation of parameter estimates (*see* **Blocking**). Sometimes if some factors cost more to adjust than others or if there

are restrictions on what randomizations are possible (see **Randomization in Experimental Designs**), a split-plot design (see **Split-Plot Designs**) may be a suitable choice to consider in order to maximize the information gained per unit of cost. For the response surface design of the part production example, it may be the case that altering the setup of the machine takes considerably longer and is more expensive than changing the temperature or the speed of the production line. In this case, the best designs to consider will be split-plot designs with the setup of the machine as a whole-plot factor, and temperature and speed as subplot factors. By not changing the setup of the machine for each observation collected, we can potentially afford to collect more data, and maximize the information gained per unit cost.

Once a model and region of interest have been specified, we can begin to consider choosing a “best” design. We can think of this as a collection of points in the region of interest that allows us to maximize the information gained once the data have been collected, while still being convenient to gather the data. We can think of considering the aspects of a “best” design by looking at several categories of criteria:

1. The ability to estimate well model parameters and make new predictions in the region of interest.
2. Robustness of the design to our choice of model, problems with running the experiment or the quality of the collected data.
3. Qualitative aspects including cost effectiveness, ease of implementation subject to constraints in **data collection**, symmetry across input factors, and inclusion of some data combinations representing best known settings.

These categories with particular characteristics are presented in Box and Draper [2] and Myers and Montgomery [3, p. 304]. In the discussion that follows, some of the trade-offs between different characteristics are considered in more detail.

Estimation and Prediction

Historically, good estimation of the model parameters and prediction for new locations in the design region have been a primary focus when selecting from among candidate designs. Several single-number summaries for the “goodness of the design”

exist and are known as *optimality criteria* (see **Optimal Design**). D- and A-optimality focus on measures of the variance of the model parameters as summarized by the moment matrix, $\mathbf{M} = \mathbf{X}'\mathbf{X}/N$, which is related to the variances and covariances of the regression coefficients of the selected model. G- and I-optimality (also known as *Q*-, *V*-, and *IV*-optimality in the literature) focus on the quality of prediction as a function of the scaled prediction variance of the form

$$SPV = N\mathbf{x}^m'(\mathbf{X}'\mathbf{X})^{-1}\mathbf{x}^m \quad (1)$$

for a completely randomized design, where N is the total number of observations in the design, $\mathbf{X}$ is the matrix for the design, and $\mathbf{x}^m$ is a location in the design region coded to match the selected model.

While these single-number summaries are helpful for making coarse comparisons between designs, often a single number may be too simplistic to capture the differences between two designs. Alternative measures may also be helpful. For example, if we wish to focus on good estimation of model parameters for the screening design to model strength in our part production example, some parameters may be of greater interest than others. For example, consider an experiment involving p factors ($X_1, X_2, \dots, X_p$) in a cuboidal region $[-1, 1]^p$ where a first-order model with two-factor interactions

$$Y = \beta_0 + \sum \beta_i X_i + \sum \sum \beta_{ij} X_i X_j + \varepsilon \quad (2)$$

is thought to be appropriate for summarizing the underlying relationship. Perhaps in some previous experiments, the main effects for each factor may have been studied quite thoroughly, and so focusing on the good estimation of the interaction terms, $\beta'_{ij}s$, is of primary interest. D_s -optimality allows for a single-number summary for the quality of estimation of a subset of the model parameters. In the part production example, if we are in the screening phase of experimentation, we are likely primarily interested in estimating the effects of our many potential factors of interest. This would enable us to discern which factors are most influential on the strength of the part, and hence focusing on D-efficiency can be advantageous.

If the focus of the experiment is to provide good prediction in the design region, G- or I-optimality, which focus on minimizing the worst-case or average scaled prediction variance values, respectively, may not be adequate to summarize all of the important

characteristics of the performance of the design. This focus is typically more common during the response surface design phases of experimentation, where optimizing the process or understanding the underlying relationship between the select few explanatory variables and the response are the key objectives. Various graphical summaries of the prediction properties have been suggested to allow for a more comprehensive look. Variance dispersion graphs (VDGs) proposed by Giovannitti-Jensen and Myers [4] and fraction of design space (FDS) plots by Zahran *et al.* [5] provide complementary two-dimensional summaries of scaled prediction variance throughout the design region. VDGs show the minimum, maximum, and average scaled prediction variance values at different distances from the center of the design space, with three separate curves. FDS plots summarize the range of scaled prediction variance values throughout the region with a single line, and are likely best for making direct comparisons between competing designs. In the FDS plot, prediction variance values are plotted versus the fraction of the design space that has prediction variance at or below the given value. Samples of VDG and FDS plots are shown in Figures 1 and 2, respectively for three designs for a second-order model for two factors on a spherical design space. The three designs – hexagonal or **central composite** designs – are common choices for a response surface methodology design, which involve either six or eight equally spaced points around the circumference of the design space with one or three center runs. All three designs are rotatable (*see Rotatable Designs and Rotatability*), which means that for the VDG, the minimum, maximum, and average scaled prediction variance values are all the same at a given radius. Notice how the VDG shows where good and bad prediction occurs, while the FDS plot allows for direct comparisons between the different designs. In this case, the central composite design with three center runs has uniformly the best proportion of the design space at any scaled prediction variance value.

Adaptations of these graphical methods have been suggested for various specialized situations when there are several groups of factors, each with distinct properties (*see Goldfarb et al.* [6, 7] and Liang *et al.* [8] for more details). Quantile plots [9] also provide another alternative to understand the characteristics of prediction variance.

In addition to considering good estimation of the model parameters and good prediction in the design

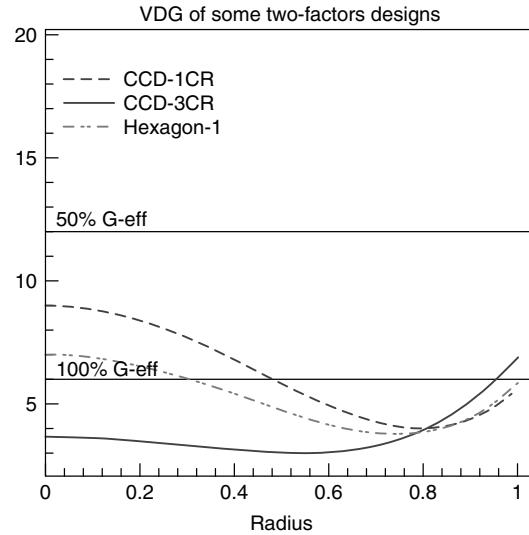


Figure 1 Variance dispersion graph to compare three rotatable designs for a second-order model with two factors on a circular design region. The designs compared are the hexagonal design with one center run, and the central composite design with one and three center runs

space, choosing a design that allows for estimation of pure experimental error, through the inclusion of some replicated design points, is advantageous. This may frequently be implemented with the inclusion of center points (*see Center Points*) in the design. This allows for a measure of the natural variability of the response around the modeled surface that is not dependent on the choice of model or the estimation procedure used. Allocating some observations for replication at one or more locations in the design space can help with the estimation of the experimental error.

Typically, estimation and prediction are considered of primary importance when selecting between competing designs, but as we will see in the following sections some additional aspects of the designs should also be considered.

Robustness of the Design

Because things often do not go exactly as planned, including considerations of robustness can be important to salvaging the results of the experiment if things do go wrong. Recall, that as we began the process of selecting a good design, we specified an

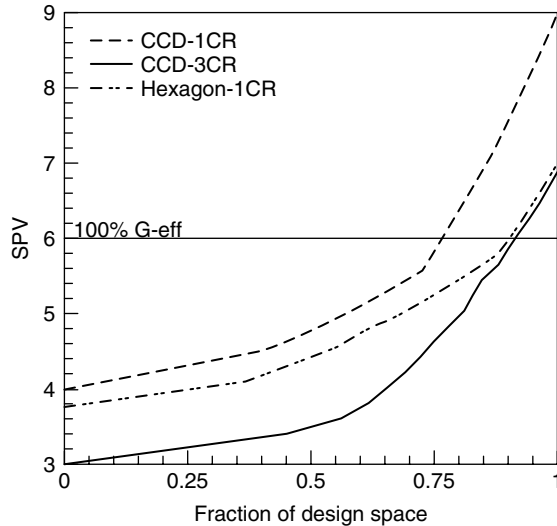


Figure 2 Fraction of design space plot to compare three rotatable designs for a second-order model with two factors on a circular design region. The designs compared are the hexagonal design with one center run, and the central composite design with one and three center runs

assumed model, which summarized the relationship between the input factors and the response. The quality of our estimation and prediction is dependent on the correctness of our assumed model. A good design will allow us to test whether our model is adequate for the observed relationship, or whether additional parameters may be needed to capture its complexity. Typically, the designs which perform best for good estimation or prediction (D-, G-, and I-optimal designs) do not have adequate ability to test for **lack of fit** of the model, and are highly nonrobust to model misspecification. So if the experimenter selects an **optimal design** for a particular model, which later turns out to be incorrect, the design may perform very poorly. Hence allocating some observations to check for the adequacy of the model is also an important consideration. DuMouchel and Jones [10] suggest a Bayesian approach for design selection that allows for the designation of both necessary and “if possible” terms in the model, which offers some protection for assessing the adequacy of the model. Alternately, we may have specified too complex a model for the underlying relationship, and exploring how well the design performs for a number of simpler forms of the model may be beneficial. Ozol-Godfrey *et al.* [11] present an adaptation of the FDS plot that allows

for comparisons of different designs for a variety of nested models. In addition, we can consider the projection properties of the design (*see Projectivity in Experimental Designs*) to assess how well we can expect the design to perform if we subsequently discover that fewer factors than originally assumed are active. For the part production example, exploring the approaches suggested by DuMouchel and Jones [10] and Ozol-Godfrey *et al.* [11] would be most helpful during the response surface methodology phase of experimentation, to protect against the second-order model being insufficiently complex or too complex, respectively, for the actual underlying relationship between temperature, setup of the machine, speed of production line and part strength.

After the data are collected, we may end up with a reduced model with fewer active input factors than initially incorporated; therefore choosing a design that has good projection properties into lower dimensions can be a helpful consideration as well. For example, if it later turns out that speed of the production line is not influential on the response, selecting a good design that has good projection properties in the temperature and setup of machinery design space would be beneficial.

A common additional assumption of the model is that the variance is homogeneous throughout the design region. If this assumption is not correct, then the quality of the estimation can suffer and lead to incorrect inference. Also, prediction precision will not be appropriately characterized. Hence, it is sometimes desirable to allocate some observations to determine the adequacy of this assumption.

Aside from assumptions made about the model, there can be difficulties during the running of the experiment. A good design can help to minimize the impact of these on results. There may be problems in setting the levels of the input factors, for example, the desired level may require a temperature set at 300 °C, but when the experiment is run, the actual temperature is 292 °C. If the experiment is run in a way that the temperature is reset multiple times, then the chance of systematic errors in the factor levels may be reduced. When the response is observed, there may be **outliers** in the response, which a well-designed experiment can help to compensate for with replicates at some or all of the locations. In addition, it may be possible that some run combinations will turn out to be impossible to run or obtain a response value. A good design should be able to cope with

some or all of these problems. It is important to consider carefully which of the above problems are most likely to be an issue in a particular experiment and then allow for additional runs to be allocated to help compensate for potential practical difficulties.

While typically secondary in importance to good estimation and prediction, considering the robustness of the design to model assumptions and the difficulties in running the experiment can help make the results obtained informative, even if everything does not turn out exactly as planned.

Qualitative Aspects of a Good Design

A number of other important aspects of a good design should also be considered. One of the most important considerations that enters into the decision process of choosing a designed experiment is its cost. We can think of cost in terms of two natural metrics, monetary or time, since these are common constraints on our ability to collect data. In many situations, we may have a fixed budget (dollar amount or time with access to the equipment needed to gather the data), and hence we choose a design that maximizes the quality of information we obtain, subject to it being within our budget requirements. For the example considered earlier with a cuboidal region for p factors and an assumed first-order model with interaction terms, we may be able to collect 28 observations. In this case, we would like to find the best design of this size that balances the important subset of the multiple criteria listed in the previous sections.

A second possible scenario for incorporating cost for choosing a best design might be to have an objective of maximizing the quality of information on a per unit cost basis. In this situation, we may choose to run an experiment of size 35 over an experiment of size 28, if we get a proportional improvement in the quality of our estimation or prediction that is greater than the 25% increase in the number of observations. It is not uncommon for the quality of the design to increase nonlinearly with the addition of extra observations, so a small increase in design size can sometimes yield a disproportionate improvement in quality. Scaled prediction variance is an example of how cost, here assumed to be proportional to the number of observations collected, is incorporated into comparisons of designs.

Increasingly, cost considerations are being included in design selection, with Bisgaard [12] and

Liang *et al.* [13] exploring how cost can be appropriately incorporated in split-plot experiments (*see Split-Plot Designs*) with different costs for the setup of whole-plot and subplot factors. There may also be different costs associated with different regions in the design space, and hence design selection may wish to reflect this consideration as well.

Another qualitative aspect of the design that may be considered is the ability to run the experiment subject to constraints on experimental conditions. For example, there may be practical limitations on the number of levels of each factor that we are able to consider. Choosing a design with a moderate number of different levels for each factor can provide cost savings. For our part production example, restricting the number of levels considered to three for each factor will considerably reduce the cost of the experiment and make it easier to run. Since a single experiment will rarely teach us everything that we need to know about the process of interest, it is often beneficial to proceed sequentially by running a series of experiments (*see Sequential Experimentation*). In this case, choosing a design in the early stages that allows for follow-up experiments that incorporate understanding gained from a previous experiment can make for more efficient data collection approaches.

Trade-Offs Between Objectives

Given the extensive list of criteria of a good design that have been outlined above, it makes sense to choose an appropriate subset of these objectives on which to focus for a particular study, based on the goal of the experiment. Many of the objectives require that the locations of observations in the design be selected specifically for that objective. Hence, a design that is D-optimal for a first-order model will likely not have the ability to test whether higher-order terms are needed, to estimate pure experimental error, or protect against difficulties with running the experiment. For our part production example, a D-optimal design might be a fractional factorial design, which will do very well for estimation of the model parameters. We may be well advised to augment the design with some center runs or additional replicates, to allow for not only good estimation but also to enhance our ability to assess various types of robustness.

In particular, many of the objectives require a larger design size, so cost effectiveness is likely at cross-purposes with most of the other criteria. Depending on the nature of the cost constraints of the experiment, working with a fixed **sample size** and looking to maximize several other objectives can be an appropriate strategy. Alternately, exploring different sized designs and determining the quality of information on a per unit cost basis may lead to desirable results.

Since good estimation or prediction are typically the dominant concerns when selecting from competing designs, it may be best to focus on this consideration initially, while reserving some additional observations to augment an optimal design to offer good performance on some of the other criteria. For example, we may construct a D-optimal design and then supplement it with additional observations to allow for estimation of pure experimental error and to assess lack of fit. For our part production example, a D-optimal design might be a fractional factorial design, which will do very well for estimation of the model parameters. A good augmentation to the design would be a number of center runs, to allow for estimation of pure error and a preliminary test of curvature. This approach of selecting a good base design with ideal properties for the primary objective, and then expanding the design to incorporate other measures of a good design can yield some excellent choices. For some situations, a design that estimates the model parameters well may not be best for prediction of new observations. Hence if only one of these objectives is of interest, then choosing a starting design based on an optimality criterion based specifically on that objective may be advantageous.

Conclusions

When selecting a design for an experiment, there are many competing objectives to consider. By prioritizing which of these aspects are important for our particular study, and then strategically choosing a design that performs well for these objectives, we can maximize the quality of information gained during the experiment. Choosing a design that is optimal for a single objective may lead to substantially less than desirable performance with respect to many other

objectives, and hence a conscious balancing on multiple objectives is likely to yield a better overall design.

References

- [1] Voelkel, J.G. (2005). The efficiencies of fractional factorial designs, *Technometrics* **47**, 488–494.
- [2] Box, G.E.P. & Draper, N.R. (1975). Robust designs, *Biometrika* **62**, 347–352.
- [3] Myers, R.H. & Montgomery, D.C. (2002). *Response Surface Methodology*, John Wiley & Sons, New York, p. 304.
- [4] Giovannitti-Jensen, A. & Myers, R.H. (1989). Graphical assessment of the prediction capacity of response surface designs, *Technometrics* **31**, 159–171.
- [5] Zahran, A.R., Anderson-Cook, C.M. & Myers, R.H. (2003). Fraction of design space to assess prediction capability of response surface designs, *Journal of Quality Technology* **35**, 377–386.
- [6] Goldfarb, H.B., Anderson-Cook, C.M., Borror, C.M. & Montgomery, D.C. (2004a). Fraction of design space to assess the prediction capability of mixture and mixture-process designs, *Journal of Quality Technology* **36**, 169–179.
- [7] Goldfarb, H.B., Borror, C.M., Montgomery, D.C. & Anderson-Cook, C.M. (2004b). Three-dimensional variance dispersion graphs for mixture-process experiments, *Journal of Quality Technology* **36**, 109–124.
- [8] Liang, L., Anderson-Cook, C.M. & Robinson, T.J. (2006). Fraction of design space plots for split-plot designs, *Quality and Reliability Engineering International* **22**, 275–289.
- [9] Khuri, A.I., Kim, H.J. & Um, Y. (1996). Quantile plots of the prediction variance for response surface designs, *Computational Statistics and Data Analysis* **22**, 395–407.
- [10] DuMouchel, W. & Jones, B. (1994). A simple Bayesian modification of D-optimal designs to reduce dependence on an assumed model, *Technometrics* **36**, 37–47.
- [11] Ozol-Godfrey, A., Anderson-Cook, C.M. & Montgomery, D.C. (2005). Fraction of design space plots for examining model robustness, *Journal of Quality Technology* **37**, 223–235.
- [12] Bisgaard, S. (2000) The design and analysis of $2^{k-p} \times 2^{q-r}$ split-plot experiments, *Journal of Quality Technology* **32**, 39–56.
- [13] Liang, L., Anderson-Cook, C.M. & Robinson, T.J. (2007). Cost penalized estimation and prediction evaluation for split-plot designs, *Quality and Reliability Engineering International* (in press).

CHRISTINE ANDERSON-COOK

Multivariate Stochastic Orders and Aging

Dynamic Multivariate IFR Notions

A random vector $\mathbf{X} = (X_1, X_2, \dots, X_n)$ is said to be smaller than the random vector $\mathbf{Y} = (Y_1, Y_2, \dots, Y_n)$ in the *multivariate ordinary stochastic order* (denoted by $\mathbf{X} \leq_{\text{st}} \mathbf{Y}$) if $E[\phi(\mathbf{X})] \leq E[\phi(\mathbf{Y})]$ for every function ϕ that is increasing with respect to the coordinate-wise order in $\mathbb{R}^n$, such that the expectations above are well defined. See Shaked and Shanthikumar [1] or Müller and Stoyan [2] for detailed studies of this order.

Let $\mathbf{X} = (X_1, X_2, \dots, X_n)$ be a nonnegative random vector with an absolutely continuous distribution function. Here it is helpful to think about $X_1, X_2, \dots, X_n$ as the lifetimes of n components $1, 2, \dots, n$ that make up some system. Suppose that an observer observes the system continuously in time and records the failure times and the identities of the components that fail as time passes. Thus, a typical “history” that the observer has observed by time $t \geq 0$ is of the form

$$h_t = \{X_I = \mathbf{x}_I, X_{\bar{I}} > t\mathbf{e}\}, \quad 0\mathbf{e} \leq \mathbf{x}_I \leq t\mathbf{e},$$

$$I \subseteq \{1, 2, \dots, n\} \quad (1)$$

(Here, for every n -dimensional vector $\mathbf{x}$ and for any set $I \subseteq \{1, 2, \dots, n\}$, $\mathbf{x}_I$ denotes the vector of the x_i 's such that $i \in I$. Also, here $\bar{I}$ denotes the complement of I in $\{1, 2, \dots, n\}$, and $\mathbf{e}$ denotes the vector of 1's.) In equation (1) I is the set of components that have already failed by time t (with failure times $\mathbf{x}_I$) and $\bar{I}$ is the set of components that are still alive at time t .

Let

$$h'_s = \{X_J = \mathbf{y}_J, X_{\bar{J}} > s\mathbf{e}\}, \quad 0\mathbf{e} \leq \mathbf{y}_J \leq s\mathbf{e},$$

$$J \subseteq \{1, 2, \dots, n\} \quad (2)$$

be another history. If $t \leq s$ and the histories h_t and h'_s are such that each component that failed in h_t also failed in h'_s , and, for components that failed in both histories, the failures in h'_s are earlier than the failures in h_t , then we say that the history h_t is *less severe* or *more pleasant* than the history h'_s and we denote it by $h_t \leq h'_s$. Note that if h_t and h'_s are as

in equations (1) and (2), then $h_t \leq h'_s$ if, and only if, $I \subseteq J$ and $\mathbf{y}_I \leq \mathbf{x}_I$.

Recall that a univariate nonnegative random variable X has the increasing failure rate (IFR) aging property if, and only if, $[X - t | X > t] \geq_{\text{st}} [X - t' | X > t']$ whenever $t \leq t'$ (see **Stochastic Orders and Aging**).

For every vector $\mathbf{a} = (a_1, a_2, \dots, a_n)$ denote by $\mathbf{a}_+$ the vector $(\max\{a_1, 0\}, \max\{a_2, 0\}, \dots, \max\{a_n, 0\})$. Generalizing the univariate IFR characterization mentioned above, we can define a nonnegative random vector $\mathbf{X}$ as multivariate increasing failure rate (MIFR) if for $t \leq s$ we have

$$[(\mathbf{X} - t\mathbf{e})_+ | h_t] \geq_{\text{st}} [(\mathbf{X} - s\mathbf{e})_+ | h'_s]$$

whenever $h_t \leq h'_s$ (3)

Another possibility is to call the nonnegative random vector $\mathbf{X}$ MIFR if for $t \leq s$ we have

$$[(\mathbf{X} - t\mathbf{e})_+ | h_t] \geq_{\text{st}} [(\mathbf{X} - s\mathbf{e})_+ | h'_s]$$

whenever h_t and h'_s coincide on $[0, t]$ (4)

The latter definition is the MIFR $|\mathcal{F}_t$ definition of Arjas [3] where $\mathcal{F}_t$ is the minimal σ field generated by $\mathbf{X}$; see Arjas [3] for more details. Obviously the two MIFR notions, defined by equations (3) and (4), satisfy

$$\text{MIFR}(3) \implies \text{MIFR}(4)$$

The two *different* definitions of MIFR, given in equations (3) and (4), have some desirable properties. For example, a vector consisting of independent IFR random variables is MIFR according to either one of these two definitions. However, perhaps the most important feature of these kinds of definitions is their intuitive interpretation. In the univariate case these two definitions coincide with the univariate definition of IFR. More properties of these MIFR notions can be found in Arjas [3] and in Shaked and Shanthikumar [4].

Let us now recall another characterization of the univariate IFR aging notion, that is, a nonnegative random variable X has the IFR aging property if, and only if, one of the following equivalent conditions hold:

$$[X - t | X > t] \geq_{\text{hr}} [X - t' | X > t'] \quad \text{whenever } t \leq t' \quad (5)$$

2 Multivariate Stochastic Orders and Aging

or

$$X \geq_{\text{hr}} [X - t | X > t] \quad \text{for all } t \geq 0 \quad (6)$$

where $\leq_{\text{hr}}$ denotes the univariate hazard rate stochastic order (see **Stochastic Orders and Aging**). Generalizing these univariate characterizations to the multivariate case we obtain other MIFR notions as described next.

Consider a typical “history” of X as in equation (1). Let $i \in \bar{I}$ be a component that is still alive at time t . Its multivariate conditional hazard rate, at time t , is defined as follows:

$$\begin{aligned} \lambda_{i|I}(t | \mathbf{x}_I) &= \lim_{\Delta t \downarrow 0} \frac{1}{\Delta t} P \{t < X_i \leq t + \Delta t | \mathbf{X}_I \\ &= \mathbf{x}_I, \mathbf{X}_{\bar{I}} > t\mathbf{e}\} \end{aligned} \quad (7)$$

where, of course, $0\mathbf{e} \leq \mathbf{x}_I \leq t\mathbf{e}$, and $I \subseteq \{1, 2, \dots, n\}$. Let $\mathbf{Y} = (Y_1, Y_2, \dots, Y_n)$ be another nonnegative random vector with an absolutely continuous distribution function. Denote its multivariate conditional hazard rate functions by $\eta_{i|I}(\cdot | \cdot)$, where the η ’s are defined analogously to the λ ’s in equation (7). Suppose that

$$\begin{aligned} \eta_{i|I \cup J}(u | \mathbf{y}_I, \mathbf{y}_J) &\geq \lambda_{i|I}(u | \mathbf{x}_I) \quad \text{whenever } J \cap I = \emptyset, \\ \mathbf{y}_I &\leq \mathbf{x}_I \leq u\mathbf{e}, \text{ and } \mathbf{y}_J \leq u\mathbf{e} \end{aligned}$$

where $i \in \overline{I \cup J}$. Then $\mathbf{Y}$ is said to be *smaller than $\mathbf{X}$ in the dynamic multivariate hazard rate order* (denoted as $\mathbf{Y} \leq_{\text{dyn-hr}} \mathbf{X}$); see Shaked and Shanthikumar [1] for more details about this multivariate stochastic order. Recalling the univariate characterization (equation (5)) of the IFR aging notion, one may define a multivariate aging property of $\mathbf{X}$ by requiring $\mathbf{X}$ to satisfy, for $t \leq s$ and histories h_t and h'_s ,

$$\begin{aligned} [(X - t\mathbf{e})_+ | h_t] &\geq_{\text{dyn-hr}} [(X - s\mathbf{e})_+ | h'_s] \\ \text{whenever } h_t &\leq h'_s \end{aligned} \quad (8)$$

Yet another possible multivariate analog of the characterization (equation (5)) is to require $\mathbf{X}$ to satisfy, for $t \leq s$,

$$\begin{aligned} [(X - t\mathbf{e})_+ | h_t] &\geq_{\text{dyn-hr}} [(X - t\mathbf{e})_+ | h'_s] \\ \text{whenever } h_t &\text{ and } h'_s \text{ coincide on } [0, t] \end{aligned} \quad (9)$$

An analog of the characterization (equation (6)) is to require $\mathbf{X}$ to satisfy equations (8) or (9) with $t = 0$; that is,

$$\begin{aligned} \mathbf{X} &\geq_{\text{dyn-hr}} [(X - s\mathbf{e})_+ | h'_s] \\ \text{for any history } h'_s, \quad s &\geq 0 \end{aligned} \quad (10)$$

It turns out that the three conditions, defined in equations (8), (9), and (10), are equivalent:

$$\text{MIFR}(8) \iff \text{MIFR}(9) \iff \text{MIFR}(10)$$

and as such they define a notion of MIFR property. Since the order $\leq_{\text{dyn-hr}}$ is stronger than the order $\leq_{\text{st}}$ (see Shaked and Shanthikumar [1]), this notion is stronger than the MIFR notions defined by equations (3) and (4), that is,

$$\text{MIFR}(8) \implies \text{MIFR}(3)$$

See more details in Shaked and Shanthikumar [4].

Arjas [3] also defined a multivariate notion of decreasing failure rate (DFR). This definition is a version of equation (4) in which the inequality is reversed, but some care is needed for such a definition – only the lifetimes of components that are alive at time s are compared.

A Dynamic Multivariate NBU Notion

Recall that a univariate nonnegative random variable X has the new better than used (NBU) aging property if, and only if, $X \geq_{\text{st}} [X - t | X > t]$ whenever $t \geq 0$ (see **Stochastic Orders and Aging**). Following the idea that led to the MIFR aging notion given in equation (4), one can define a multivariate new better than used (MNBU) aging notion. Explicitly, one can call the nonnegative random vector $\mathbf{X}$ MNBU if

$$\mathbf{X} \geq_{\text{st}} [(X - t\mathbf{e})_+ | h_t] \quad \text{for all } t \geq 0 \quad (11)$$

This is the MNBU $|\mathcal{F}_t$ definition of Arjas [3] where $\mathcal{F}_t$ is the minimal σ field generated by $\mathbf{X}$; see Arjas [3] for more details. Obviously the MIFR notion in equation (4) implies this NBU notion.

Arjas [3] also defined a multivariate notion of new worse than used (NWU). This definition is a version of equation (11) in which the inequality is reversed, but some care is needed for such a definition – only

the lifetimes of components that are alive at time t are compared.

Multivariate Aging Notions Based on Comparisons with Multivariate Exponential Distributions

The univariate aging notions of IFR, IFRA (increasing failure rate average), and NBU, can be characterized, respectively, by the univariate convex transform order $\leq_c$, by the univariate star order $\leq_*$, and by the univariate superadditive order $\leq_{su}$ (see **Stochastic Orders and Aging**). Explicitly, if X is a nonnegative random variable, and Exp denotes an exponential random variable (no matter what the mean is) then

$$X \text{ is IFR} \iff X \leq_c \text{Exp} \quad (12)$$

$$X \text{ is IFRA} \iff X \leq_* \text{Exp} \quad (13)$$

$$X \text{ is NBU} \iff X \leq_{su} \text{Exp} \quad (14)$$

Motivated by equations (12–14), one may define MIFR, multivariate increasing failure rate average (MIFRA), and MNBU aging notions. In order to do that one needs multivariate analogs of the univariate stochastic orders $\leq_c$, $\leq_*$, and $\leq_{su}$, and one needs to choose a suitable multivariate exponential distribution to replace Exp in equations (12–14). Below we first describe multivariate analogs of the orders $\leq_c$, $\leq_*$, and $\leq_{su}$.

Roy [5] considered the following multivariate extensions of the univariate stochastic orders $\leq_c$, $\leq_*$, and $\leq_{su}$. Let $\mathbf{X} = (X_1, X_2, \dots, X_n)$ and $\mathbf{Y} = (Y_1, Y_2, \dots, Y_n)$ be two nonnegative random vectors, with survival functions $\bar{F}$ and $\bar{G}$, respectively; that is, $\bar{F}(\mathbf{x}) = P\{\mathbf{X} > \mathbf{x}\}$ and $\bar{G}(\mathbf{x}) = P\{\mathbf{Y} > \mathbf{x}\}$ for $\mathbf{x} = (x_1, x_2, \dots, x_n) \in \mathbb{R}_+^n$. For $i = 1, 2, \dots, n$, and $\mathbf{x} \in \mathbb{R}_+^n$, denote by $\bar{G}_i^{-1}$ the inverse of $\bar{G}(\mathbf{x})/\bar{G}(x_1, \dots, x_{i-1}, 0, x_{i+1}, \dots, x_n)$ as a function of x_i , and also denote $\bar{F}_i(\mathbf{x}) = \bar{F}(\mathbf{x})/\bar{F}(x_1, \dots, x_{i-1}, 0, x_{i+1}, \dots, x_n)$. Then $\mathbf{X}$ is said to be less than $\mathbf{Y}$ in the multivariate convex transform (star, superadditive) order (denoted by $\mathbf{X} \leq_c$ ($\leq_*$, $\leq_{su}$) $\mathbf{Y}$) if for each $i = 1, 2, \dots, n$, it holds that $\bar{G}_i^{-1}(\bar{F}_i(x_1, x_2, \dots, x_n))$ is convex (star-shaped, superadditive) in x_i for any given $x_1, \dots, x_{i-1}, x_{i+1}, \dots, x_n$.

Next, let **Exp** denote any random vector that has a multivariate Gumbel exponential distribution with

the survival function

$$\begin{aligned} \bar{G}(x_1, x_2, \dots, x_n) = \exp & \left[- \sum_i \lambda_i x_i \right. \\ & \left. - \sum_{i < j} \lambda_{ij} x_i x_j \cdots - \lambda_{1,2,\dots,n} x_1 x_2 \cdots x_n \right], \\ (x_1, x_2, \dots, x_n) \in \mathbb{R}_+^n \end{aligned} \quad (15)$$

where $\lambda_{i_1, i_2, \dots, i_k} > 0$ for each $\{i_1, i_2, \dots, i_k\} \subseteq \{1, 2, \dots, n\}$; see Kotz *et al.* [6, p. 406].

Now one can define MIFR, MIFRA, and MNBU aging notions for a random vector $\mathbf{X}$ by replacing Exp in equations (12–14) by **Exp**, using the multivariate stochastic orders $\leq_c$, $\leq_*$, and $\leq_{su}$ described above. However, historically, proper definitions of MIFR, MIFRA, and MNBU aging notions were introduced first, and then analogs of equations (12–14) were obtained. So this is the way we will proceed here.

Roy [7] proposed multivariate extensions of the IFR, IFRA, and NBU notions as described below. Let $\mathbf{X} = (X_1, X_2, \dots, X_n)$ be an absolutely continuous nonnegative random vector with survival function $\bar{F}$. For $i = 1, 2, \dots, n$, denote by $r_i(\mathbf{x})$ and $A_i(\mathbf{x})$ the failure rate and the failure rate average of the random variable $[X_i | \bigcap_{j \neq i} X_j > x_j]$ at the point x_i , where $\mathbf{x} = (x_1, x_2, \dots, x_n) \in \mathbb{R}_+^n$; that is, writing $R(\mathbf{x}) = -\log \bar{F}(\mathbf{x})$, we have

$$r_i(\mathbf{x}) = \frac{\partial}{\partial x_i} R(\mathbf{x})$$

and

$$A_i(\mathbf{x}) = \frac{1}{x_i} \int_0^{x_i} r_i(x_1, \dots, x_{i-1}, y, x_{i+1}, \dots, x_n) dy$$

According to Roy [7], $\mathbf{X}$ has the

1. MIFR property if $r_i(\mathbf{x})$ is increasing in x_i for all $(x_1, \dots, x_{i-1}, x_{i+1}, \dots, x_n) \in \mathbb{R}_+^{n-1}$ and each i .
2. MIFRA property if $A_i(\mathbf{x})$ is increasing in x_i for all $(x_1, \dots, x_{i-1}, x_{i+1}, \dots, x_n) \in \mathbb{R}_+^{n-1}$ and each i ,
3. MNBU property if

$$\begin{aligned} & \bar{F}(x_1, \dots, x_{i-1}, x_i + y, x_{i+1}, \dots, x_n) \\ & \times \bar{F}(x_1, \dots, x_{i-1}, 0, x_{i+1}, \dots, x_n) \end{aligned}$$

4 Multivariate Stochastic Orders and Aging

$$\begin{aligned} &\leq \bar{F}(x_1, \dots, x_{i-1}, x_i, x_{i+1}, \dots, x_n) \\ &\times \bar{F}(x_1, \dots, x_{i-1}, y, x_{i+1}, \dots, x_n), \end{aligned}$$

for all $y \geq 0$ and for all $\mathbf{x} \in \mathbb{R}_+^n$ and each i .

Roy [5] provided the following characterizations

$$X \text{ is MIFR} \iff X \leq_c \mathbf{Exp} \quad (16)$$

$$X \text{ is MIFRA} \iff X \leq_* \mathbf{Exp} \quad (17)$$

$$X \text{ is MNBU} \iff X \leq_{su} \mathbf{Exp} \quad (18)$$

where $\mathbf{Exp}$ denotes any random vector with survival function given in equation (15). Roy [5] also noted that equations (16–18) hold if $\mathbf{Exp}$ is replaced by a vector of independent exponential random variables.

References

- [1] Shaked, M. & Shanthikumar, J.G. (2007). *Stochastic Orders*, Springer, New York.
- [2] Müller, A. & Stoyan, D. (2002). *Comparison Methods for Stochastic Models and Risks*, John Wiley & Sons, New York.
- [3] Arjas, E. (1981). A stochastic process approach to multivariate reliability systems: notions based on conditional stochastic order, *Mathematics of Operations Research* **6**, 263–276.
- [4] Shaked, M. & Shanthikumar, J.G. (1991). Dynamic multivariate aging notions in reliability theory, *Stochastic Processes and Their Applications* **38**, 85–97.
- [5] Roy, D. (2002). Classification of multivariate life distributions based on partial ordering, *Probability in the Engineering and Informational Sciences* **16**, 129–137.
- [6] Kotz, S., Balakrishnan, N. & Johnson, N.L. (2000). *Continuous Multivariate Distributions, Volume 1: Models and Applications*, 2nd Edition, John Wiley & Sons, New York.
- [7] Roy, D. (1994). Classifications of life distribution under multivariate model, *IEEE Transaction on Reliability* **43**, 224–229.

Related Article

Aging and Positive Dependence.

FÉLIX BELZUNCE AND MOSHE SHAKED

Scan Statistics

Introduction

Scan statistics have been used extensively in several areas of science and technology to test for occurrences of clusters of rare events. In this article, we survey results in the area of scan statistics that are applicable to **quality control** and reliability theory. In reliability theory, scan statistics have been used in the analysis and design of a k -within-consecutive- m -out-of- n : F system of n components arranged in a linear or circular fashion. This system fails if k components fail within any segment of m consecutive components. In the section titled “ k -within-Consecutive- m -out-of- n : F Systems”, we survey results for approximations and bounds for the system reliability of such systems. Two-dimensional scan statistics have been used in the analysis and design of a k -within-consecutive- (m_1, m_2) -out-of- (n_1, n_2) : F system of $n_1 n_2$ components arranged in a rectangular or cylindrical lattice. This system fails if k components fail within any m_1 by m_2 rectangular subregion. In the section titled “ k -within-Consecutive- (m_1, m_2) -out-of- (n_1, n_2) : F Systems” of this article, we survey results for the approximations and bounds for the system reliability of these two-dimensional systems. In quality control, scan statistics have been used in the design and analysis of **acceptance sampling** schemes. In the section titled “Quality Control and Acceptance Sampling Plans”, we survey results that are useful in evaluating the operating characteristics of these acceptance sampling schemes. In material science, reliability of certain materials can be compromised by the appearance of a large number of microcracks in a small area or volume. In the section titled “Reliability of Materials”, we present a survey of two- and three-dimensional scan statistics that have been used to analyze the occurrences of such microcracks. References include related results and references that are not discussed in this article.

k -within-Consecutive- m -out-of- n : F Systems

A k -within-consecutive- m -out-of- n : F system is defined as a system of n components arranged in

a linear or circular sequence. The system fails if and only if k components fail in any consecutive segment of m components [1–5]. Let $X_1, \dots, X_n$ be a sequence of n 0–1 Bernoulli trials, with $P(X_i = 1) = p_i$ and $P(X_i = 0) = 1 - p_i$, $0 < p_i < 1$. If $X_i = 1$ we say that the i th component has failed and if $X_i = 0$ we say that it is functioning. We review here recent results for discrete scan statistics that are relevant to evaluating the reliability of consecutive k -within- m -out-of- n : F systems. We first discuss the linear system. For integers $2 \leq m < n$ and $k \geq 2$ a linear discrete scan statistic is given by:

$$S_m = \max\{X_i + \dots + X_{i+m-1}; 1 \leq i \leq n - m + 1\} \quad (1)$$

We are interested in evaluating $R(k; m, n, p) = P(S_m \leq k - 1)$, the reliability of a k -within-consecutive- m -out-of- n : F system. Most of the research has been focused on the case of independent and identically distributed (i.i.d.) Bernoulli trials with $p_i = p$, $1 \leq i \leq n$. Approximations and inequalities for $P(S_m \leq k - 1)$ are discussed in [6–8; 9, Chapter 13; 10, Chapter 9; 11, Chapter 5]. Exact results for $P(S_m \leq k - 1)$ for restricted values of the parameters have been derived using a combinatorial approach in [12, 13]. Recently, in [14], using a finite Markov chain imbedding method, an exact formula for evaluating $P(S_m \leq k - 1)$ has been derived for a wide range of the parameters. This Markov chain approach for studying reliability systems has been introduced by Chao and Fu [15, 16]. Additional applications of this method are presented in [10, 11, 17, 18].

For the case of independent nonidentically distributed Bernoulli trials and Markov-dependent Bernoulli trials, evaluation of $P(S_m \leq k - 1)$ has been discussed in [10, 19–21].

The following related problem in system reliability is of interest as well. Suppose $P(X_i = 1) = p$ is unknown and we have observed that a components, in the linear system discussed above, have failed and $n - a$ components are functioning. We seek to approximate

$$\begin{aligned} R^*(k; m, n, a) &= P\left(S_m \leq k - 1 \mid \sum_{i=1}^n X_i = a\right) \\ &= P(S_m(a) \leq k - 1) \end{aligned} \quad (2)$$

the reliability of a k -within-consecutive- m -out-of- n : F linear system with a failed components, where

2 Scan Statistics

S_m is defined in equation (1). $S_m(a)$ is called a *conditional scan statistic*.

For integers $k, m, l \geq 2$ and $n = ml$ an exact formula for $R^*(k; m, n, a)$ has been derived in [13, Theorem 1]. If n, m, l are large and $k < a/2$, computing $R^*(k; m, n, a)$ using [13, Theorem 1] becomes impractical. Moreover, the finite Markov chain approach in [14] or the martingale method in [21] cannot be extended to evaluate $R^*(k; m, n, a)$. Approximations for $R^*(k; m, n, a)$ have been investigated in [7].

We now survey relevant results on scan statistics that are applicable for the circular system. For integers $2 \leq m < n$ and $k \geq 2$ a circular discrete scan statistic is given by:

$$S_m^c = \max\{X_i + \dots + X_{i+m-1}; 1 \leq i \leq n\} \quad (3)$$

For the circular system one is interested in evaluating $R^c(k; m, n, p) = P(S_m^c \leq k - 1)$. Accurate approximations for $P(S_m^c \leq k - 1)$ have been derived in [22]. Tight bounds for the system reliability $R^c(k; m, n, p)$ can be obtained from the bounds for the circular discrete scan statistic in [8].

If a components in the circular system discussed above have failed and $n - a$ components are functioning, then one is interested in evaluating

$$R^{*c}(k; m, n, a) = P\left(S_m^c \leq k - 1 \mid \sum_{i=1}^n X_i = a\right) \quad (4)$$

the reliability of a k -within-consecutive- m -out-of- n : F circular system with a failed components. For this system, [8, 22] provide accurate approximations and inequalities, respectively, for $R^{*c}(k; m, n, a)$.

k -within-Consecutive- (m_1, m_2) -out-of- (n_1, n_2) : F Systems

A k -within-consecutive- (m_1, m_2) -out-of- (n_1, n_1) : F system is a system of $n_1 n_2$ components arranged in a rectangular or cylindrical fashion. A k -within-consecutive- (m_1, m_2) -out-of- (n_1, n_2) : F system fails if and only if k or more components fail within any rectangular grid of size m_1 by m_2 within the n_1 by n_2 rectangular or cylindrical lattice of components [1, 23–30]. We first discuss the system with a rectangular arrangement for its $n_1 n_2$ components. We assume that the system consists of identical components and the failure of one or more components does

not affect the failures of other components in the system. The following probability model has been employed for this system. Let $X_{i,j}$, $i = 1, \dots, n_1$ and $j = 1, \dots, n_2$, be i.i.d 0–1 Bernoulli trials with $P(X_{i,j} = 1) = p$, $0 < p < 1$, and $P(X_{i,j} = 0) = 1 - p$. The event $X_{i,j} = 1$ denotes that the i, j component is functioning, while $X_{i,j} = 0$ denotes that the i, j component has failed. Let

$$Y_{i_1, i_2} = \sum_{j=i_2}^{i_2+m_2-1} \sum_{i=i_1}^{i_1+m_1-1} X_{i,j} \quad (5)$$

where $1 \leq i_1 \leq n_1 - m_1 + 1$ and $1 \leq i_2 \leq n_2 - m_2 + 1$, denote the number of components that have failed in the m_1 by m_2 rectangular region with a southwest corner located at (i_1, i_2) . A two-dimensional scan statistic is defined as:

$$S_{m_1, m_2} = \max\{Y_{i_1, i_2}; 1 \leq i_1 \leq n_1 - m_1 + 1, 1 \leq i_2 \leq n_2 - m_2 + 1\} \quad (6)$$

Approximations and bounds for $R(k; m, m, n, n, p) = P(S_{m, m} \leq k - 1)$ have been derived in [28–34].

We now discuss a scan statistic that is relevant for the two-dimensional cylindrical system. For integers $i = 1, 2$, $2 \leq m_i < n_i$ and $k \geq 2$ a cylindrical two-dimensional discrete scan statistic is defined as follows:

$$S_{m_1, m_2}^c = \max\{Y_{i_1, i_2}; 1 \leq i_1 \leq n_1, 1 \leq i_2 \leq n_2\} \quad (7)$$

where Y_{i_1, i_2} is given in equation (5). For the cylindrical system with $n_1 = n_2 = n$ and $m_1 = m_2 = m$, one is interested in evaluating $R^c(k; m, m, n, n, p) = P(S_{m, m}^c \leq k - 1)$. The methods used in evaluating the reliability of the rectangular system can be extended to the cylindrical system.

Suppose that, p the probability of a failure of a component in the system is unknown and the entire system contains a failed components. Then, the reliability of the two-dimensional rectangular lattice system is given by

$$\begin{aligned} R(k; m, m, n, n, a) &= P\left(S_{m, m} \leq k - 1 \mid \sum_{j=1}^{n_2} \sum_{i=1}^{n_1} X_{i,j} = a\right) \\ &= P(S_{m, m}(a) \leq k - 1) \end{aligned} \quad (8)$$

where $S_{m,m}$ is defined in equation (6). Algorithms for evaluating quite accurate approximations for $P(S_{m,m}(a) \leq k-1)$ have been derived in [35]. Similar approximations for a two-dimensional cylindrical system, with a failed components, can be derived using the approach in [35].

Approximation for the rectangular and cylindrical two-dimensional systems discussed above can be extended to the case where $n_1 \neq n_2$ and $m_1 \neq m_2$. Extending these approximations for systems with nonidentical components is computationally complex. Approximations for reliability of two-dimensional systems with nonindependent components is an open problem.

Quality Control and Acceptance Sampling Plans

The use of scan statistics in quality control and acceptance sampling plans has been reported in the statistical literature since 1958 [36], where an i.i.d. Bernoulli trial model is employed. Let $X_1, \dots, X_n, \dots$ be a sequence of 0–1 i.i.d. Bernoulli trials, with $P(X_i = 1) = p$ and $P(X_i = 0) = 1 - p$, $0 < p < 1$. The event $X_i = 1$ denotes that the i th inspected item is defective, while $X_i = 0$ denotes that the i th inspected item is nondefective. Acceptance sampling plans have been designed in which inspection is suspended when k or more defective items are observed in m consecutive inspected items [37]. Let

$$W_{k,m} = \inf \left\{ n \geq m; \sum_{i=n-m+1}^n X_i \geq k \right\} \quad (9)$$

be the waiting time until k or more defective items are observed in m consecutive inspected items. The expected waiting time for $W_{k,m}$, denoted by $E(W_{k,m})$, is called the **average run length** of this inspection plan until suspension. For this model, Troxell [37] developed acceptance sampling plans and evaluated their average run length and the operating characteristic curve for $k = 2$ and 3 and $3 \leq m \leq 11$. Additional results for this model have been derived in [12, 13, 38]. Note that

$$P(W_{k,m} \geq n) = P(S_m \leq k-1) \quad (10)$$

where S_m is defined in equation (1). References for evaluating $P(S_m \leq k-1)$ and deriving accurate approximations and bounds have been given in

the section titled “ k -within-Consecutive- m -out-of- n : F Systems”.

In [37], a sample inspection suspension rule is designated as (k, m) , if k lots are rejected in m or less consecutive inspected lots. Binomial and Poisson models can be used for the numbers of defective items in each of the lots. A more general acceptance sampling rule can be developed on the basis of exceeding a preassigned number of defective items in m or less consecutive inspected lots. In this case, a binomial or a Poisson model is used for the X_i 's in equation (9). Glaz and Naus [6], Chen and Glaz [7], and Karwe and Naus [39] discuss approximations and inequalities for evaluating $P(W_{k,m} \geq n)$ and $E(W_{k,m})$.

Recently, Lachenbruch *et al.* [40] investigated a potential use of the scan statistic defined in equation (1), in a quality control scheme for monitoring blood product manufacturing. Both Bernoulli and binomial models, as well as their conditional counterparts (when the parameter p is unknown and the numbers of defective samples has been recorded), can be utilized in this important application.

Control charts based on scan statistics or moving sums of i.i.d. normal random variables have been discussed in [41–44]. Asymptotic approximations and bounds for $P(W_{k,m} \geq n)$ and $E(W_{k,m})$ have been derived for $m = 2, 3$. In [43, 44], $P(W_{k,m} \geq n)$ and $E(W_{k,m})$ are evaluated *via* simulation. Accurate approximations and bounds for moderate or large values of m have not been developed yet.

Reliability of Materials

In certain materials microcracks develop, on the surface or within the material, at random *via* a **Poisson process** with some intensity λ [45–47]. A large number of microcracks in a small area (volume) on the surface (in the body) of the material can develop into a crack, resulting in deterioration of the material. The following two- and three-dimensional scan statistics play an important role in detecting material fatigue *via* analyses of occurrences of such microcracks.

Let $\mathbf{X}$ be a two-dimensional Poisson process with intensity λ . Following [47], the two-dimensional scan statistic is defined as

$$S_W = S_W(\lambda, A) = \max\{\mathbf{X}(W(\mathbf{x}) \cap A); \mathbf{x} \in R^2\} \quad (11)$$

where $A = [0, T_1] \times [0, T_2]$ is the entire scanned region, W is the scanning set and $W(x)$ is the translate of W by a point $x \in R^2$. For a rectangular scanning set $W = [0, w_1] \times [0, w_2]$, the scan statistic given in equation (11) is given by

$$S_W = \max\{\mathbf{X}([t_1, t_1 + w_1] \times [t_2, t_2 + w_2]); \\ 0 \leq t_i \leq T_i - w_i, i = 1, 2\} \quad (12)$$

In [45–47], accurate approximations for $P(S_W \geq k)$ have been derived for rectangular, circular, and triangular scanning sets. Approximations for the following three-dimensional scan statistic with a rectangular scanning set are derived in [47]:

$$S_W = \max\{\mathbf{X}([t_1, t_1 + w_1] \times [t_2, t_2 + w_2] \\ \times [t_3, t_3 + w_3]); 0 \leq t_i \leq T_i - w_i, \\ i = 1, 2, 3\} \quad (13)$$

where $\mathbf{X}$ is a three-dimensional Poisson process with intensity λ .

If the **intensity** of a Poisson process is unknown, then conditional on the number of observed events, these events are uniformly distributed in the scanning region A . For the two-dimensional case with a rectangular scanning set, Loader [48] derives approximations for $P(S_W \geq k)$. These approximations have not been extended to the three-dimensional case.

References

- [1] Papastavridis, S.G. & Koutras, M.V. (1993). Bounds for reliability of consecutive-k-within-m-out-of-n systems, *IEEE Transactions on Reliability* **42**, 156–160.
- [2] Papastavridis, S.G. & Koutras, M.V. (1994). Consecutive-k-out-of-n systems, in *New Trends in System Reliability Evaluation*, K.-B. Misra, ed, Elsevier Science Publishers, Amsterdam, pp. 228–248.
- [3] Koutras, M.V., Papadopoulos, G.K. & Papastavridis, S.G. (1994). Circular overlapping success runs, in *Runs and Patterns in Probability*, A.-P. Godbole & S.-G. Papastavridis, eds, Kluwer Academic Publishers, Dordrecht, pp. 287–305.
- [4] Koutras, M.V., Papadopoulos, G.K. & Papastavridis, S.G. (1995). Runs on a circle, *Journal of Applied Probability* **32**, 396–404.
- [5] Chao, M.T., Fu, J.C. & Koutras, M.V. (1995). Survey of reliability studies of consecutive-k-out-of-n: F & related systems, *IEEE Transactions on Reliability* **44**, 120–127.
- [6] Glaz, J. & Naus, J.I. (1991). Tight bounds and approximations for scan statistic probabilities for discrete data, *Annals of Applied Probability* **1**, 306–318.
- [7] Chen, J. & Glaz, J. (1999). Approximations for the distribution and moments of discrete scan statistics, in *Scan Statistics and Applications (Statistics for Industry and Technology)*, J. Glaz & N. Balakrishnan, eds, Birkhäuser, Boston, pp. 67–96.
- [8] Chen, J., Glaz, J., Naus, J. & Wallenstein, S. (2001). Bounds for conditional scan statistics, *Statistics and Probability Letters* **53**, 67–77.
- [9] Glaz, J., Naus, J. & Wallenstein, S. *Scan Statistics*, Springer, New York, 2001.
- [10] Balakrishnan, N. & Koutras, M.V. (2002). *Runs and Scans with Applications*, John Wiley & Sons, New York.
- [11] Fu, J.C. & Lou, W.Y.W. (2003). *Distribution Theory of Runs and Patterns and its Applications: A Finite Markov Chain Imbedding Approach*, World Scientific, Singapore.
- [12] Saperstein, B. (1972). The generalized birthday problem, *Journal of the American Statistical Association* **67**, 425–428.
- [13] Naus, J.I. (1974). Probabilities for a generalized birthday problem, *Journal of the American Statistical Association* **69**, 810–815.
- [14] Fu, J.C. (2001). Distribution of scan statistics for a sequence of bi-state trials, *Journal of Applied Probability* **38**, 1–9.
- [15] Chao, M.T. & Fu, J.C. (1989). A limit theorem of certain repairable systems, *Annals of the Institute of Statistical Mathematics* **41**, 809–818.
- [16] Chao, M.T. & Fu, J.C. (1991). The reliability of large series systems under a Markovian structure, *Advances in Applied Probability* **23**, 894–908.
- [17] Koutras, M.V. & Alexandrou, V.A. (1996). Runs, scans and urn model distributions: a unified Markov chain approach, *Annals of the Institute of Statistical Mathematics* **47**, 743–766.
- [18] Koutras, M.V. (1996). On a Markov chain approach for the study of reliability structures, *Journal of Applied Probability* **33**, 357–367.
- [19] Glaz, J. (1983). Moving window detection for discrete data, *IEEE Transactions of Information Theory* **29**, 457–462.
- [20] Fu, J.C. & Chang, Y.M. (2002). On probability generating functions for waiting time distributions of compound patterns in a sequence of multistate trials, *Journal of Applied Probability* **39**, 70–80.
- [21] Glaz, J., Kulldorff, M., Pozdnyakov, V. & Steele, J.M. (2006). Gambling teams and waiting times for patterns in two-state Markov chains, *Journal of Applied Probability* **43**, 127–140.
- [22] Chen, J., Glaz, J. (1999). Approximation for discrete scan statistics on the circle, *Statistics and Probability Letters* **44**, 167–176.
- [23] Salvia, A.A. & Lasher, W.C. (1990). 2-dimensional consecutive-k-out-of-n: F models, *IEEE Transactions on Reliability* **39**, 382–385.
- [24] Koutras, M.V., Papastavridis, S.G. & Papadopoulos, G.K. (1993). Reliability of 2-dimensional consecutive-k-out-of-n: F systems, *IEEE Transactions on Reliability* **42**, 658–661.

- [25] Fu, J.C. & Koutras, M.V. (1994). Poisson approximations for 2-dimensional patterns, *Annals of the Institute of Statistical Mathematics* **46**, 179–192.
- [26] Fu, J.C. & Koutras, M.V. (1995). Reliability bounds for coherent structures with independent components, *Statistics and Probability Letters* **22**, 137–148.
- [27] Koutras, M.V., Papastavridis, S.G. & Petakos, K.I. (1996). Bounds for coherent reliability structures, *Statistics and Probability Letters* **26**, 285–292.
- [28] Makri, F.S. & Psillakis, Z.M. (1997). Bounds for reliability of k -within-connected- (r,s) -out-of- $(m,n):F$ failure systems, *Microelectronics and Reliability* **37**, 1217–1224.
- [29] Akiba, T. & Yamamoto, H. (2001). Reliability of a 2-dimensional k -within-consecutive- $r \times s$ -out-of- $m \times n:F$ system, *Naval Research Logistics* **48**, 625–637.
- [30] Yamamoto, H. & Akiba, T. (2005). Evaluating methods for the reliability of a large 2-dimensional rectangular k -within-consecutive- (r,s) -out-of- $(m,n):F$ system, *Naval Research Logistics* **52**, 243–252.
- [31] Chen, J. & Glaz, J. (1996). Two dimensional discrete scan statistics, *Statistics and Probability Letters* **31**, 59–68.
- [32] Boutsikas, M.V. & Koutras, M.V. (2000). Reliability approximation for Markov chain imbeddable systems, *Methodology and Computing in Applied Probability* **2**, 393–411.
- [33] Boutsikas, M.V. & Koutras, M.V. (2002). Bounds for the distribution of two-dimensional binary scan statistics, *Probability Engineering in the Information Sciences* **17**, 509–525.
- [34] Lin, D. & Zuo, M.J. (2000). Reliability evaluation of a linear k -within- (r,s) -out-of- $(m,n):F$ lattice system, *Probability Engineering in the Information Sciences* **14**, 435–443.
- [35] Chen, J. & Glaz, J. (2002). Approximations for a conditional two-dimensional scan statistic, *Statistics and Probability Letters* **58**, 287–296.
- [36] Roberts, S.W. Properties of control chart zone tests, *Bell System Technical Journal* 1958. **37**, 83–114.
- [37] Troxell, J.R. (1980). Suspension systems for small sample inspection, *Technometrics* **22**, 517–533.
- [38] Huntington, R.J. (1976). Mean recurrence time for k successes within m trials, *Journal of Applied Probability* **13**, 604–607.
- [39] Karwe, V.V. & Naus, J. (1997). New recursive methods for scan statistics probabilities, *Computational Statistics and Data Analysis* **23**, 389–404.
- [40] Lachenbruch, P.A., Foulkes, M.A., Williams, A.E. & Epstein, J.S. (2005). Potential use of the scan statistics for quality control in blood product manufacturing, *Journal of Biopharmaceutical Statistics* **15**, 353–366.
- [41] Roberts, S.W. (1966). A comparison of some control chart procedures, *Technometrics* **8**, 411–430.
- [42] Lai, T.L. (1974). Control charts based on weighted sums, *Annals of Statistics* **2**, 134–147.
- [43] Bauer, P. & Hackl, P. (1978). The use of MOSUMS for quality control, *Technometrics* **20**, 431–436.
- [44] Bauer, P. & Hackl, P. (1980). An extension of the MOSUMS technique for quality control, *Technometrics* **22**, 1–7.
- [45] Alm, S.E. (1997). On the distribution of scan statistics of a two dimensional Poisson process, *Advances in Applied Probability* **29**, 1–18.
- [46] Alm, S.E. (1998). Approximation and simulation of the distributions of scan statistics for Poisson process in higher dimensions, *Extremes* **1**, 111–126.
- [47] Alm, S.E. (1999). Approximations of the distribution of scan statistics of Poisson processes, in *Scan Statistics and Applications (Statistics for Industry and Technology)*, J. Glaz & N. Balakrishnan, eds, Birkhäuser, Boston, pp. 113–139.
- [48] Loader, C.R. (1991). Large deviation approximations to the distribution of scan statistics, *Advances in Applied Probability* **23**, 751–771.

Further Reading

- Boutsikas, M.V. & Koutras, M.V. (2000). Generalized probability bounds for coherent structures, *Journal of Applied Probability* **37**, 778–794.
- Fu, J.C. & Koutras, M.V. (1994). Distribution theory of runs: a Markov chain approach, *Journal of the American Statistical Association* **89**, 1050–1058.

Related Article

k -out-of- n Systems.

JOSEPH GLAZ

Causality

On the Causality Concept

Notions about causes and effects, or causality, must have been around as long as human beings have lived on earth. If human beings, when they came into existence, had not perceived physical dangers as potential causes of harm and death, humanity would not have survived. It seems obvious that when human beings started to construct tools and other equipment, notions of causes and effects had influenced the ideas behind the working mechanisms of these things. In prehistoric time, notions of causes and effects probably ranged from superstitious ones to rational ones, as we know they have in historic time. Even the great Greek philosopher and scientist Aristotle (384–322 B.C.) had superstitious causal ideas. He claimed among other things that the suction pump worked because of *horror vacui*, that is, nature's fear of empty space. This notion was dominating in science until Blaise Pascal (1623–1662) in 1647 performed experiments that provided convincing evidence that the suction pump worked because of the atmospheric pressure.

There is no widely accepted definition of causality, neither a general one nor one in relation to statistics. In spite of this, there is a vast literature on causality, also in relation to statistics (see e.g., the references [1–4]). In [1, p. 97], is given a list of some different ways in which the cause concept has been understood:

- as some essential aspect of a prior state that jointly determines an event, the *effect*, and without which the effect would not occur;
- as any condition tending to increase the probability of the effect;
- as any quantity whose prior variation would result in the subsequent variation of another quantity; and
- almost universally, in connection with the idea that feasible or imaginable interventions in natural systems change, or would change, the distribution features in populations of those features – were the sun to vanish, the planets would go off in nearly straight lines, were drinking water purified, people would live longer.

Causal relationships range from very simple to extremely complicated. A simple example is putting

up some dominoes in a row such that when the first is tilted so that it falls against the second, then all the dominoes are tilted in turn, without further intervention. A complicated example is what happens when a drug cures an illness. The cause-and-effect mechanisms involved in drug treatments are usually only partially understood, so that causal models are insufficient or impossible. This is one reason why it is compulsory to perform **clinical trials** to collect empirical evidence regarding effects of the drug on the illness in question, and on unwanted side effects and often on interactions with other drugs, before the drug is possibly registered for regular use.

There are still difficult and unsolved philosophical problems related to the concept of causality, some of which may well remain unsolved. One such problem is whether everything in the universe is subject to causal laws. Views on causality in science have ranged between extremes. One extreme is the so-called deterministic worldview, stating that everything that will happen in the future in the universe, down to the least detail, is predetermined for all future, owing to causal effects. On the other extreme is the view, based on results in quantum mechanics, that there is no cause for an individual quantum event, and that the universe is fundamentally unpredictable. This view is discussed in [5, p. 743]. Going into the important epistemological and philosophical problems related to the above views on causality is however far beyond the scope of this article.

It should be noted that in spite of the vast literature that exists on causality, causality is usually not an explicit part of the syllabus in science and technological studies at universities. It seems that an intuitive and pragmatic attitude to the concept is usual in such studies, where it is understated as obvious what causal mechanisms mean in different contexts. This seems to be the case also in research in the same areas. Nevertheless, causal thinking has been necessary for the development of natural science and technology. Scientific knowledge consists, to a large extent, of causal explanations and models, often in the form of mathematical or statistical models.

On Establishing Causal Relationships

The assumption that causal relationships exist in nature and that many such relationships can be revealed by human effort, either fully or partly, is

a central part of the basis for natural science and technology.

On Causality and Statistics

In statistical literature it is frequently emphasized that one cannot just from statistical association between phenomena found by **observational studies**, validly conclude that a causal relationship is present. This emphasizing is well founded. A classical example is that only from the fact that lung cancer is more widespread among smokers than nonsmokers, one cannot validly conclude that smoking causes lung cancer. In [6, p. 387], it is stated that: “It may be that persons whose genetic constitution predisposes them to smoke, also predisposes them to lung cancer or even that incipient lung cancer causes people to smoke.” The first possibility corresponds to the classical case of a “third variable” which influences the two variables among which an association is observed. The second possibility implies that the “causal direction” is opposite to what is intuitively reasonable to believe. In [6, p. 388], it is stated that: “The fact that a clear concept of causation requires widely applicable models of a class of phenomena implies that by far the surest, if not the only way to establish (non-spurious) **correlation** is by a series of experiments in which situations to which the model is supposed to apply are set up, and the relevant input variables are deliberately modified by the experimenter. If over a long series of repetitions under widely different circumstances, changes in the input variables are found to be associated with changes in the output variables, then causal relations can be said to be established.” The view that causal relations can be established by experiments as described is widely accepted in science. According to [2, p. v], Albert Einstein (1879–1955) has stated: “Development of Western science is based on two great achievements: the invention of the formal logical system (in Euclidian geometry) by the Greek philosophers and the discovery of the possibility to find out causal relationships by systematic experiment (during the Renaissance).” Accordingly, a reliable way to find out if smoking causes lung cancer would be to perform a randomized trial involving a large group of young people, where half of them commit themselves to smoke at least a certain amount per day for a large number of years, and the other half commit themselves not to smoke during the same period. Such

an experiment would however be neither ethically acceptable nor practically possible.

There is a high degree of general agreement that the existence of a causal relationship often can be concluded from results of experiments, including statistical association, while this is not so for observational studies. However, statistical association found by observational studies can often be used as a clue to a possible causal relation, which has to be investigated closer by additional means.

Empirical versus Mechanistic Causal Models

Regarding causal relationships in science and technology, the ideal is a detailed, subject matter-based understanding and description. Such understanding is termed *mechanistic*, whether mechanical devices are involved or not. Mechanistic understanding of a causal relationship gives deeper insight than a purely empirically based causal relationship established through experiments. In [7] the differences between empirical (causal) models and mechanistic (causal) models are discussed. An example of a type of empirical causal model can be, linear or nonlinear, regression models with one or more explanatory variables that are used for prediction purposes. A necessary presumption for a regression model to be classified as a causal model is that the data on which it is based stem from experiments, with deliberate interventions in the phenomenon, or system, which is modeled. In accordance with the discussion in the section titled “On Causality and Statistics”, a regression model based on data from an observational study should not be considered a causal model. An example of a mechanistic causal model, as opposed to an empirical causal model, is Newton’s law of movement, $f = ma$. This law expresses that when a body of mass m is exposed to a net force of size f , then the body is subject to an acceleration of size a . This model is also an example of a deterministic model, i.e. a model containing no stochastic elements, as opposed to for example a regression model. However, compilation of statistical data was part of Newton’s work when he established the law. According to [8, p. 80], Isaac Newton (1643–1727) was “a master at manipulating observations so that they exactly fitted his calculations” and according to [9], he often reported data with a precision, which was unlikely to be achieved with observational techniques at his time. So, it seems likely that Newton arrived at the law of

movement by means of manipulated data, which fitted the law much better than his actually observed data. However, for all practical purposes, the law can be considered as exactly valid except for elementary particles, which are accelerated up to velocities so large that their mass is noticeably increased according to the theory of relativity. This happens when a body attains a velocity that is not negligible compared to the velocity of light. The above mentioned is related to two interesting perspectives, which are accounted for shortly in the sections titled “Agreement within Acceptable Margins of Error” and “Karl Popper’s Concept of Falsifiability”.

Agreement within Acceptable Margins of Error

It is known for certain that Newton and many other world famous scientists during the history have manipulated experimental data to fit the causal laws they reported (see e.g., [10], especially page 23). According to [8, pp. 81–82], “The concept of agreement within acceptable margins of error did not exist until statistical theory of testing of hypothesis was developed. It was thought that a closer agreement with data implied a more accurate theory and a more convincing evidence for acceptance by the peers. We are now aware – due to the emergence of statistical ideas – that too close an agreement with data might imply a spurious theory!” Researchers nowadays who want to manipulate, or fabricate, data to fit a theory, sometimes put constructed variation into their faked data in order to disguise the faking. This can be designated *second-order faking*. It is, however, often difficult to fabricate realistic variation in data and in many situations, statistical methods for detecting such fabrication exist or they can be constructed. Faking of data, and how to reveal it, is treated briefly in [8, Section 2.3], with references to more comprehensive treatments.

Karl Popper’s Concept of Falsifiability

The science theorist Karl Popper (1902–1994) advocates an interesting point of view on the possibility of establishing scientific laws, including causal laws. He claims that a scientific law cannot be verified, only falsified. This implies that if predictions are made by means of a hypothesized law, and experiments are performed in order to check if the predictions are correct (within acceptable error margins), the

law cannot be verified no matter under how many different conditions predictions and experiments are done without deviations from the predictions. The law is however falsified when the experimental results do not agree with predictions, provided no failures have been done. Popper claims however that one has more reason to have faith in a potential law the more efforts are made to falsify it without succeeding. The concept of falsifiability is central in his teaching. He claims that if a statement cannot be made subject to falsifiability, then it is not a scientific statement. For an account of these ideas, see [11, Chapter 1].

The fact that the mass of a body increases with increased velocity is a consequence of the theory of relativity, which was established more than 230 years after Newton’s law of movement. This new knowledge implied that Newton’s law of movement had to be modified for large velocities. This need for modification indicates that Popper’s skepticism of the possibility of final verification of scientific laws, including causal laws, is sound. Popper’s view is in accordance with the following statement by Einstein: “No amount of experiments can ever prove me right; a single experiment may at any time prove me wrong” (quoted from [12, p. 715]).

Empirical Causal Relationships May Turn Mechanistic with New Knowledge

According to [2, p. 335], “Snell’s law, Hooke’s law, Ohm’s law, and Joule’s law are examples of purely empirical generalizations that were discovered and used much before they were explained by more fundamental principles.” Thus these laws were changed from empirical causal relationships to mechanistic ones by new knowledge. A situation between a purely empirically based causal relationship and a full mechanistic understanding is common, especially when the causal relationship is complicated and entangled. An example: In 1747, James Lind discovered by experiment that oranges and lemons prevented ship passengers from developing scurvy. This led the British Admiralty to mandate the provision of lime juice to all sailors, thereby eliminating scurvy from the navy (described in [13, p. 44]). At that time, the causal relationship in question was purely empirical. Now we know that the working agent in question is vitamin C, which gives some mechanistic understanding. However, a total mechanistic

understanding of why vitamin C effectively prevents against scurvy has not yet been reached, and will possibly never be. It is however likely that a deeper understanding will be reached than what physiology provides today. In [3, p. 156], it is stated that "... even today our mechanistic knowledge is very limited or even nonexistent in many fields. Typically, it is quite good in physics, quite limited in biology, and almost nonexistent in the social sciences." In the next paragraph, it is stated that "Statistics plays its most important role in connection with experience-based causality. In fact, much modern use of statistics may be seen as an extensive, formalized collection and analysis of experience, where mechanistic understanding often plays little role." The authors of [3] use the term *black-box causality* in relation to (purely) experience-based causal relations, and the term *mechanistic causality* in relation to understanding of causal mechanisms. Gaining more mechanistic understanding corresponds to revealing some of the content in the black box.

A Little on Causes and Effects in Quality and Reliability

It is well known among quality experts that variation in production processes, and the implied variation in output from such processes, is a usual and very important cause of reduced quality and reliability of products. This is easy to explain to people with no formal background in quality improvement. For example, almost everybody can easily imagine that if the parts going into a car motor do not have dimensions that are close to their nominal values, this will cause the motor to be subject to fast wear and low reliability, and thus low quality. It is therefore important to keep the variation of the part dimensions small. Statistical process control (SPC) is an important tool for controlling and reducing variation in, among other things, production processes, and thus output from such processes. The concepts *common causes of variation* and *special causes of variation* are important in SPC, and constitute part of the standard terminology. These concepts, and their use related to reduction of variation in processes and products are explained in the chapter on SPC.

Another quality tool where the term *cause* is used is the cause-and-effect diagram. Other names for the diagram are fishbone diagram and Ishikawa diagram.

This diagram is mainly used by competent people in a subject matter area to map perceived causes of unwanted circumstances, and relations/interactions between such causes, in a visual and clearly setup way. It can be used as an idea generating tool, as well as a diagnostic and preventive tool, and is well suited for discussions also.

Reference [14] is an expository paper outlining a Bayesian approach to reliability analysis of network systems comprising a web of interconnected components, which are subject to interacting (or dependent) failures like the power grid and other infrastructure systems. Such systems are often prone to paralyzing collapse caused by a succession of rapid failures, a phenomenon referred to as *cascading failures*. The purpose of the paper is to articulate on aspects of infrastructure reliability, in particular in relation to the notions of chance, cause, interaction, and cascading. It is emphasized that assessing the reliability of an infrastructure system is a key step in its design. Some additional aspects of the concept of cascading failures are treated in [15] in Section 4.8, titled "Causality and models for cascading failures".

Generally, there seems to exist little literature on causality related to statistics in quality and reliability, and on causality related to statistics in technology.

Further Reading – Especially on Newer Treatments of Causality and Statistics

At the beginning of modern statistics, around 1900, the prevailing attitude was that causality should not be treated in relation to statistics. According to [2, pp. 338–339], the seminal reason for this was the discovery of correlation at that time, and that correlation need not imply causality. The strong reluctance of statisticians to concern themselves with the concept of causality has changed gradually, but slowly, over time. There have however been considerable differences among statisticians in attitude toward causality. An interesting discussion of the development of causality in statistics is found in [2, pp. 338–342].

Dawid [16] gives an overview of different ways the concepts of probability and causality have been perceived and treated in the statistical literature and gives a synthesis of some of these ways. The article is highly philosophical, rather theoretical and parts of it are relatively mathematically advanced.

During the past few decades, quite a lot of literature on new theory on causality and statistics has emerged. The treatment is far more formal, even mathematical, than the informal discussion common earlier. Different types of problems, arising in different subject matter areas, have been attacked with different types of models, though without defining causality explicitly. The book by Pearl [2] is considered an authoritative source. Articles by Aalen and Frigessi [3] and Cox and Wermuth [4] are good overview articles on modern approaches to causality and statistics. Glymour's article [1] is also a good source of such information.

Three major formal approaches in the new area are counterfactual models, graphical models and the so-called Granger causality.

Counterfactual models are, to state simply, concerned with analysis of what would have been observed under other circumstances than those which actually were present in a statistical investigation; e.g., what would have been observed in a randomized trial with a treatment group and a control group if the two groups had been exchanged, and what conclusions can be drawn from such analysis. Counterfactual thinking is fundamental in causality.

Graphical models have a long history, beginning in 1921 (see [17]). Important treatments of modern graphical models are given in, among other sources, [2] and [18–22].

Counterfactual modeling started as early as 1923 (see [23]). Important treatments are found in, among other sources, [2] and [24–26].

The seminal paper on Granger causality is from 1969 (see [27]). (Granger shared the Nobel prize in economics in 2003, mainly for other ideas.) Granger causality is defined for two **time series**, $X(t)$ and $Y(t)$. One says that X does not Granger cause Y if prediction of Y based on all predictors is no better than prediction based on all predictors with X omitted. Further developments of Granger causality, with references, are described in [3, pp. 159–162].

References

- [1] Glymour, C. (1998). Causation (Update), in *Encyclopedia of Statistical Sciences*, S. Kotz, C.B. Read & D.L. Banks, eds, John Wiley & Sons, New York, Update Vol. 2, pp. 97–109.
- [2] Pearl, J. (2000). Causality, in *Models, Reasoning and Inference*, Cambridge University Press, Cambridge.
- [3] Aalen, O.O. & Frigessi, A. (2007). What can statistics contribute to a causal understanding? *Scandinavian Journal of Statistics* **34**, 155–168.
- [4] Cox, D.R. & Wermuth, N. (2004). Causality: a statistical view, *International Statistical Review* **72**, 285–305.
- [5] Zeilinger, A. (2005). The message of the quantum, *Nature* **438**, 743.
- [6] Barnard, G.A. (1982). Causation, in *Encyclopedia of the Statistical Sciences*, S. Kotz, N. Johnson & C.B. Read, eds, John Wiley & Sons, New York, pp. 387–389.
- [7] Box, G.E.P., Hunter, J.S. & Hunter, W.G. (2005). *Statistics for experimenters*, John Wiley & Sons, Hoboken.
- [8] Rao, C.R. (1989). Statistics and Truth, *Putting Chance to Work*, World Scientific, Singapore.
- [9] Westfall, R.S. (1973). Newton and the fudge factor, *Science* **179**, 751–758.
- [10] Broad, W. & Wade, N. (1985). *Betrayers of the Truth – Fraud and Deceit in Science*, Oxford University Press, Oxford.
- [11] Popper, K.R. (1959). *The Logic of Scientific Discovery*, Hutchinson & Company, London.
- [12] Jaynes, E.T. (2003). *Probability Theory – The Logic of Science*, Cambridge University Press, Cambridge.
- [13] Greene M.F., ed. (2000). Looking Back on the Millennium in Medicine, *New England Journal of Medicine* **342**(1), 42–49.
- [14] Lindley, D.V. & Singpurwalla, N.D. (2002). On exchangeable, causal and cascading failures, *Statistical Science* **17**(2), 209–219.
- [15] Singpurwalla, N.D. (2006). Reliability and risk, *A Bayesian Perspective*, John Wiley & Sons, Chichester.
- [16] Dawid, P.A. (2004). Probability, causality and the empirical world: a Bayes–de Finetti–Popper–Borel synthesis, *Statistical Science* **19**(1), 44–57.
- [17] Wright, S. (1921). Correlation and causation, *Journal of Agricultural Research* **20**, 557–585.
- [18] Lauritzen, S.L. (1996). *Graphical models*, Clarendon, Oxford.
- [19] Lauritzen, S.L. (2000). Causal inference from graphical models, in *Complex Stochastic Systems*, O.E. Barndorff-Nielsen D.R. Cox & C. Klüppelberg, Chapman and Hall, London, pp. 63–107.
- [20] Lauritzen, S.L. & Wermuth, N. (1989). Graphical models for associations between variables, some of which are qualitative and some quantitative, *Annals of Statistics* **17**, 31–57.
- [21] Wermuth, N. & Lauritzen, S.L. (1990). On substantive research hypotheses, conditional independence graphs and graphical chain models (with discussion), *Journal of the Royal Statistical Society. Series B* **52**, 21–72.
- [22] Robins, J.M. (1987). A graphical approach to the identification and estimation of causal parameters in mortality studies with sustained exposure periods, *Journal of Chronic Diseases* **40**(Suppl. 2), 139–161.
- [23] Neyman, J. (1923). On the application of probability theory to agricultural experiments, Essay on principles. Section 9. (In French.) English translation of excerpts by

- D. Dabrowska and T. S. Speed, *Statistical Science* 1990 **5**, 463–472.
- [24] Rubin, D.B. (1974). Estimating causal effects of treatments in randomized and nonrandomized studies, *Journal of Educational Psychology* **66**, 688–701.
- [25] Ginsberg, M.L. (1986). Counterfactuals, *Artificial Intelligence* **30**, 35–79.
- [26] Balke, A. & Pearl, J. (1994). Probabilistic evaluation of counterfactual queries, in *Proceedings of the Twelfth National Conference on Artificial Intelligence (AAAI-94)*. Seattle, WA, July 31–August 4 1994. MIT Press, Cambridge, Vol. 1, pp. 230–237.
- [27] Granger, C.W.J. (1969). Investigating causal relations by econometric models and cross-spectral methods, *Econometrica* **37**, 424–438.

ØYSTEIN EVANDT

Process Capability Indices, Alternatives to

Introduction

The entire issue of the journal of quality technology (JQT) in October 1992 JQT was devoted to the topic of process capability indices (PCIs). The five articles covered recent developments of PCIs, confidence bounds on PCIs, relationships to squared error loss, and the distributional and inferential properties of PCIs [4–8]. Additionally, the editor, Peter Nelson discussed the various problems with PCIs and further stated the reason for the publication was “to give those unfortunate enough to have to use capability indices a clear understanding of the assumptions involved and of what the indices are actually measuring” [9]. Finally, with the publication of that issue of JQT, *everything* you ever wanted to know about PCIs had been discussed and PCIs would no longer be abused. Unfortunately, the articles were probably too heavily involved mathematically and the messages contained in the JQT issue were not fully appreciated. Many companies still rely on PCIs to measure and compare the quality of their vendors’ processes without understanding the underlying assumptions of PCIs and the properties of the *estimators* of these indices.

Preceding the JQT issue, the topic of capability indices had been a source of much interest and debate from their popularization in the early 1980s. The indices C_p and C_{pk} were regularly taught in most introductory statistical process control (SPC) classes and usually without caveats about their shortcomings. Efforts to improve PCIs continue as evidenced by the fact that in a recent issue of JQT [10], an article by Wang *et al.* compared three multivariate PCIs.

A multitude of indices exist and C_p , C_{pk} , C_{pm} , C_{pmk} , C_{psk} , C_{pw} , P_p , P_{pk} , $C_p M$, MC_{pM} , and so on, are a few examples. Many improvements have been proposed such as the use of **confidence intervals** to show the uncertainties in the calculated estimates of the indices. Unfortunately, these are not practical for large sets of indices and the widths of these intervals increase as the estimate increases. It is not clear how to demonstrate that a required index has been met. Must the required value exceed the lower confidence limit or is it sufficient to be contained

in the interval? Another suggestion has been to use tolerance regions (see Wang [10]), but how can these be used in a practical way for large sets of indices? In addition, there is some difficulty in interpretation of these regions.

The purpose of this paper is to introduce a basic approach to meet most of the goals of PCIs that are presently being used without some of the associated problems that have been addressed in the previous literature. The intention is to create a foundation and direction for future expansion rather than to formulate a total solution. In cases where simple assumptions (normality and symmetric specifications) are not justified, the usual modifications can easily be introduced to meet the requirements of the application.

Specification Limits and Process Capability Indices

PCIs are often used in two specific phases. In phase I, the capability of a process to meet required specifications is determined. Three approaches are often used and are explained below.

The first approach is to determine the process specifications by needs (product performance, process optimization, customer requirements, etc.). A technology development group then determines the process distribution. The PCIs are calculated and if they are not met, the process is returned to the development group for the needed improvement to meet the required PCIs.

The second alternative is to determine the process specifications by needs as in the initial approach. A set of PCIs is specified and the technology group determines the process distribution. If the process does not meet the PCIs, a new set of indices is calculated that can be met.

The third technique always works the first time. The technology group determines the process distribution. A set of PCIs are specified and then the process specifications are determined from the combination of the process distribution and the PCIs. Unfortunately, this method is used more often than expected particularly when the specifications are on intermediate process steps and not on final product parameters.

In phase II, the goal is to measure changes in the process (either degradations or improvements) over time.

2 Process Capability Indices, Alternatives to

Before continuing, perhaps a few comments on how specification limits might best be determined are worthwhile. When the limits are determined for process and product development, the following ideas are often utilized. The design requirements of the product are determined, simulation models are run, engineering judgment is factored into the picture, and last, but not least, customer agreement is needed. These limits are then validated by data from product or process window characterizations and/or by historical data if it exists. Specifications used in manufacturing are often obtained by this data validation approach.

An Alternative Approach to Process Capability Indices in Phase I

The alternative is to measure whether the process distribution meets the required specifications by matching the mean (or a measure of the center) and the variance (or a measure of the spread) against those values assumed by the specifications. The criteria for the matching of the mean and standard deviation will be based on ratios that encompass both the sample size of the data set as well as a level of confidence associated with the ratio. Consequently, it would be feasible, if desired, to make comparisons between requirements, time periods, processes, and so on, in a statistically meaningful procedure. The usual caveats applied to PCIs must still hold. The process must be stable and the specifications must be technically based. In addition, a set of initial assumptions will be made. For this discussion it will be assumed that the process is normally distributed with mean μ and variance σ^2 . The specifications will be assumed to be symmetric about the mean and the specification width reflects $k\sigma$. These assumptions can be modified for future work. The matching of the variance will be addressed first.

Let $x_1, x_2, \dots, x_n$ represent a sample of n data points from the normal distributed process under study with $\bar{X}$ and S^2 the unbiased estimators of the mean and variance. Thus,

$$\frac{(n-1)S^2}{\sigma^2} \sim (\chi_{n-1}^2) \text{ so that} \quad (1)$$

$$\Pr \left[\frac{(n-1)S^2}{\sigma^2} < \chi_{1-\alpha, n-1}^2 \right] = 1 - \alpha \text{ and} \quad (2)$$

$$\Pr \left[S^2 < \frac{(\chi_{1-\alpha, n-1}^2)\sigma^2}{n-1} \right] = 1 - \alpha \quad (3)$$

Hence,

$$\Pr \left[S < \sigma \sqrt{\frac{(\chi_{1-\alpha, n-1}^2)}{n-1}} \right] = 1 - \alpha \quad (4)$$

Similarly, at the lower tail of the distribution,

$$\Pr \left[\sigma \sqrt{\frac{(\chi_{\alpha, n-1}^2)}{n-1}} < S \right] = 1 - \alpha \quad (5)$$

Denote σ_{upper} and σ_{lower} as follows:

$$\sigma_{\text{upper}} = \sigma \sqrt{\frac{(\chi_{1-\alpha, n-1}^2)}{n-1}} \quad (6)$$

and

$$\sigma_{\text{lower}} = \sigma \sqrt{\frac{(\chi_{\alpha, n-1}^2)}{n-1}} \quad (7)$$

Using the sample size of the data set, a predetermined probability level α , and the value of σ reflected by the specification, the values of σ_{lower} and σ_{upper} can be determined. From process data, an unbiased estimate of σ^2, s^2 , may be obtained.

Now define the ratios:

$$S_{\text{ru}} = -\frac{s}{\sigma_{\text{upper}}} \quad (8)$$

and

$$S_{\text{rl}} = \frac{\sigma_{\text{lower}}}{s} \quad (9)$$

With these definitions,

$$\Pr[S_{\text{ru}} < -1] = \alpha \text{ and } \Pr[S_{\text{rl}} > 1] = \alpha \quad (10)$$

If $S_{\text{ru}} < -1, s > \sigma$ with α significance (based on the sample size) then the process needs improvement with respect to σ .

If $S_{\text{rl}} > 1, s < \sigma$ with α significance (based on the sample size) then the process is performing better than expected and is clearly capable with respect to σ .

If neither condition exists, the difference between s and σ is not statistically significant and no action is necessary.

A similar approach may be used for matching the mean of the process distribution. However, with large sample sizes, a statistical significance may not reflect the practical or engineering significance. Therefore, it is necessary to define a “minimum engineering significant difference” (MESD) that is of real importance.

As stated earlier, $\bar{X}$ and S^2 are the unbiased estimators of the mean and variance. Thus,

$$\sqrt{n} \frac{\bar{X} - \mu}{S} \sim (t_{n-1}) \quad (11)$$

so that

$$\Pr \left[-t_{\alpha/2, n-1} < \sqrt{n} \frac{\bar{X} - \mu}{S} < t_{\alpha/2, n-1} \right] = 1 - \alpha \quad (12)$$

and

$$\Pr \left[-1 < \sqrt{n} \frac{\bar{X} - \mu}{t_{\alpha/2, n-1} S} < 1 \right] = 1 - \alpha \quad (13)$$

Define the mean ratio (MR) as

$$\sqrt{n} \frac{\bar{X} - \mu}{t_{\alpha/2, n-1} S} \quad (14)$$

Hence,

If $MR > 1$ and $(\bar{X} - \mu) > MESD$, then $\bar{X} > \mu$, with significance α .

If $MR < -1$ and $(\bar{X} - \mu) < -MESD$, then $\bar{X} < \mu$, with significance α .

Otherwise, no significant difference exists.

Phase II – Measure Changes in the Process Over Time

The approach to measure changes over time uses the two-sample analogs of phase I, which are to test the significance of the differences in the variance (spread) and the mean (center) from the previous time period. The initial assumptions are similar to those of the one-sample case with respect to symmetry, spec width, and the underlying distributions for the present (new) and the past (old) time periods. Additionally, assume no change has taken place between the period so that $\sigma_{old} = \sigma_{new}$ and $\mu_{old} = \mu_{new}$. As in the earlier case the comparison of variances will be addressed first.

Now,

$$\left(\frac{S_{new}^2}{S_{old}^2} \right) \sim (Fn_{new} - 1, n_{old} - 1) \quad (15)$$

so that

$$\Pr \left[\left(\frac{S_{new}^2}{S_{old}^2} \right) < (F_{1-\alpha}, n_{new} - 1, n_{old} - 1) \right] = 1 - \alpha \quad (16)$$

This leads to

$$\Pr[S_{new}^2 > (S_{old}^2)(F_{1-\alpha}, n_{new} - 1, n_{old} - 1)] = \alpha \quad (17)$$

and

$$\Pr[S_{new}^2 < (S_{old}^2)(F_{\alpha}, n_{new} - 1, n_{old} - 1)] = \alpha \quad (18)$$

As with the one-sample case, denote S_{upper} and S_{lower} as follows:

$$S_{upper} = (S_{old})\sqrt{(F_{1-\alpha}, n_{new} - 1, n_{old} - 1)} \quad (19)$$

and

$$S_{lower} = (S_{old})\sqrt{(F_{\alpha}, n_{new} - 1, n_{old} - 1)} \quad (20)$$

From the past and present process data calculate s_{old}^2 and s_{new}^2 . Using the sample sizes of the two sets of data and a predetermined probability level, α , S_{lower} , and S_{upper} can be determined.

Now define the ratios:

$$S_{cru} = -\frac{s_{new}}{s_{upper}} \quad (21)$$

and

$$S_{crl} = \frac{s_{lower}}{s_{new}} \quad (22)$$

With these definitions,

$$\Pr[S_{cru} < -1] = \alpha \text{ and } \Pr[S_{crl} > 1] = \alpha \quad (23)$$

If $S_{cru} < -1$, $s_{new} > s_{old}$ with α significance based on the sample sizes then the process needs improvement with respect to s_{old} .

If $S_{crl} > 1$, $s_{new} < s_{old}$ with α significance based on the sample sizes then the process is performing better than expected and is clearly capable with respect to s_{old} .

If neither condition exists, the difference between s_{old} and s_{new} is not statistically significant and no action is necessary.

4 Process Capability Indices, Alternatives to

Continuing with the same assumptions previously stated and using an analog to the one-sample case, define the comparative mean ratio (CMR) as follows:

$$\frac{(\bar{X}_{\text{new}} - \bar{X}_{\text{old}})}{(s^*)(t_{\alpha/2, \nu})} \quad (24)$$

where s^* is the appropriate **standard error** of the estimate determined after the variances have been tested, using a pooled estimate of σ if no change is indicated and the calculated values of s_{old}^2 and s_{new}^2 if a change is indicated, and ν is the **degrees of freedom** for the combined sample.

Hence,

If $CMR > 1$ and $(\bar{X}_{\text{new}} - \bar{X}_{\text{old}}) > MESD$, then $\bar{X}_{\text{new}} > \bar{X}_{\text{old}}$ with significance α .

If $CMR < -1$ and $(\bar{X}_{\text{new}} - \bar{X}_{\text{old}}) < -MESD$, then $\bar{X}_{\text{new}} < \bar{X}_{\text{old}}$ with significance α .

Otherwise, no significant difference exists.

Example

The example is discussed in Spiring's article in JQT 29 [11]. The process is a blow-molding procedure for plastic bottles where the outside lip diameter is the measure under study. He states that the upper specification limit (USL) is 0.846, lower specification limit (LSL) is 0.814, and the target is 0.830. A sample of 100 observations were collected with $\bar{X} = 0.8254$ and $\hat{\sigma}$ (m.l.e.) = 0.005. The author then calculates the following estimates and 95% confidence intervals:

$$\hat{C}_p = 1.067 \text{ (0.914, 1.209)} \quad (25)$$

$$\hat{C}_{pk}^* = 0.760 \text{ (0.688, 0.881)} \quad (26)$$

$$\hat{C}_{pm} = 0.785 \text{ (0.686, 0.836)} \quad (27)$$

What has changed? Does a problem exist with this process? The first index confirms that the process spread is not a problem, but does not address the possibility that it is better than the specification suggests. The second and third indices and their confidence intervals point to a problem in the mean of the process.

Assuming that the process is normally distributed and that the specification is defined as $\mu \pm 3\sigma$, then $\mu = 0.830$ and $\sigma = 0.016/3 = 0.0053$.

Since $\hat{\sigma}$ (m.l.e.) = 0.005, then $s = 0.005025$ which is less than the hypothesized σ . Therefore,

$S_{\text{rl}} = \frac{\sigma_{\text{lower}}}{s}$ is the appropriate statistic to evaluate. For this example,

$$\begin{aligned} \sigma_{\text{lower}} &= \sigma \sqrt{\frac{(\chi^2_{\alpha, n-1})}{n-1}} \\ &= (0.0053) \sqrt{\frac{77.046}{99}} = 0.0047 \end{aligned} \quad (28)$$

and $S_{\text{rl}} = 0.9353 < 1$. Consequently, based on the sample size of 100 and the one-sided 95% confidence interval, there is no significant difference in the standard deviation of the process.

Looking at the $MR = \sqrt{n} \frac{\bar{X} - \mu}{t_{\alpha/2, n-1} s}$ and assuming the MESD is approximately one σ , the

$$MR = \sqrt{100} \left(\frac{0.8254 - 0.830}{(1.98422)(0.005025)} \right) = -4.614 \quad (29)$$

Now, $MR = -4.614 < -1$, but since $(\bar{X} - \mu) = -0.0046$ and the MESD was assumed to be $\sigma = 0.0053$, no action would be taken despite the statistical significance of the result. However, if the MESD was less than 0.0046, the statistical significance would also be of engineering significance.

Other Issues

This paper has addressed only the relatively simple case in which the specifications are two-sided symmetrical and the process is assumed to be normally distributed. However, because the technique uses the distribution parameters in the metric, it may easily be adapted to applications involving other assumed distributions and their parameters such as discrete and nonsymmetric distributions without an intermediate transformation or fitting to a **normal distribution**.

In cases where these conditions are not true, the use of medians, interquartile ranges, and other standard robust procedures may be used in a similar manner.

Comments and Conclusions

The advantages of this technique include the following:

1. Since the sample size and confidence levels are built directly into the ratios, they may be sorted to

- prioritize improvement efforts unlike the present PCIs.
2. The ratios can be used to measure changes over time.
3. The concepts can be expanded to other distributions in a direct manner.
4. The ratios add statistical thinking to the measurement processes so that meaningful changes are recognized.
5. Many of the problems associated with the traditional PCIs have been eliminated.

Although the advantages are clear, the implementation of this procedure to the areas where PCIs are presently used will be difficult to accomplish. A great deal of effort will be required to change the mindset with respect to the *status quo* and this must occur at all levels within a company from top management to the worker on the factory floor.

References

- [1] Gunter, B. (1989). The use and abuse of Cpk: parts 1–4, *Quality Progress* (January): **22**(1), 72–73; (March): **22**(3), 108–109; (May): **22**(5), 79–80; (July): **22**(7), 86–87.
- [2] Benson, E., Kotz, S. & Alt, F.B. (1994). The next generation of PCI's, in *Proceedings Joint Statistical Meetings*, August 18, Toronto.
- [3] Process Capability Indices Issue, *Journal of Quality Technology* **24**, (October).
- [4] Rodriguez, R.N. (1992). Recent developments in process capability analysis, *Journal of Quality Technology* **24**, 176–187. Also, an extensive reference list appears in this article.
- [5] Kushler, R.H. & Hurley, P. (1992). Confidence bounds for capability indices, *Journal of Quality Technology* **24**, 188–195.
- [6] Franklin, L.A. & Wasserman, G.S. (1992). Bootstrap lower confidence limits for capability indices, *Journal of Quality Technology* **24**, 196–210.
- [7] Johnson, T. (1992). The relationship of Cpm to squared error loss, *Journal of Quality Technology* **24**, 211–215.
- [8] Pearn, W.L., Kotz, S. & Johnson, N.L. (1992). Distributional and inferential properties of process capability indices, *Journal of Quality Technology* **24**, 216–231.
- [9] Nelson, P. (1992). Editorial, *Journal of Quality Technology* **24**, 175.
- [10] Wang, F.K., Hubele, N.F., Lawrence, F.P., Miskulin, J.D. & Shahriari, H. (2000). Comparison of three multivariate process capability indices, *Journal of Quality Technology* **32**, 263–275.
- [11] Spiring, F.A. (1997). A unifying approach to process capability indices, *Journal of Quality Technology* **29**, 49–58.

Further Reading

- Dovich, R.A. (1991). Statistical terrorists, *Quality in Manufacturing* 14–15.
- Johnson, N.L. & Kotz, S. (1993). *Process Capability Indices*, Chapman & Hall, London.
- Potter, R.W. & Pantula, S.G. (1991). *Confidence Intervals for Capability Indices in Nested Experiments*, SEMATECH, Austin.

RICHARD I. POST

Statistics in Data and Information Quality

Introduction

This article discusses statistical issues and opportunities in data and information quality. Just as managing and improving the quality of manufactured goods depends on statistics for its foundations and makes ample use of statistical methods, so too for data and information. But data and information differ from manufactured products in some critical respects. We focus on these differences and their implications for **quality management**.

After providing a brief historical perspective and describing what we mean by data and information, we focus on the following topics:

- Measurement issues, which stem from the fact that, unlike manufactured products, data and information are intangible.
- Opportunities to automatically correct data errors, which stem from the fact that most data are highly redundant.
- Statistical control issues, which stem from the fact that “interesting” data differ from one another.
- Inconsistency issues, which stem from the fact that data and information can be copied so easily. But it is very difficult to ensure that subsequent changes in one copy are reflected in all other copies.
- Matching issues, which stem from the fact that real-world entities may be identified in different ways.
- Privacy and security issues, which stem, in part, from the fact that “stolen data” are not actually missing.

High-Quality Data and Information Are Increasingly Important

Quality professionals have always depended on measurement data and information to bring manufactured goods into closer conformance to requirements, to identify and eliminate root causes, to improve customer satisfaction, and to drive business results. Similarly, generals, managers, scientists, and professionals

have always depended on data and information to craft strategy, to keep operations humming along, to test and advance theory, and to develop new insights. So, data and information have always been important, at least to an elite few.

However, a number of trends are conspiring to make data and information more and more important to virtually everyone. First is the spectacular and sustained development of information technologies that create, store, process, and communicate data and information. Second is the spectacular and sustained growth in the sheer quantity, variety, and complexity of new data and information. Third is the penetration of data, information, and the associated information technologies into virtually all human endeavors, be they public, private, or personal (in the developed world certainly and more slowly in the developing world). From a quality standpoint, this means more customers, with more needs, more diverse needs, and, in many cases, little training on what data mean, their nuances, or proper ways to interpret them.

Fourth and not surprisingly, data and information are responsible for an increasing share of the economy and especially new economic growth. For example, in the shipping industry it is estimated that, by the year 2010, half the value associated with delivering a container halfway around the world will stem from the data and information associated with the container (Taggart [1]). Fifth and unfortunately, recent data and information quality “disasters”, from the disputed year 2000 presidential election, to the corporate collapses brought on by misstated financial results, to the case for the Iraq War, to privacy breaches have raised the importance of data and information quality in the eye of public.

All of these underlying trends appear to be accelerating, so we should expect data and information quality to grow even more important.

Data and Information Defined

We noted above that practitioners are familiar with “data” and “information” in the context of quality management and experimentation. But since data and information are important in so many settings, we will first define these two terms.

Importantly, “data” is a multifaceted notion, and philosophers, computer scientists, and information systems engineers have proposed many competing

approaches to defining the term. We have selected the definition used herein (Fox *et al.* [2], Tsichritzis and Lochovsky [3]) because it best exposes how data are created and used in organizations, which is essential if they are to be managed effectively.

Thus, data (or a data collection, dataset, etc.) consist of two related components: a *data model* and *data values*. Data models are abstractions of the real world that define what the data are all about. Models define entity classes, the collections of things (physical or ideal) of interest, and attributes and relationships, which specify pertinent features of the entities. For example, the reader is an entity. His or her employer is interested in the entity class “EMPLOYEES” and attributes such as NAME, DEPARTMENT, SALARY, and MANAGER. Each attribute must have a well-defined domain of possible values. For example in the United States, all possible nine-digit numbers define the domain of possible values for SOCIAL SECURITY NUMBER (SSN). Data values are to be drawn from the domain and assigned to attributes for specific entities. A single datum is thus a triple:

Datum = <Entity, ATTRIBUTE, value>

An example datum is

<John Q., Public, SSN = 123-45-6789>

Note that data, as defined above are abstract. One never actually sees data. Instead, what we actually see (and computers process) are *data records* on paper, *via* computer applications, in tables or charts, and so on. Finally, note that the same datum may be recorded and presented in a virtually unlimited number of formats.

For the purposes of this entry, we will define information as “the portion of a collection of data (or, more generally a ‘signal’) used in a specific decision-making process”. As an example, when the datum

<Consumer Price Index, 4th QUARTER 2006,
0.7%>

is used (with other data) to adjust social security payments, it is information. In other contexts it is not. This definition adheres closely enough to the formal approaches used in communications, thermodynamics, and statistics (Resnikoff [4]) and to the

less formal way in which many business people use the term.

Measurement

Every quality practitioner knows that one cannot manage what one cannot measure. One can measure the length of a manufactured product, the viscosity of a liquid, the impedance of an electrical component, and so forth. “Length”, “viscosity”, and “impedance” are all physical properties and all can be measured directly. But data and information are intangible, so they have no physical properties. Hence, no direct measurement.

There are quite literally, hundreds of “dimensions” of data and information quality (Wang and Strong [5]) ranging from relevancy to comprehensiveness, to ease of interpretation. Here, we will focus on *accuracy* because it is critical to almost all data customers and because it illustrates the approaches to dealing with the lack of physical properties. We define accuracy as “the degree to which data and information agree with a standard (usually but not always the real world) assumed to be correct” [2].

Thus, even though direct measurement of accuracy is impossible (i.e., there is no “accurometer”), practitioners have devised a number of ways to get around this problem (Redman [6]). The first involves comparing data to their real-world counterparts. For example, a company may contact customers to check that the addresses in its database are correct. This approach is powerful, but comparing data to the real world is time consuming and expensive. So sampling must be employed to make the job manageable.

The second method is called *data tracking*. It involves “tracking” a sample of data as the selected records wind their way across a process, and noting unexpected changes in data values or formats. Figure 1 below is an example of a tracked record. The Figure also features several such unexpected changes, which usually correspond to data errors. “Counts” of the fraction of records make it across the process correctly and types of errors advise the process owner how well the process is performing and point out the likely sources of error.

The most popular method to measure accuracy takes advantage of the aforementioned domain of allowed values. Properly defined, a data value may only take on certain values. For example, there are

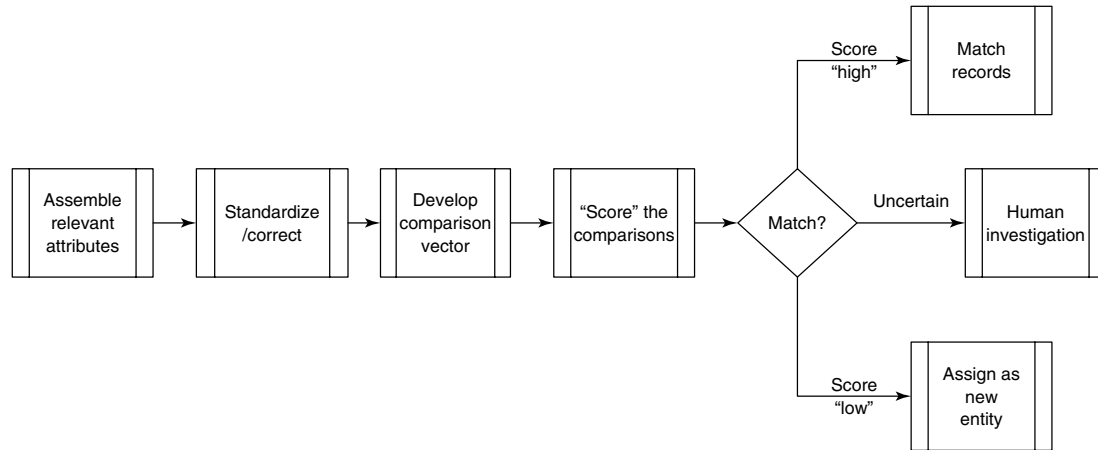


Figure 1 The matching process

only a few possible values of the attribute SEX. Values outside the domain of allowed values cannot be correct. Importantly, the domain of allowed values may be large and complex, but virtually all properly defined data are so constrained.

This point extends to multiple data elements. Consider the pair

<STREET ADDRESS = 12Main Street>
<CITY = Any town>

Not every town has a Main Street, so offending pairs are easy to identify. Adding STATE, POSTAL CODE, and other data elements increases the potential for error detection.

So-called business rules (sometimes called *edits*) specify the domains of allowed data values and failures of business rules that correspond to erred data. Counting the failures often leads to useful measurements of accuracy. The method has two important advantages: first, there are a whole host of good business rules and methods for developing new ones (Ross [8] and Olson [9]). Second, there are plenty of good off-the-shelf software programs to execute them. The principal shortcomings are that erred data need not fail a business rule and that additional work is needed to determine which data elements are erred.

Error Correction

Another way of looking at business rules is that they stem from redundancies that occur in even the

simplest collection of data. As noted above, this redundancy makes it easier to find errors and in many cases, to fix them. As an example, there is no place with the address

<12 Monmouth Ave., Rumson, NJ, USA, 07706>

so the record will fail good business rules. But there is a

<12 Monmouth Avenue, Rumson, NJ, USA, 07760>

Note two differences: first, Ave. and Avenue. While “Ave.” is not really incorrect, changing it to “Avenue”, is an example of “data standardization”. The second difference is that the last two digits in the POSTAL CODE are transposed. It is the likely error. Of course, other attributes, even the entire record, may be erred. So error correction is essentially a probabilistic exercise.

Data standardization and error correction can be quite sophisticated when there are lots of data and good business rules. ADDRESS is probably the best example. Generally though, business rules are not defined clearly and one must resort to the real world to determine the correct data values.

Interestingly, in telecommunications, algorithms that digitize and transmit data (including voice conversations) use parity bits to ensure faithful end-to-end transmission. More sophisticated error-correcting algorithms take the concept a step further. They “build-in” redundancy to provide error correction (Pierce and Noll [10]). Except in the design of some

identifiers (attributes used to uniquely identify individuals), these concepts are yet to be fully exploited in data and information quality.

Statistical Control

The twin notions of interchangeable parts and reducing variation are the heart and soul of the industrial revolutions and quality management. Further, since its discovery by Shewhart [11], statistical control has proved unmatched in helping reduce variation by helping sort out sources of variation:

- inherent to the process (common causes)
- external to it (special causes).

But the notion of interchangeable parts simply does not apply to data. Quite the opposite. “Interesting” data are novel and different. No one bothers to include an attribute PLANET OF RESIDENCE in a customer database, as the value will be the same for each individual (today at least). So it is not at all clear whether statistical control will apply.

Three insights help resolve the problem. First, the insight that a more constructive way to view data is not “as the stuff stored away in corporate and other databases”, but as “organic”, coming into existence, moving around the organization and being modified as it goes, being stored and processed, and, from time to time being put to some value-creating use (Levitin and Redman [12]). Second, the insight that “control” is best applied to *process*, and it matters not whether a process produces components, final assemblies, or data. And finally, the insight that accuracy measurements, made in-process, lend themselves beautifully to so-called *p* charts (see **Control Charts, Overview; Control Charts for the Mean**).

Inconsistency in Data and Information

Data are inconsistent when two or more sets of data do not agree. For example if Sally Jones’ primary residence is given as

<STREET ADDRESS = 10 Main Street>
in one database and
<STREET ADDRESS = 14 Maple Avenue>

in a second, these data are inconsistent. Both cannot be correct. This issue is exactly the same as noted above in which a pair of attributes within a data record could not simultaneously be correct.

Inconsistencies arise for many reasons. One is that it is easy to copy data from one place to another, but very difficult to synchronize the copies when changes occur. A second is that each department (even person!) can develop its own database, for its own reasons and processes for keeping up with changes. Synchronization is near impossible!

Note that there is no analog to data inconsistency in manufacturing, and it is an important and expensive problem. Billions are wasted sorting out inconsistencies. And everyone can cite examples where important decisions were delayed because the “facts” did not agree.

Of course, good process design to keep data values current helps. So does “concurrency control”, an important technique in database management systems (Date [13]). Also, data tracking, as described, has proved effective at sniffing out inconsistencies as data wind their way across processes, but overall, the body of theory and technique for addressing inconsistency is neither powerful enough nor consistently applied.

Unique Identification and Matching

One of the most important dimensions of data quality and principles of database design is “identifiability”, the notion that each entity should be uniquely labeled. EMPLOYEE ID and SOCIAL SECURITY NUMBER represent attempts to establish unique identifiers. An especially virulent form of inconsistency arises when identifiability is not met. It is sometimes called the *matching problem*.

The problem arises in many ways. Determining whether a person is on any “terrorist watch lists” is one example. Another involves the sequential mergers of four companies into one. The new company needs to understand the totality of its relationships with customers and so wants to merge the contents of customer databases in the former companies. As an example, four customer names are:

Former company 1: IBM
Former company 2: Intl Bus Mach
Former company 3: International Business Machines

Former company 4: Intergalactic Bell Modules.

The question is “Which of these records indicate the same customer?”

Statisticians have worked on this problem for decades and, as computing power has grown, there have been many improvements (Bishop *et al.* [14], Winkler and Thibaudeau [15], and Loshin [16]). Methods for answering the question proceed as in Figure 2.

Step 1: Assemble relevant attributes: For matching people, HOME ADDRESS, AGE, and SEX are useful. None applies for matching companies, though TAX ID and CEO may prove useful.

Step 2: Standardize/correct: To the degree possible, standardize data values and make corrections, as described above.

Step 3: Assemble a comparison vector. Determine the degree to which attributes match. Exact matches are easy. One may assign a “partial match” to values that are close, but not exact matches, such as <LAST NAME = Shackelford > and <LAST NAME = Shackleford>.

Step 4: Score the comparison vector. Doing so usually involves assigning weights to each attribute and summing to obtain an overall “score”.

Step 5: Make matches/assign new entities/human

investigation. On the basis of the scores, one now matches data records with a high enough score. Similarly, unmatched records are recognized as unique entities. Finally, some data records may fall into a middle area, where further investigation by human beings is warranted.

Privacy and Security

The last couple of years have witnessed a spate of “data breaches”, in which personal and other private data have been compromised [17]. While protecting physical assets is tough enough, protecting data and information may be even more problematic. For, as we noted above, since data can be digitized they can be easily copied. This means that, for example, people can put copies on their laptops and take them home. It also means that to “steal data”, criminals need not actually take the data. They can make a copy and take it instead. And those in charge of protecting the data will not detect any change.

Final Remarks

The statistical foundations of quality management have proved remarkably fungible to data and information quality, but not without effort. For data

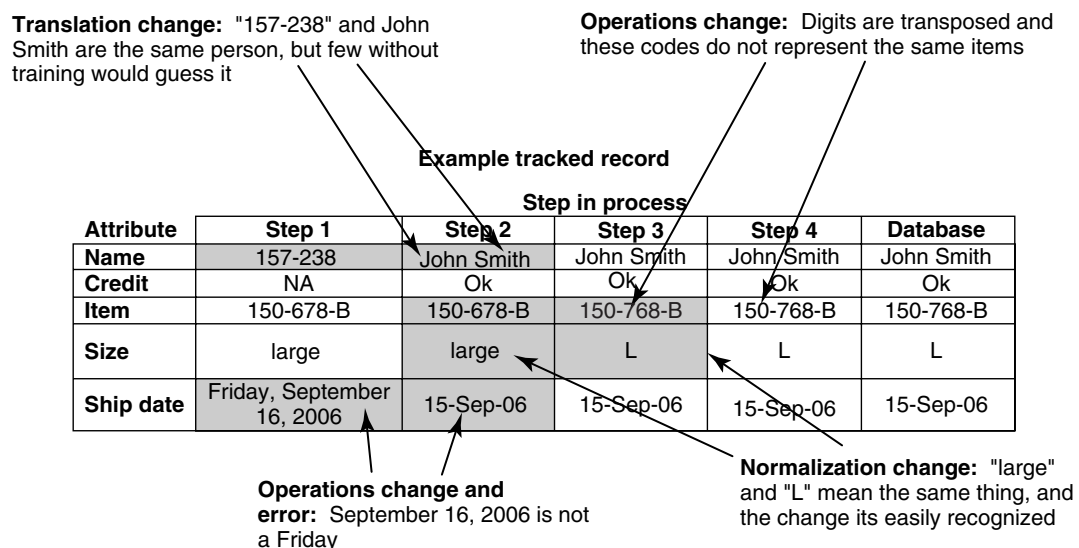


Figure 2 An example of a tracked record. Each of the shaded pairs of attributes highlights a change in the tracked record as it winds its way through the process. From Redman [7]

and information have no physical properties, complicating measurement; the most important data are novel and unique, complicating statistical control; it is extremely difficult to properly identify “entities” uniquely, complicating management, and a “bad guy” need not physically remove data and information to steal them, complicating privacy and security. On the other hand, data and information are often highly redundant, making it easier to identify, and sometimes correct errors.

Acknowledgments

The author would like to thank Frank Guess of the University of Tennessee and David Loshin of Knowledge Integrity for discussions that improved this article.

References

- [1] Taggart, S. (1999). The 20-Ton packet, *Wired* 7.
- [2] Fox, C.J., Levitin, A.V. & Redman, T.C. (1994). The notion of data and its quality dimensions, *Information Processing and Management* 30, 9–19.
- [3] Tsichritzis, D.C. & Lochovsky, F.H. (1982). *Data Models*, Prentice-Hall, Englewood Cliffs.
- [4] Resnikoff, H.L. (1987). *The Illusion of Reality*, Springer-Verlag, New York.
- [5] Wang, R.Y. & Strong, D.M. (1996). Beyond accuracy: what data quality means to consumers, *Journal of Management Information Systems* 4, 5–34.
- [6] Redman, T.C. (2005). Measuring data accuracy: A framework and review, in *Information Quality, Advances in Management Information Systems Series*, M. E. Sharpe, New York, pp. 21–36.
- [7] Redman, T.C. (2001). *Data Quality: The Field Guide*, Butterworth-Heinemann, Boston.
- [8] Ross, R.G. (2003). *Principles of the Business Rule Approach*, Addison-Wesley, Boston.
- [9] Olson, J.E. (2003). *Data Quality: The Accuracy Dimension*, Morgan Kaufmann Publishers, Amsterdam.
- [10] Pierce, J.R. & Noll, A.M. (1990). *Signals: The Science of Telecommunications*, Scientific American Library, New York.
- [11] Shewhart, W.A. (1986). *Statistical Method from the Viewpoint of Quality Control*, Dover Publications, New York. Original: Graduate School of the Department of Agriculture, Washington: 1939.
- [12] Levitin, A.V. & Redman, T.C. (1993). A model of data (life) cycles with applications to quality, *Information and Software Technology* 35, 217–224.
- [13] Date, C.J. (1987). *An Introduction to Database Systems*, Addison-Wesley, Reading, Vol. 1.
- [14] Bishop, Y.M.M., Feinberg, S.E. & Holland, P.W. (1975). *Discrete Multivariate Analysis*, MIT Press, Cambridge.
- [15] Winkler, W.E. & Thibaudeau, Y. (1991). *An Application of the Fellegi-Sumter Model on Record Linkage to the 1990 U.S. Decennial Census, Statistical Research Report Series RR91/09*, U.S. Bureau of the Census, Washington, D.C.
- [16] Loshin, D. (2000). Value-added data: merge ahead, *Intelligent Enterprise* 3.
- [17] Privacy Rights Clearinghouse. *A chronology of data breaches*, <http://www.privacyrights.org/ar/ChronData-Breaches.htm> Updated 29 November 2006.

THOMAS C. REDMAN

Interlaboratory Comparisons

Introduction

An interlaboratory comparison for a measurement procedure is an exercise carried out by a group of laboratories to compare their performance or assess a measurement standard. The basic approach for carrying out this type of comparison, also commonly referred to as a round-robin comparison, is to circulate an artifact to be measured among the laboratories, collect the results from each laboratory, compare the results from the group of laboratories, and report the results to the participants. As usual, however, there are many important details to be considered when planning an interlaboratory comparison or analyzing interlaboratory data.

The details to be considered start with the exact goals of the comparison. Interlaboratory comparisons are typically used for one of three main purposes:

1. to assess the random variation in measurement results across a population of laboratories,
2. to determine the systematic differences in results among a fixed set of laboratories, or
3. to determine the value of a physical property of an artifact or a population of artifacts.

On the basis of the purpose and results of the comparison, different conclusions or actions may be taken. For example, if the primary goal is to assess the variation in the results of a standard measurement procedure and the variation is found to be reasonably low, then the results may be used to establish an uncertainty statement that will become part of the published measurement procedure. If, on the other hand, the variation between laboratories is higher than acceptable, the details of the measurement standard or the individual measurement processes used at each laboratory may be revisited to learn how they can be improved.

Design of Interlaboratory Comparisons

Naturally, the goals of an interlaboratory comparison will impact its design. For example, if the goal is to assess the between-laboratory variation in a

population of laboratories, then a random sample of laboratories from the relevant population will be needed. When a random sample of laboratories is not available, a representative sample that can be treated as random is often used as a surrogate. When the differences in the results from a fixed set of laboratories are of interest then there is no need to randomly select which laboratories will participate in the comparison.

In a study whose goal is to determine the value of a physical property of an artifact, a randomly selected or fixed set of laboratories may be used. For example, in the production of reference materials for measurement assurance or **calibration** or in the analysis of environmental samples a random sample of typical laboratories may be selected to determine a consensus value and associated uncertainty for the artifact property [1, 2]. Alternatively, two or three laboratories that use measurement methods based on different underlying scientific principles that have biases that can be bounded [3] or can be assumed to bracket the artifact value [4] may be chosen.

The cost and stability of the artifacts and the time frame for the comparison will also affect the design. In a simple, circular design a single artifact is circulated, as illustrated in Figure 1(a). This design has the advantage that the results from each laboratory are directly comparable, but requires a stable artifact whose value will not change as it travels. Use of a circular design is also relatively time consuming since the laboratories cannot make measurements in parallel.

If the stability of the artifact is a concern, a variation on a figure-eight design, shown in Figure 1(b), can be used. This design still uses one artifact, but it is remeasured by the organizing laboratory (Lab 1 in this case) about halfway through the exercise to check for drift or other changes to the artifact's value. Depending on how much the value of the artifact has changed, the design of the comparison can be modified or appropriate corrections can be made in the analysis of the data.

If the use of multiple artifacts is feasible, then the figure-eight design can be run in parallel, as demonstrated in Figure 1(c) for artifacts A1 and A2. This design can be used to assess artifact drift like the single-artifact figure-eight design, but may require the use of the results from Lab 1 to make comparisons among labs that measured different artifacts, if the artifacts cannot be assumed to be identical. The

2 Interlaboratory Comparisons

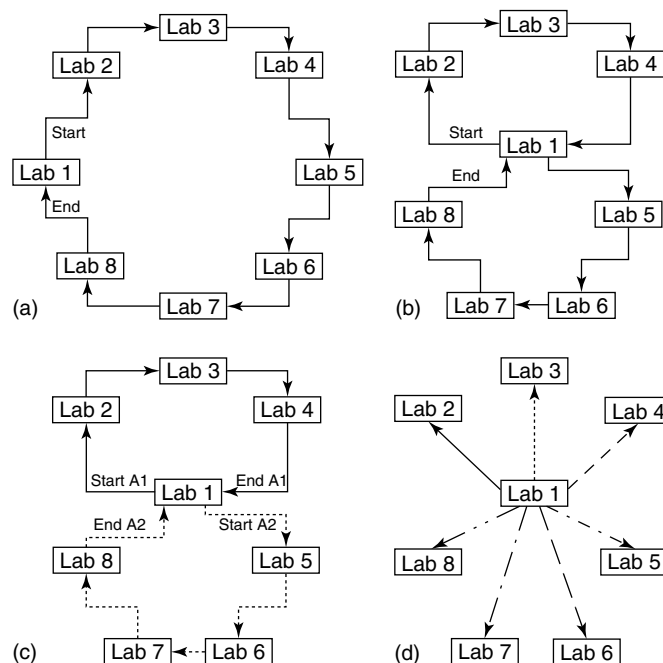


Figure 1 High-level experiment designs for interlaboratory comparisons. (a) A circular design, (b) a figure-eight design with one artifact, (c) a figure-eight design with artifacts A1 and A2, (d) a star design that requires an artifact for each participant

figure-eight design can be extended to include more loops to assess artifact drift more frequently or shorten the study by increasing parallelism. In the limit, each laboratory may receive its own artifact to measure, as shown in the star-shaped design in Figure 1(d). This design is used frequently in comparisons of chemical properties of solutions, where a master solution can easily be divided into aliquots with essentially identical properties.

Finally, as for most experiment designs, the models hypothesized to describe the data and the properties of the analysis methods will also have a major impact. For example, a very simple and useful graphical analysis of interlaboratory data is described in a group of papers published by Youden starting in the 1950s [5–8]. For this analysis two artifacts with similar, but not identical, values are circulated. Each laboratory measures each artifact once, or provides a prespecified number of measurements for each that can be summarized by the organizer to obtain a pair of results. The decision to use a pair of artifacts is convenient for the graphical analysis of the data (discussed in the next section), provides data that can

be used to assess the random variation within and between laboratories, and also provides a reasonable level of blinding to reduce screening of the results prior to reporting.

A data set for the residual cement caught in a 45- μm sieve using Youden's two-artifact design is given in Table 1. The proportion of fine particles in a cement formulation is an important parameter that affects its physical properties and is monitored for quality control. ASTM standard test method C430 [9] is used for the determination of this cement property.

As for the graphical analysis, to use a hierarchical, analysis of variance (**ANOVA**)-like model to describe the data the design of the comparison must incorporate the necessary features. In particular, replicate measurements at each laboratory are likely to be needed and it is desirable to have each laboratory make an equal number of measurements. If additional sources of variation other than between- and within-laboratory sources are important then each laboratory may need to make measurements on a prescribed number of days, with a particular number of operators, and so on.

Table 1 Data from an interlaboratory comparison to assess the performance of ASTM test method C430 for the residual cement left on a 45- μm sieve

Lab	Residue mass fraction (%)	
	Cement 1	Cement 2
1	8.838	7.039
2	7.574	5.927
3	8.677	7.254
4	8.494	7.047
5	6.760	5.191
6	7.545	6.209
7	8.434	6.749
8	7.985	6.401
9	9.365	7.871
10	8.351	6.870
11	9.639	7.806
12	9.311	7.723
13	8.065	6.603
14	8.119	6.449

Of course, while it is natural, and potentially necessary, to pin down these design details when planning an interlaboratory study, it is also important to keep in mind their impact on the scope of the results. One of the goals of most comparisons is to reflect typical conditions, thus a good design must accommodate trade-offs between ideal statistical properties and realistic day-to-day testing practices.

Analysis of Interlaboratory Data

As mentioned above, Youden's graphical analysis [5–8] of artifact pairs is easy to implement, intuitive, and often successfully communicates the essential information in the data for comparisons designed to assess the variation between laboratories. The graphical analysis is carried out by simply making a scatter plot of the pairs of results for each laboratory, using the means of the measurements on each artifact if replicate measurements were made. The plot is then augmented with median lines for each artifact and a 45° reference line (relative to the intersection of the median lines) to help the analyst interpret how well the results from different laboratories agree. In some cases a circle whose radius is chosen to encompass a specified fraction of laboratories on average, under the assumption that there is no between-laboratory variation, is also added to the plot to aid interpretation

of the results. Figure 2 shows a Youden plot for the cement data from Table 1 with a circle sized to encompass about 95% of the laboratories, if there were no significant variation between laboratories.

Although Youden's analysis does not rely on many formal statistical inferences, there is a wealth of information in this simple plot about how well the laboratories agree and how their performance can be improved. For example, points that are arrayed together in an elliptical pattern along the 45° line rather than being randomly scattered inside the circle, as in Figure 2, indicate the presence of systematic differences in how the laboratories are implementing the measurement procedure. Points along the 45° line that are well away from their peers, like the lower left point in Figure 2, indicate laboratories that may not understand the standard procedure or have serious difficulties with its implementation. Points that fall well outside the circle but away from the 45° line identify laboratories that may have larger within-laboratory variation than their peers or may not be in statistical control.

The results obtained from random-effects hierarchical models, fit using ANOVA or Bayesian methods, allow the analyst to assess the between-laboratory, within-laboratory, and potentially other,

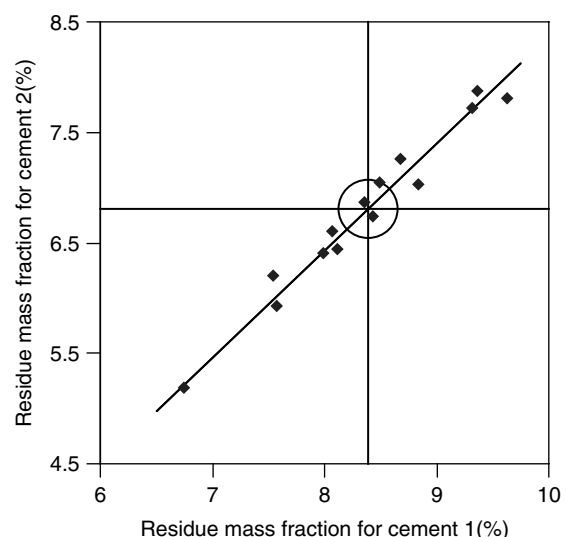


Figure 2 A Youden plot of the cement data from Table 1 showing large variation between laboratories relative to the within-laboratory variation

4 Interlaboratory Comparisons

sources of random variation associated with a particular test method. In addition to standard statistics texts that discuss these types of analyses, a description of a one-way, random-effects ANOVA specific to interlaboratory comparisons is given in ASTM standard practice E691 [10]. Approximate **confidence intervals** for the between-laboratory and within-laboratory **variance components** and their sum can also be obtained by building on basic ANOVA results or *via* maximum likelihood [11, 12].

Other more specialized models and methods extend the results of these workhorse analysis methods. These include the use of robust statistical methods to better handle data that contains **outliers** [13, 14] and models to simultaneously assess a measurement method over a range of analyte levels [2, 15, 16].

Probabilistic intervals for pairwise differences of between laboratories are often used to assess the agreement among a specific set of laboratories. In many cases, when the difference between one specific pair of laboratories is of interest, use of a simple confidence interval based on Student's *t* distribution (or other analogous probabilistic interval) is often most appropriate. When the goal is to make several simultaneous comparisons between laboratories with a stated error rate, however, procedures for multiple comparisons must be used [17].

When analyzing data from an interlaboratory comparison with a random-effects model to determine a consensus mean for the value of an artifact, there are many possible methods ranging from simple to complex. For example, in cases in which the data comes from a random sample of laboratories each with the same level of within-laboratory measurement variation and each about the same number of results, the consensus value and its uncertainty can simply be based on a standard Student's *t* distribution confidence interval computed using the laboratory means. This simple interval will also work well when the results from each laboratory are not equally precise as long as the between-laboratory variation dominates the within-laboratory variation in the least precise result. When the individual laboratory results do not all have the same level of within-laboratory variation and the between-laboratory variation does not dominate, however, other methods will give better results. As long as results are available from a relatively large number of laboratories, the maximum-likelihood methods described in [18–20] work well.

Iyer, Wang, and Mathew present a **Monte Carlo** method using generalized pivot quantities that gives results similar to the maximum-likelihood results and works reasonably well even for small numbers of laboratories [21]. The same general method can be used to estimate artifact values in comparisons based on fixed sets of laboratories as well. The authors compare confidence intervals based on generalized pivot quantities with those based on the minimax approach in [3] that uses known bounds on the laboratory biases. They also comment on the intervals used in [4] that assume the laboratory biases bound the true artifact value.

References

- [1] (2006). *Reference Materials – General and Statistical Principles for Certification*, International Organization for Standardization, Geneva, pp. 30–47, 59–62.
- [2] Bhaumik, D. & Gibbons, R. (2005). Confidence regions for random-effects calibration curves with heteroscedastic errors, *Technometrics* **47**, 223–230.
- [3] Eberhardt, K., Reeve, C. & Spiegelman, C. (1989). A minimax approach to combining means with practical examples, *Chemometrics and Intelligent Laboratory Systems* **5**, 129–154.
- [4] Levenson, M., Banks, D. & Eberhardt, K., Gill L., Guthria W., Liu H.-K., Vangel M., Yen J., Zhang N.-F. (2000). An approach to combining results from multiple methods motivated by the ISO GUM, *Journal of Research of the National Institute of Standards and Technology* **105**, 571–579.
- [5] Youden, W. (1959). Graphical diagnosis of interlaboratory test results, *Industrial Quality Control* **15**, 24–28.
- [6] Youden, W. (1959). Evaluation of chemical analyses on two rocks, *Technometrics* **1**, 409–417.
- [7] Youden, W. (1960). The sample, the procedure, and the laboratory, *Analytical Chemistry* **32**, 23–37.
- [8] Youden, W. (1975). Statistical techniques for collaborative tests, in *Statistical Manual of the Association of Official Analytical Chemists*, W. Youden & E. Steiner, eds, Association of Official Analytical Chemists, Gaithersburg, pp. 1–63.
- [9] *Standard Test Method for Fineness of Hydraulic Cement by the 45- μ m (no. 325) Sieve C430-96*, (2003). ASTM International, West Conshohocken, pp. 1–3.
- [10] *Standard Practice for Conducting an Interlaboratory Study to Determine the Precision of a Test Method E 691-05*, (2003). ASTM International, West Conshohocken, pp. 1–23.
- [11] Burdick, R. & Graybill, F. (1992). *Confidence Intervals on Variance Components*, Marcel Dekker, New York, pp. 58–116.

-
- [12] Crowder, M. (1992). Interlaboratory comparisons: round robins with random effects, *Applied Statistics* **41**, 409–425.
- [13] Rocke, D. (1983). Robust statistical analysis of interlaboratory studies, *Biometrika* **70**, 421–483.
- [14] Lischer, P. (1996). Robust statistical methods in interlaboratory analytical studies. in *Robust Statistics, Data Analysis, and Computer Intensive Methods, Lecture Notes in Statistics*, H. Rieder, ed, Springer-Verlag, New York, pp. 251–265.
- [15] Mandel, J. & Lashof, T. (1959). The interlaboratory evaluation of testing methods, *ASTM Bulletin* **239**, 53–61.
- [16] Gibbons, R. & Bhaumik, D. (2001). Weighted random-effects regression models with application to interlaboratory calibration, *Technometrics* **43**, 192–198.
- [17] Miller, R. (1981). *Simultaneous Statistical Inference*, Springer-Verlag, New York, pp. 37–108.
- [18] Rukhin, A. & Vangel, M. (1998). Estimation of a common mean and weighted mean statistics, *Journal of the American Statistical Association* **93**, 303–308.
- [19] Rukhin, A. & Vangel, M. (1999). Maximum likelihood analysis for heteroscedastic one-way random effects ANOVA in interlaboratory studies, *Biometrics* **55**, 129–136.
- [20] Rukhin, A., Biggerstaff, B. & Vangel, M. (2000). Restricted maximum likelihood estimation of a common mean and the Mandel-Paule algorithm, *Journal of Statistical Planning and Inference* **83**, 319–330.
- [21] Iyer, H., Wang, J. & Mathew, T. (2004). Models and confidence intervals for true values in interlaboratory trials, *Journal of the American Statistical Association* **99**, 1060–1071.

WILLIAM F. GUTHRIE

Quadrature and Numerical Integration

Introduction

A numerical integration technique is an algorithm to calculate the numerical value of a definite integral,

$$\int_a^b f(x) dx \quad (1)$$

Numerical integration problems go back at least to Greek antiquity when e.g. the area of a circle was obtained by successively increasing the number of sides of an inscribed polygon. In the seventeenth century, the invention of calculus resulted in a new development of the subject, leading to the basic numerical integration rules. In the following centuries, the field became more sophisticated and, with the introduction of computers in the recent past, many classical and new algorithms have been implemented, leading to very fast and accurate results. Consequently, there is a vast amount of relevant literature on solving numerical integration problems consisting of books, articles, software packages, etc. Some essential references are [1–6].

Numerical integration is sometimes called *quadrature*. A quadrature rule is a numerical approximation of an integral using a weighted sum of some values of the integrand,

$$\int_a^b f(x) dx \approx \sum_{i=1}^n \omega_i f(x_i) \quad (2)$$

where $x_1, \dots, x_n$ are the points or abscissas, usually chosen to be in the interval of integration, and $\omega_1, \dots, \omega_n$ are the weights associated to these points.

The simplest quadrature rules are the Riemann sums. Given a set of ordered abscissas, $a = x_1 < x_2 < \dots < x_n = b$, the left-handed Riemann sum is obtained with,

$$\int_a^b f(x) dx \approx \sum_{i=1}^{n-1} h_i f(x_i) \quad (3)$$

where $h_i = (x_{i+1} - x_i)$, the length of the i th interval between points. Observe that the area under the curve in the i th interval is approximated by a rectangle

of width h_i and height $f(x_i)$. Analogously, one can define the right-handed and midpoint Riemann sums by evaluating the function on the maximum or midpoint value, respectively, of each i th interval. The Riemann sums are also known as *rectangular rules* because of the use of rectangles to approximate the integral.

The rest of this paper is organized as follows. The section titled “Newton–Cotes Formulas” describes the Newton–Cotes formulas which are based on evaluating the integrand at a number of equally spaced points. The section titled “Gaussian Quadrature” introduces the Gaussian quadrature rules where the abscissas are chosen optimally to give the most accurate approximations possible. This section also includes a numerical example comparing different approaches. Some comments on the extensions of the Gaussian quadrature are also included. The section titled “Discussion and Remarks” concludes with some discussion and remarks.

Newton–Cotes Formulas

Assume that we divide the interval $[a, b]$ into n equally spaced points such that,

$$x_i = x_1 + ih, \quad \text{for } i = 0, 1, \dots, n-1 \quad (4)$$

where $x_1 = a$, $x_n = b$ and $h = (b - a)/n$. Let us denote $f_i = f(x_i)$, for $i = 1, \dots, n$. The main idea in Newton–Cotes formulas is to approximate the function $f(x)$ using polynomials that pass through the points (x_i, f_i) and integrate them to approximate the integral from x_1 to x_n . Newton–Cotes formulas are of *closed type* when the endpoints, x_1 and x_n , are included in the set of abscissas and they are called *open type* when the endpoints are not used to approximate the integral.

The basic closed Newton–Cotes formula is the *trapezoidal rule* which uses two points, x_1 and x_2 , and approximates the function $f(x)$ on $[x_1, x_2]$ using a straight line joining (x_1, f_1) and (x_2, f_2) . Then, the function is approximated by the following first-order polynomial,

$$f(x) \approx \left(\frac{f_2 - f_1}{h} \right) x + \frac{x_2 f_1 - x_1 f_2}{h} \quad (5)$$

2 Quadrature and Numerical Integration

This polynomial can be integrated to give the following approximation,

$$\int_{x_1}^{x_2} f(x) dx \approx \int_{x_2}^{x_2} \frac{f_2 - f_1}{h} x + \frac{x_2 f_1 - x_1 f_2}{h} dx$$

$$dx = \frac{h}{2}(f_1 + f_2) \quad (6)$$

Observe that this is called the trapezoidal rule because the area under $f(x)$ is approximated by a trapezoid with bases f_1 and f_2 and height $h = (x_2 - x_1)$, see Figure 1. The error of this approximation can be shown to be given by, see e.g. [2],

$$E_T = -\frac{h^3}{12} f''(\xi) \quad (7)$$

where ξ is some point in the interval $[x_1, x_2]$. Then, the amount of error will not be larger than the maximum value of the second derivative of the function in the interval and it will increase with the curvature of the function. The trapezoidal rule is exact for polynomials up to and including degree 1.

The trapezoidal rule can be used repeatedly to approximate the integral over the whole interval, $[x_1, x_n]$. This is called the extended or composite trapezoidal rule and is given by,

$$\int_{x_1}^{x_n} f(x) dx = \int_{x_1}^{x_2} f(x) dx + \dots + \int_{x_{n-1}}^{x_n} f(x) dx$$

$$\approx h \left(\frac{f_1}{2} + f_2 + \dots + f_{n-1} + \frac{f_n}{2} \right) \quad (8)$$

whose approximation error is equal to $-(n-1)h^3 f''(\xi)/12$. A powerful extension of the trapezoidal

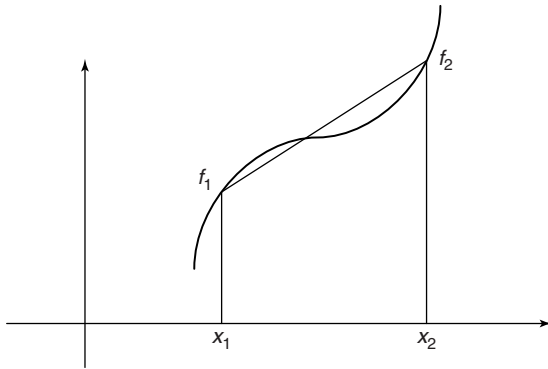


Figure 1 Trapezoidal rule

rule is the Romberg integration, see e.g. [5], which uses a sequence of refinements of the extended trapezoidal rule by increasing carefully the number of subintervals.

The three-point closed Newton–Cotes formula is called the *Simpson's rule* and uses a quadratic polynomial joining the points (x_1, f_1) , (x_2, f_2) and (x_3, f_3) to approximate the function $f(x)$ on the interval $[x_1, x_3]$,

$$f(x) \approx \alpha x^2 + \beta x + \gamma \quad (9)$$

where the coefficients α , β and γ are given by,

$$\alpha = \frac{f_1 - 2f_2 + f_3}{2h^2} \quad (10)$$

$$\beta = \frac{-2x_1(f_1 - 2f_2 + f_3) - h(3f_1 - 4f_2 + f_3)}{2h^2} \quad (11)$$

$$\gamma = \frac{x_1^2(f_1 - 2f_2 + f_3) + hx_1(3f_1 - 4f_2 + f_3) + 2h^2 f_1}{2h^2} \quad (12)$$

Then, we can integrate this second-order polynomial to obtain the Simpson's approximation,

$$\int_{x_1}^{x_3} f(x) dx \approx \frac{h}{3}(f_1 + 4f_2 + f_3) \quad (13)$$

The error in this case can be shown to be given by,

$$E_S = -\frac{h^5}{90} f^{(iv)}(\xi) \quad (14)$$

where $f^{(iv)}(\xi)$ is the fourth derivative of $f(x)$ evaluated at some point ξ such that $x_1 \leq \xi \leq x_3$. Simpson's rule is exact for polynomials up to and including degree 3. This is surprising because it is designed to be exact for quadratic polynomials but, due to the symmetry of the formula, it is also exact for cubic polynomials.

As before, we can use Simpson's rule repeatedly to obtain the following extended or composite Simpson's rule, where the number of points, n , must be odd,

$$\int_{x_1}^{x_n} f(x) dx = \int_{x_1}^{x_3} f(x) dx + \dots + \int_{x_{n-2}}^{x_n} f(x) dx$$

$$\approx \frac{h}{3}(f_1 + 4f_2 + 2f_3 + 4f_4$$

$$+ 2f_5 \dots + 4f_{n-1} + f_n) \quad (15)$$

The approximation error is equal to $-(n-1)h^5 f^{(iv)}(\xi)/180$. Note that Simpson's rule will in general imply better accuracy than the trapezoidal rule if the function $f(x)$ is smooth enough (with finite fourth derivative).

The four-point closed Newton–Cotes formula is called *Simpson's 3/8 rule* and is based on a cubic polynomial approximation,

$$f(x) \approx \alpha x^3 + \beta x^2 + \gamma x + \delta \quad (16)$$

where α , β , γ , and δ are constants such that the polynomial passes through the points $(x_1, f_1), \dots, (x_4, f_4)$. These coefficients can be calculated using the Lagrange interpolation formulas, see e.g. [2]. Then, the Simpson's 3/8 rule can be obtained to be,

$$\int_{x_1}^{x_3} f(x) dx \approx \frac{3h}{8}(f_1 + 3f_2 + 3f_3 + f_4) \quad (17)$$

whose approximation error is given by,

$$E_{\bar{S}} = -\frac{3h^5}{80} f^{(iv)}(\xi) \quad (18)$$

As expected, Simpson's 3/8 rule is exact for polynomials up to and including degree 3. Note that the approximation error for Simpson's 3/8 rule is larger than the ordinary Simpson's error given in equation (14). However, Simpson's 3/8 rule has the advantage that the number of points, n , is not required to be odd when it is used repeatedly to approximate the integral over the whole interval, $[x_1, x_n]$. Then, if n is even, one possibility is using the 3/8 Simpson's rule for the first four points and the ordinary Simpson's rule for the remaining points, whose number (including the fourth point) is definitely odd.

The five-point closed Newton–Cotes formula is called the *Boole's rule* and is given by,

$$\int_{x_1}^{x_5} f(x) dx = \frac{2h}{45}(7f_1 + 32f_2 + 12f_3 + 32f_4 + 7f_5) \quad (19)$$

which can be obtained using the fourth-order Lagrange polynomial that passes through $(x_1, f_1), \dots, (x_5, f_5)$. The error is $-8h^7 f^{(vi)}(\xi)/945$ and it is exact for polynomials of degree 5 or less. Boole's rule is also known as *Bode's rule*, as in e.g. [7], due to an early typo.

Further point closed Newton–Cotes formulas can be obtained by integrating the Lagrange polynomials that pass through the considered points, see e.g. [8].

In general, observe that higher-point formulas will imply higher accuracy only when the function can be well approximated by a polynomial.

Open Newton–Cotes formulas approximate the integral from $a = x_1$ to $b = x_n$ without using the interior points, $x_2, x_2, \dots, x_{n-1}$. Some of these formulas are shown in Table 1. Newton–Cotes formulas of open type can be useful, for example, when the function takes infinite values at the endpoints or when one or both endpoints are infinity. However, these are not very common in practice as it is not reasonable to use them repeatedly to obtain extended open rules. Also, the Gaussian quadrature rules, which will be described in the next section, lead always to more accurate approximations than open Newton–Cotes formulas.

Gaussian Quadrature

Suppose now that we have the freedom to choose the abscissas at which to evaluate the function $f(x)$ rather than being equally spaced points. The Gaussian quadrature rules provide careful choices of these points in order to obtain much more accuracy in approximating the required integral. These are open formulas in the sense that it is not required to evaluate the function at the endpoints. Furthermore, it is possible to define a weighting function, $W(x)$, such that the approximation of the integral will be exact for polynomials times this weight function instead of only for polynomials. Then, the numerical

Table 1 Newton–Cotes formulas of open type

Approximation formula	Error
$\int_{x_1}^{x_3} f(x) dx \approx 2hf_2$	$\frac{h^3}{24} f''(\xi)$
$\int_{x_1}^{x_4} f(x) dx \approx \frac{3h}{2}(f_2 + f_3)$	$\frac{h^3}{4} f''(\xi)$
$\int_{x_1}^{x_5} f(x) dx \approx \frac{4h}{3}(2f_2 - f_3 + 2f_4)$	$\frac{28h^5}{90} f^{(iv)}(\xi)$
$\int_{x_1}^{x_6} f(x) dx \approx \frac{5h}{24} \times (11f_2 + f_3 + f_4 + 11f_5)$	$\frac{95h^5}{144} f^{(iv)}(\xi)$
$\int_{x_1}^{x_7} f(x) dx \approx \frac{6h}{20}(11f_2 - 14f_3 + 26f_4 - 14f_5 + 11f_6)$	$\frac{-41h^7}{140} f^{(vi)}(\xi)$

4 Quadrature and Numerical Integration

approximation will be as follows,

$$\int_a^b W(x) f(x) dx \approx \sum_{i=1}^n \omega_i f(x_i) \quad (20)$$

The advantage of incorporating a weighting function, $W(x)$, is that singularities or difficult terms can be removed from the function to be integrated.

An n -point Gaussian quadrature rule is constructed to be exact for polynomials of degree $(2n - 1)$, by a suitable choice of the points, x_i , and weights, ω_i . In fact, it can be shown (see e.g., [5]) that the abscissas used for the n -point Gaussian quadrature formulas are optimal and given precisely by the roots of the n th orthogonal polynomial, $p_n(x)$, for the same interval, $[a, b]$, and weighting function, $W(x)$. A set of polynomials $\{p_n(x)\}$ is orthogonal if they satisfy the condition,

$$\int_a^b W(x) p_n(x) p_m(x) dx = 0, \quad \text{for } n \neq m \quad (21)$$

The simplest form of Gaussian quadrature uses uniform weighting, $W(x) = 1$, and the reference interval is $[a, b] = [-1, 1]$. The orthogonal polynomials for this weighting function and interval are the Legendre polynomials, see e.g. [5], which are used to approximate the integrand $f(x)$ over the interval $[-1, +1]$ as follows,

$$\int_{-1}^1 f(x) dx \approx \sum_{i=1}^n \omega_i f(x_i) \quad (22)$$

where x_i are the roots of the Legendre polynomials and ω_i are the corresponding weights, which are well known for this case and are given in Table 2. If we are interested in approximating the integral defined on an arbitrary interval, $[a, b]$, we can use a simple change of variable as follows,

$$\begin{aligned} \int_a^b f(x) dx &= k \int_{-1}^1 f(c + kx) dx \\ &\approx k \sum_{i=1}^n \omega_i f(c + kx_i) \end{aligned} \quad (23)$$

where $c = (b + a)/2$ and $k = (b - a)/2$. This is called the *Gauss–Legendre quadrature rule*. The following example illustrates and compares it with the Simpson's approximation.

Example Assume that we are interested in the approximation of the following integral,

$$\int_0^{\pi/2} \cos(x) dx \quad (24)$$

which is known analytically to be one. Firstly, we use the two-point Gauss–Legendre quadrature, given in equation (23), to obtain the following approximation,

$$\begin{aligned} \int_0^{\pi/2} \cos(x) dx &\approx \frac{\pi}{4} \left[1.0 \times \cos\left(\frac{\pi}{4} + \frac{\pi}{4\sqrt{3}}\right) + 1.0 \right. \\ &\quad \left. \times \cos\left(\frac{\pi}{4} - \frac{\pi}{4\sqrt{3}}\right) \right] = 0.99847 \end{aligned} \quad (25)$$

which is very close to the true value of one and gives an approximation error of 0.00153. Now, we approximate the same integral using the three-point Simpson's quadrature, given in equation (13),

$$\begin{aligned} \int_0^{\pi/2} \cos(x) dx &\approx \frac{\pi}{12} \left[\cos(0) + 4 \cos\left(\frac{\pi}{4}\right) \right. \\ &\quad \left. + \cos\left(\frac{\pi}{2}\right) \right] = 1.0023 \end{aligned} \quad (26)$$

This gives an approximation error of 0.0023 which is larger than the error obtained previously with the Gauss–Legendre quadrature that was based only on two points.

The Gaussian quadrature is also superior when it is compared with the composite Simpson's rule. For example, the four-point Gauss–Legendre quadrature gives the following approximation for the same integral,

$$\begin{aligned} \int_0^{\pi/2} \cos(x) dx &\approx \frac{0.65215\pi}{4} \times \cos\left(\frac{\pi}{4} + \frac{0.3399\pi}{4}\right) \\ &\quad + \frac{0.65215\pi}{4} \times \cos\left(\frac{\pi}{4} - \frac{0.3399\pi}{4}\right) \\ &\quad + \frac{0.34785\pi}{4} \times \cos\left(\frac{\pi}{4} + \frac{0.8611\pi}{4}\right) \\ &\quad + \frac{0.34785\pi}{4} \times \cos\left(\frac{\pi}{4} - \frac{0.8611\pi}{4}\right) \\ &= 0.999999977 \end{aligned} \quad (27)$$

Table 2 Gauss–Legendre abscissas and weights

n^o points	2	3	4	5				
x_i	$\pm\sqrt{\frac{1}{3}}$	0.0	± 0.7746	± 0.3399	± 0.8611	0.0	± 0.5385	± 0.9062
ω_i	1.0	0.88889	0.55556	0.65215	0.34785	0.56889	0.47863	0.23693

which is a much better approximation than that obtained with the following composite Simpson’s quadrature based on five points, as given in equation (15),

$$\begin{aligned} \int_0^{\frac{\pi}{2}} \cos(x) \, dx &\approx \frac{\pi}{24} \cos(0) + 4 \cos\left(\frac{\pi}{8}\right) + 2 \cos\left(\frac{\pi}{4}\right) \\ &\quad + 4 \cos\left(\frac{3\pi}{8}\right) + \cos\left(\frac{\pi}{2}\right) \\ &= 1.000134585 \end{aligned} \quad (28)$$

Alternatively to the Gauss–Legendre quadrature, other Gaussian rules can be developed using different classical orthogonal polynomials. Table 3 shows some of these polynomials together with their associated intervals and weighting functions. A large amount of information about these polynomials and their properties can be found in [7] and [9]. See also [5] and [10] for numerical procedures on the calculation of the abscissas and weights.

When we have an unusual choice for $W(x)$, we can develop our own Gaussian quadrature although this task usually becomes more difficult. Firstly, we need to obtain the set of orthogonal polynomials for the considered interval and weighting function. This can be done for example using a recurrence relation as described in [5]. Then, we need to calculate the zeros of these polynomials which will be the points, x_i , at which to evaluate the function. These can be obtained using a root-finding algorithm like Newton’s method or faster procedures as those described in [5].

Table 3 Classical orthogonal polynomials with their intervals and weighting functions

Interval	$W(x)$	Symbol	Polynomial
$[-1, 1]$	1	$P_n(x)$	Legendre
$[0, \infty)$	e^{-x}	$L_n(x)$	Laguerre
$(-\infty, \infty)$	e^{-x^2}	$H_n(x)$	Hermite
$[-1, 1]$	$(1 - x^2)^{-1/2}$	$T_n(x)$	First-kind Chebyshev
$[-1, 1]$	$(1 - x^2)^{1/2}$	$U_n(x)$	Second-kind Chebyshev

Finally, we need to calculate the weights ω_i which can be obtained analytically by integrating the orthogonal polynomials for which the approximation is exact. There are also alternative formulas, which are more efficient for the calculation of these weights, see e.g. [5] and [10].

The Gaussian quadrature can be modified in order to incorporate one or both endpoints in the set of abscissas, leading to the Radau and Lobatto quadrature formulas respectively, see e.g. [10]. These formulas use a uniform weighting function $W(x) = 1$ and the free abscissas are the roots of some polynomials related to the Legendre orthogonal polynomials. Radau and Lobatto quadratures are slightly less optimal than the Gaussian quadrature.

Another important extension of the Gaussian quadrature is the Gauss–Kronrod algorithm, see e.g. [5] and [10]. This is an adaptive Gaussian method where the abscissas are suitably selected such that they can be reused in the next iterations reducing the number of function evaluations. This approach is implemented in various software packages such as Mathematica (Wolfram Research Inc.).

Discussion and Remarks

The generalization of the described quadrature rules to the multidimensional case is not straightforward. The main reason is that the number of function evaluations for an n -dimensional integral increases to the power of n and then, it becomes very expensive even for low-dimensional integrals. An alternative approach for these cases is the use of Monte Carlo integration (*see Monte Carlo Methods; Markov Chain Monte Carlo, Introduction*) which gives reasonable approximations for multidimensional integrals defined over complicated regions.

It is frequent in practice to have a tabulated function at given points, x_i , obtained from experimental data. In these cases, the abscissas are predetermined and cannot be chosen at will. One possibility is to

obtain the Lagrange polynomial which interpolates the observed points and integrate it to approximate the integral, see e.g. [10]. One alternative is the use of splines which are piecewise polynomial functions with some finite derivatives that can be arranged to interpolate the data points, see e.g. [2].

Finally, note that the problem of evaluating an integral such as that given in equation (1) is equivalent to solving the differential equation $y'(x) = f(x)$ with $y(a) = 0$. There are many numerical procedures in the literature for solving differential equations that could be applied for this problem, see e.g. [5]. These would be specially well suited when the function to integrate is concentrated around one or various peaks.

References

- [1] Davis, P.J. & Rabinowitz, P. (1984). *Methods of Numerical Integration*, 2nd Edition, Academic Press, New York.
- [2] Evans, G. (1993). *Practical Numerical Integration*, John Wiley & Sons, Chichester.
- [3] Hildebrand, F.B. (1956). *Introduction to Numerical Analysis*, McGraw-Hill, New York.
- [4] Krommer, A.R. & Ueberhuber, C.W. (1994). *Numerical Integration on Advanced Computer Systems*, Springer-Verlag, Berlin.
- [5] Press, W.H., Teukolsky, S.A., Vetterling, W.T. & Flannery, B.P. (2002). *Numerical Recipes in C++: The Art of Scientific Computing*, 2nd Edition, Cambridge University Press, Cambridge.
- [6] Ueberhuber, C.W. (1997). *Numerical Computation: Methods, Software, and Analysis*, Springer-Verlag, Berlin.
- [7] Weisstein, E.W. (2007). *Newton-Cotes Formulas*. From *MathWorld—A Wolfram Web Resource*. <http://mathworld.wolfram.com/Newton-CotesFormulas.html>.
- [8] Abramowitz, M. & Stegun, I.A. (eds) (1972). *Handbook of Mathematical Functions with Formulas, Graphs, and Mathematical Tables*, Dover Publications, New York.
- [9] Stroud, A.H. & Secrest, D. (1966). *Gaussian Quadrature Formulas*, Prentice Hall, Englewood Cliffs.
- [10] Weisstein, E.W. (2007). *Gaussian Quadrature*. From *MathWorld—A Wolfram Web Resource*. <http://mathworld.wolfram.com/GaussianQuadrature.html>.

Related Articles

Convergence and Mixing in Markov Chain Monte Carlo; Laplace Approximations in Bayesian Lifetime Analysis; Monte Carlo Methods; Markov Chain Monte Carlo, Introduction.

M. CONCEPCIÓN AUSÍN

Bayesian Networks in Reliability

Introduction

Definition

Mathematically, a Bayesian network (BN) is a compact representation of a joint statistical distribution function. A BN encodes the **probability density function** governing a set of variables by specifying a set of conditional independence statements together with a set of conditional probability functions (CPFs). More specifically, a BN consists of a qualitative part, a *directed acyclic graph* where the nodes mirror the variables, and a quantitative part, the set of CPFs. An example of a BN is shown in Figure 1, only the qualitative part is given. This BN models the risk of an explosion in a process system. *Explosions* occurs if there is a *Leak* of chemical substance that is not detected by the gas detection system. This system detects all leaks unless it is in its failed state (*GD Failed*). The *Environment* will influence the probability of both leaks and a failure of the gas detection system. Finally, an explosion may lead to a number of *Casualties*.

There is a similar example of a BN in a article named **Bayesian Networks**, where inference and learning in BNs are considered. Here we focus on how BNs may be built, and their use in decision making.

For notational convenience, we consider the variables $\{X_1, \dots, X_n\}$ when we make some general definitions about BNs in the following. Now, we call the nodes with outgoing edges pointing into a specific node the *parents* of that node, and say that X_j is a *descendant* of X_i if and only if there exists a directed path from X_i to X_j in the graph. In Figure 1, *Leak?* and *GD fails?* are the parents of *Explosion?*, written $\text{pa}(\text{Explosion?}) = \{\text{Leak?}, \text{GD fails?}\}$ in short. Furthermore, $\text{pa}(\text{Casualties?}) = \{\text{Explosion?}\}$. Since there are no directed paths from *Casualties?* to any of the other nodes, the descendants of *Casualties?* are given by the empty set and, accordingly, its non-descendants are $\{\text{Environment?}, \text{GD fails?}, \text{Leak?}, \text{Explosion?}\}$. The edges of the graph represent the assertion that a variable is conditionally independent of its nondescendants in the graph given its parents

in the same graph. The graph in Figure 1 does for instance assert that for all distributions compatible with it, we have that *Casualties?* is conditionally independent of $\{\text{Environment?}, \text{GD fails?}, \text{Leak?}\}$ when conditioned on $\{\text{Explosion?}\}$.

The graphical structure has an intuitive interpretation as a model of causal influences. Although this interpretation is not necessarily entirely correct, it is helpful when the BN structure is to be elicited from experts. Furthermore, it can also be defended if some additional assumptions are made.

When it comes to the quantitative part, each variable is described by the CPF of that variable *given the parents* in the graph, that is, the collection of CPFs $\{f(x_i | \text{pa}(x_i))\}_{i=1}^n$ is required. The underlying assumptions of conditional independence encoded in the graph allow us to calculate the joint probability function as

$$f(x_1, \dots, x_n) = \prod_{i=1}^n f(x_i | \text{pa}(x_i)) \quad (1)$$

and this is in fact the main point when working with BNs: Assume that a distribution function $f(x_1, \dots, x_n)$ factorizes according to equation (1). This defines the parent set of each X_i , which in turn defines the graph, and from the graph we can read off the conditional independence statements encoded in the model. As we have seen, this also works the other way around, as the graph defines that the joint distribution *must* factorize according to equation (1). Thus, the graphical representation is the bridging of the gap between the (high-level) conditional independence statements we want to encode in the model and the (low-level) constraints this enforces on the joint distribution function.

After having established the full joint distribution over $\{X_1, \dots, X_n\}$ (using equation (1)), any marginal distribution $f(x_i, x_j, x_k)$, as well as any conditional distribution $f(x_i, x_j | x_k, x_\ell)$, can be calculated using extremely efficient algorithms.

Use of BNs in Reliability

BNs originated in the field of artificial intelligence, where they were used as a robust and efficient framework for reasoning with uncertain knowledge. The history of BNs in reliability can be traced back (at least) to Barlow [1] and Almond [2]; the first real attempt to merge the efforts of the two communities is

2 Bayesian Networks in Reliability

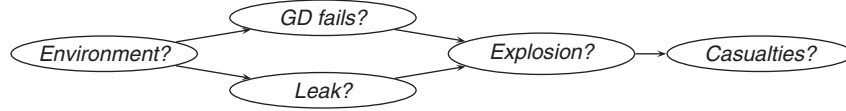


Figure 1 An example BN describing a gas leak scenario. Only the qualitative part of the BN is shown

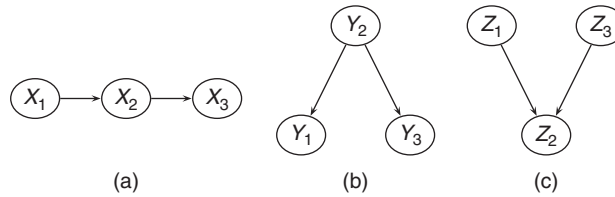


Figure 2 Three small network fragments describing different structural situations: serial (a), diverging (b), and converging (c)

probably the work of Almond [2], where he proposes the use of the *Graphical-Belief* tool for calculating reliability measures concerning a low-pressure coolant injection system for a nuclear reactor (a problem originally addressed by Martz and Waller [3]).

BNs constitute a modeling framework that is particularly easy to use in interaction with domain experts, also in the reliability field [4]. BNs have found applications in, for example, fault detection and identification, monitoring, software reliability, troubleshooting systems, and **maintenance optimization**.

Further Information on Bayesian Networks

Pearl [5] gives the first full account of BNs, whereas Charniak [6] gives a reader-friendly introduction to the field. The recent book by Jensen and Nielsen [7] gives a comprehensive review of the theory and methods of Bayesian networks and decision graphs. The causal perspective toward BNs is thoroughly analyzed in [8].

Several standard books on statistical modeling and decision theory contain chapters on BNs, and some books on risk and reliability analysis include discussions of these issues, see for example [9]. A review of the usage of BNs in reliability analysis can be found in [10]. The annual *Conference on Uncertainty in Artificial Intelligence* and, in particular, the co-located *Bayesian Modeling Applications Workshops*, are good references for the continuous development in the community.

Many BN tools are available to the practitioners. Examples include Hugin (www.hugin.com), BayesiaLab (www.bayesia.com), and Netica (www.norsys.com). For a comprehensive list of software tools, see the one maintained by Kevin Murphy at www.cs.ubc.ca/~murphyk/Software/BNT/bnsoft.html.

Conditional Independence

The driving force when making BN models is the set of *conditional independence* statements the model encodes. It is common to use the notation $\mathcal{X} \perp\!\!\!\perp \mathcal{Y} \mid \mathcal{Z}$ to denote that the variables in the two sets $\mathcal{X}$ and $\mathcal{Y}$ are conditionally independent given the variables in $\mathcal{Z}$. If $\mathcal{Z}$ is the empty set, we simply write $\mathcal{X} \perp\!\!\!\perp \mathcal{Y}$. Finally, we use $\mathcal{X} \not\perp\!\!\!\perp \mathcal{Y} \mid \mathcal{Z}$ to make explicit that $\mathcal{X}$ and $\mathcal{Y}$ are conditionally *dependent* given $\mathcal{Z}$.

We saw the first example of how we can read conditional independence statements from the graph when we concluded that $\{Casualties?\} \perp\!\!\!\perp \{Environment?, GD\ fails?, Leak?\} \mid \{Explosion?\}$ in Figure 1. The general analysis of conditional independence is a bit broader, and uses the concept of *d-separation*, which centers around the three categories of network fragments shown in Figure 2.

The *serial connection* (Figure 2a) encodes that $X_1 \perp\!\!\!\perp X_3 \mid X_2$, but $X_1 \not\perp\!\!\!\perp X_3$. For example, let X_1 denote the *planned* preventive maintenance (PM) program for a given component, let X_2 be the *implemented* PM, and X_3 the life length of the component. The conditional independence statements encoded

in this model tell us that if we do not know the implemented PM program, then the planned PM can tell us something about the life length of the component ($X_1 \perp\!\!\!\perp X_3$). However, as soon as we learn about the implemented PM program, the *plans* are irrelevant for the life length: $X_1 \perp\!\!\!\perp X_3 \mid X_2$.

Figure 2(b), a *diverging connection*, dictates similar properties: $Y_1 \perp\!\!\!\perp Y_3 \mid Y_2$, but $Y_1 \not\perp\!\!\!\perp Y_3$. This BN can for instance model the quality of the production from an assembly line. Let Y_1 be the quality of the first item that was produced at this line, and Y_3 be the quality of the second item. Finally, let Y_2 be a measure of how well the assembly line operates overall. Now, if Y_2 is unknown to us, information about Y_1 being good (bad) would make us infer that the quality of the line as such (Y_2) was good (bad) as well, and therefore we would expect that Y_3 would be good (bad), too. Thus, $Y_1 \not\perp\!\!\!\perp Y_3$. On the other hand, if the quality of the production line is known, this model says that the quality of each produced item may be seen as independent of each other ($Y_1 \perp\!\!\!\perp Y_3 \mid Y_2$).

Finally, the *converging connection* in Figure 2(c) encodes that $Z_1 \perp\!\!\!\perp Z_3$, but $Z_1 \not\perp\!\!\!\perp Z_3 \mid Z_2$. To understand this part, one can think of Z_1 as the quality of an assembly line, Z_3 as the environmental conditions of the production (temperature, humidity, etc.), and Z_2 as the quality of a product coming from the assembly line. In this model, the quality of the assembly line is *a priori* independent of the environmental conditions ($Z_1 \perp\!\!\!\perp Z_3$), however, as soon as we observe the quality of the product, we can make inference regarding the quality of the line from what is known about the environmental conditions: If the quality of the product is low, although it was produced under good environmental conditions, it is reasonable to assume that this may be because the overall quality of the assembly line is poor ($Z_1 \not\perp\!\!\!\perp Z_3 \mid Z_2$).

Building Bayesian Networks

Building BNs from expert input can be a difficult and time-consuming task. This is typically an assignment given to a group of specialists. A BN expert guides the model building, asks relevant questions, and explains the assumptions that are encoded in the model to the rest of the group. The domain experts, on the other hand, supply their knowledge to the BN expert in a structured fashion. In our experience, it will pay off to start the model building

by familiarization: the BN expert should study the domain, and the reliability analysts require basic knowledge about BNs. As soon as this is established, model building will proceed through a number of phases:

Step 0: Decide what to model. Select the boundary for what to include in the model, and what to leave aside.

Step 1: Defining variables. Select the important variables in the model. The *range* of the continuous variables and the *states* of the discrete variables are also determined at this point.

Step 2: The qualitative part. Next, we define the graphical structure that connects the variables. In this phase it can be beneficial to consider the edges in the graph as *causal*, but the trained domain expert may also be confident about (conditional) dependencies/independencies to include in the model. In the latter case, the *PC algorithm* [11] is a principled way of asking the most efficient questions about conditional independencies. Domain experts can often be very eager to incorporate an unpractically large number of links in the structure in an attempt to “get it right”. The BN expert’s task is in this setting to balance the model’s complexity with the modeling assumptions the domain experts are willing to accept. Often, a post-processing of the structure may reveal void edges (e.g., those creating triangles in the graph [12]).

Step 3: The quantitative part. To define the quantitative part, one must select distributional families for all variables, and fix parameters to specify the distributions. If the BN structure has not been carefully elicited (and pruned) this may be a formidable task. Luckily, the consistency problems common when eliciting large probability distribution functions are tackled by a “divide-and-conquer” strategy here: If each CPF $f(x_i \mid \text{pa}(x_i))$ is defined consistently, then this implies global consistency as well. To elicit the quantitative part from experts, one must acquire all CPFs $\{f(x_i \mid \text{pa}(x_i))\}_{i=1}^n$ in equation (1), and once again the causal interpretation can come in as a handy tool.

To fully specify the set of CPFs, we must (a) select parametric families for each $f(x_i \mid \text{pa}(x_i))$ and (b) determine values for all parameters of each CPF. To do the last part, we must use either statistical data or expert judgment, but since both these sources of information can have low quality, as well as come with a cost, one would like the BN formalism to

minimize the number of parameters required. This is exactly what it attempts to do; it represents the joint distribution in a cost-efficient manner (through its factorized representation in equation (1)). Let us consider the domain represented by $\{X_1, \dots, X_n\}$, and let the nodes be labeled such that $\text{pa}(X_i) \subseteq \{X_1, \dots, X_{i-1}\}$ for $i = 1, \dots, n$ (this can always be obtained through relabeling). Although it is not essential for the following argument, we will for now assume that all variables take on values in $\{0, 1\}$ (these values could for instance signify a component that has either failed or is operating). According to equation (1) we will require CPFs of the type $f(x_i | \text{pa}(x_i))$; these are in this case fully determined by the probability $P(X_i = 1 | \text{pa}(x_i))$. The number of parameters required to specify $f(x_i | \text{pa}(x_i))$ for a given i is thus equal to $2^{|\text{pa}(X_i)|}$, and the total number of parameters is given by $\sum_{i=1}^n 2^{|\text{pa}(X_i)|}$. This number is by construction not higher than $2^n - 1$, that is the number of parameters required to define the full joint distribution (without using the conditional independence statements in the BN structure). Hence, when we use a BN, we are never worse off than if the full joint distribution is defined directly, and we are bound to improve as long as at least one conditional independence assumption can be made. As an example, if all variables of the domain in Figure 1 are binary, then the full joint distribution will require 31 parameters to be specified. Alternatively, the BN representation uses only 11 parameters.

Step 4: Verification. Verification should be performed both through sensitivity analysis as well as by testing how the model behaves when analyzing well-known scenarios. Typically, this step gives need for refinement/redefinition of the model, and this is repeated until the work put into improving the model does not lead to substantial benefits. As pointed out by, for example, [13], the sensitivity with respect to BN structure is relatively high, and the graph is thus

the most vital part. Sensitivity with respect to the parameters is by large dependent on the application.

Lately, some tools that are aimed at guiding model building have emerged. These tools attempt to enable a domain expert to build a BN without interacting with a BN expert. For instance, [14] describes a system that can be used to build troubleshooter systems efficiently.

Decision Graphs

One of the most intriguing properties of the BN framework is that it can easily be extended to represent *decisions*. The basis for the representation is utilities, which are quantified measures for preference. That is, a real number is attached to each possible scenario in question. Let $S = \{s_1, \dots, s_t\}$ be the possible scenarios with the probability distribution $(p_1, \dots, p_t)$, and let $(u_1, \dots, u_t)$ be the corresponding utilities. The expected utility of S is then $EU(S) = \sum_i u_i p_i$.

Example

Assume for the scenario in Figure 1 that we wish to consider the possibility of a structural safety measure (*SSM*). The *SSM* has an impact on the number of casualties in case of an explosion, and the cost of the *SSM* shall be balanced with the expected reduction in costs due to casualties. To represent this, the BN is extended with two types of variables, *decision nodes* (rectangular shaped), representing the various decision options, and *utility nodes* (diamond shaped), representing preferences on a numerical scale. This is done in Figure 3.

The edge from *SSM* to *Casualties?* indicates the impact of an *SSM*, and the quantitative specification is the conditional probability $f(\text{Casualties?} | \text{SSM}, \text{Explosion?})$. The edges to the *Cost* nodes indicate the variables relevant for the cost.

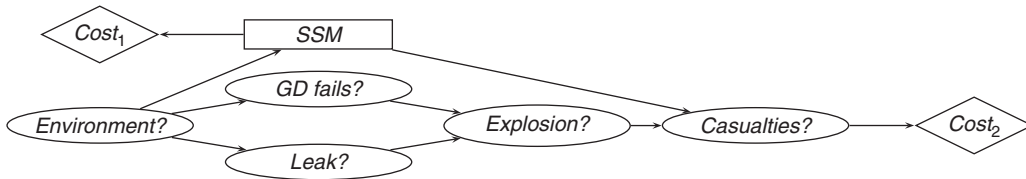


Figure 3 Extending the BN in Figure 1 to optimally choose the level of structural protection (*SSM*)

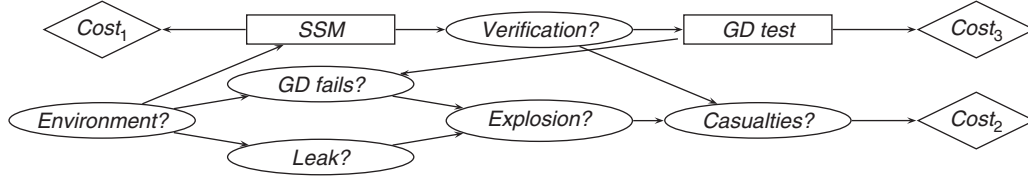


Figure 4 An influence diagram (ID) extending the one-decision scenario in Figure 3. It can be used to solve the problem with several decisions

That is, $Cost_1$ is a function of SSM , whereas $Cost_2$ is a consequence of the number of casualties (if any). The edge from *Environment?* to SSM indicates that the state of the environment is known when the decision on $SSMs$ is made.

Exploiting probability updating in the BN (Figure 1), it is now easy to calculate the expected utility for each decision option given the environment. The result is an optimal policy, which for any state of the environment selects a decision of maximal expected utility.

Several Decisions

The framework also allows the representation of decision problems involving more than one decision. Assume in the previous decision problem that we also wish to decide a testing regime for the gas detection system. We introduce a decision node GD test to represent the decision on testing frequency (see Figure 4). At the time when this is decided, we know the decision on SSM , and we have also run an on-site verification to quantify the effectiveness of the implemented $SSMs$ (*Verification?*). We use the latter as input to decide the test interval, hence the edge from *Verification?* to GD test.

The result is an optimal policy for SSM as well as for GD Test. The optimal policy for GD test is a function of *Verification?*, *Environment?*, and SSM . The calculation of an optimal policy for SSM is in this situation more complex than in the previous example. When deciding on SSM we need to have an idea of the GD test decision. Therefore, the first thing to calculate is an optimal policy for GD test, and assuming that we will follow this policy, we can determine an optimal policy for SSM . The two policies together form an optimal strategy for the decision problem.

Influence Diagrams

When the decision problem consists of a predefined sequence of observations and decisions (as in the

example above) it can be represented as an influence diagram [15]. An influence diagram is an extension of a BN with decision nodes and utility nodes. An edge from a node A to a decision node D indicates that the state of A is known before the decision D is taken. In order not to overload the graph with edges, the *no-forgetting* property is usually assumed: If the state of A is known when taking decision D , then it is also known when taking all subsequent decisions. To reflect the premise of a predefined sequence of observations and decisions, there is graphical constraint for influence diagrams: there must be a directed path encompassing all decision nodes.

An influence diagram is primarily used to calculate the expected utilities for the first decision variable. To do this, we have to predict the decisions we will take in the future, and in order to decide on the present decision we calculate policies for the future decisions based on the principle of maximizing the expected utility [16–18].

Although influence diagrams are a powerful modeling tool, the framework has its shortcomings. The decision problem must have a finite time horizon, and it must be *symmetric*: the next nodes to observe, and the next decisions to take, must be independent of the current observations and decisions. Work has been done to overcome these shortcomings [19–21], but these languages are much more complex to work with.

Further Pointers to the Literature

BNs are *directed* graphical models. Sometimes, a domain expert may be uncomfortable directing an edge between two nodes. A *chain graph* is a graphical model closely related to BNs, but which accepts both directed and undirected edges. The interested reader is referred to [22].

In this article we have discussed how BNs could be elicited from domain experts. Another possibility is to learn the BNs from historical data. This is an active research topic, see, for example, [23].

Modeling *dynamics* is often a critical part of a reliability application. As an example, consider a fault detection system, which continuously monitors the environment looking for indications that it is about to become unstable. Obviously, modeling development over time is crucial in this situation. There are several special classes of BNs developed for such situations, Smyth [24] uses one of these, *Hidden Markov Models*, for detecting failures in satellite antennas.

We see a preference for *discrete* variables in the BN community, mainly due to the technicalities of the calculation scheme used to calculate the posterior probabilities. We note that the applicability of BNs in risk and reliability analysis would be enormously limited if one would consider only discrete variables, and this is therefore the most important feature that is missing before BNs provide a full framework for such tasks. Numerous efforts have been taken up to solve this problem, ranging from inference by sampling [25], via exact calculations in restricted structures [26, 27], to approximations of the CPFs of the continuous variables [28, 29].

References

- [1] Barlow, R.E. (1988). Using influence diagrams. In *Accelerated Life Testing and Experts' Opinions in Reliability*, Vol. 102 in *Enrico Fermi International School of Physics*, C.A. Clarotti & D.V. Lindley, eds, Elsevier Science Publishers, BV, North-Holland, pp. 145–157.
- [2] Almond, R.G. (1992). An extended example for testing GRAPHICAL-BELIEF, Technical Report 6, Statistical Sciences, Seattle.
- [3] Martz, H.F. & Waller, R.A. (1990). Bayesian reliability analysis of complex series/parallel systems of binomial subsystems and components, *Technometrics* **32**(4), 407–416.
- [4] Sigurdsson, J.H., Walls, L.A. & Quigley, J.L. (2001). Bayesian belief nets for managing expert judgment and modeling reliability, *Quality and Reliability Engineering International* **17**, 181–190.
- [5] Pearl, J. (1988). *Probabilistic Reasoning in Intelligent Systems: Networks of Plausible Inference*, Morgan Kaufmann Publishers, San Mateo.
- [6] Charniak, E. (1991). Bayesian networks without tears, *AI Magazine* **12**(4), 50–63.
- [7] Jensen, F.V. & Nielsen, T.D. (2007). *Bayesian Networks and Decision Graphs*, Springer-Verlag, Berlin.
- [8] Pearl, J. (2000). *Causality-Models, Reasoning, and Inference*, Cambridge University Press, Cambridge.
- [9] Bedford, T.J. & Cooke, R.M. (2001). *Probabilistic Risk Analysis: Foundations and Methods*, Cambridge University Press.
- [10] Langseth, H. & Portinale, L. (2007). Bayesian networks in reliability, *Reliability Engineering and System Safety* **92**(1), 92–108.
- [11] Spirtes, P., Glymour, C. & Scheines, R. (1993). *Causation, Prediction, and Search*, Springer-Verlag, New York.
- [12] Abramson, B., Brown, J., Edwards, W., Murphy, A. & Winkler, R.L. (1996). Hailfinder: a Bayesian system for forecasting severe weather, *International Journal of Forecasting* **12**, 57–71.
- [13] Druzdzel, M. & van der Gaag, L. (2000). Building probabilistic networks: where do the numbers come from? -a guide to the literature, Technical Report UU-CS-2000-20, Institute of Information & Computing Sciences, University of Utrecht, Utrecht.
- [14] Skaaning, C. (2000). A knowledge acquisition tool for Bayesian-network troubleshooters, *Proceedings of the Sixteenth Conference on Uncertainty in Artificial Intelligence*, Morgan Kaufmann Publishers, San Francisco, pp. 549–557.
- [15] Howard, R.A. & Matheson, J.E. (1984). Influence diagrams, in *Readings on the Principles and Applications of Decision Analysis II*, R.A. Howard & J.E. Matheson, eds, Strategic Decision Group, Menlo Park, pp. 719–762.
- [16] Shachter, R.D. (1986). Evaluating influence diagrams, *Operations Research* **34**(6), 871–882.
- [17] Shenoy, P.P. (1992). Valuation-based systems for Bayesian decision analysis, *Operations Research* **40**(3), 463–484.
- [18] Jensen, F., Jensen, F.V. & Dittmer, S.L. (1994). From influence diagrams to junction trees, *Proceedings of the Tenth Conference on Uncertainty in Artificial Intelligence*, Morgan Kaufmann Publishers, San Francisco, pp. 367–373.
- [19] Kaelbling, L.P., Littman, M.L. & Cassandra, A.R. (1998). Planning and acting in partially observable stochastic domains, *Artificial Intelligence* **101**, 99–134.
- [20] Jensen, F.V. & Vomlelova, M. (2002). Unconstrained influence diagrams, *Proceedings of the Eighteenth Conference on Uncertainty in Artificial Intelligence*, Morgan Kaufmann Publishers, San Francisco, pp. 234–241.
- [21] Jensen, F.V., Nielsen, T.D. & Shenoy, P.P. (2006). Sequential influence diagrams: a unified framework, *International Journal of Approximate Reasoning* **42**(1), 101–118.
- [22] Lauritzen, S.L. & Richardson, T.S. (2002). Chain graph models and their causal interpretation (with discussion), *Journal of the Royal Statistical Society. Series B* **64**, 321–361.
- [23] Neapolitan, R.E. (2003). *Learning Bayesian Networks*, Prentice Hall, Upper Saddle River.

- [24] Smyth, P. (1994). Markov monitoring with unknown states, *IEEE Journal of Selected Areas in Communications, Special Issue on Intelligent Signal Processing for Communications* **12**(9), 1600–1612.
- [25] Gilks, W., Thomas, A. & Spiegelhalter, D.J. (1994). A language and program for complex Bayesian modelling, *The Statistician* **43**, 169–178.
- [26] Lauritzen, S.L. & Wermuth, N. (1989). Graphical models for associations between variables, some of which are quantitative and some qualitative, *The Annals of Statistics* **17**, 31–57.
- [27] Lerner, U., Segal, E. & Koller, D. (2001). Exact inference in networks with discrete children of continuous parents, *Proceedings of the Seventeenth Conference on Uncertainty in Artificial Intelligence*, Morgan Kaufmann Publishers, San Francisco, pp. 319–328.
- [28] Jordan, M.I., Ghahramani, Z., Jaakkola, T.S. & Saul, L.K. (1999). An introduction to variational methods for graphical models, *Machine Learning* **37**, 183–233.
- [29] Moral, S., Rumí, R. & Salmerón, A. (2001). Mixtures of truncated exponentials in hybrid Bayesian networks,

Sixth European Conference on Symbolic and Quantitative Approaches to Reasoning with Uncertainty, Volume 2143 of Lecture Notes in Artificial Intelligence, Springer-Verlag, Berlin, pp. 145–167.

Related Articles

Bayesian Control Charts; Bayesian Networks; Bayesian Reliability Analysis; Bayesian Reliability Demonstration; Causality; Decision Search via Simulation; Dependence; Expert Opinion in Reliability; Fault Trees; Masked Failure Data; Bayesian Modeling; Markov Chain Monte Carlo, Introduction; Prior Distribution Elicitation; Software Reliability; Bayesian Analysis; Statistical Process Control, Bayesian.

HELGE LANGSETH AND FINN VERNER JENSEN

Computational Issues in Network Reliability

Introduction

Networks are pervasive and familiar structures in our technological society, ranging from fiber-optic telecommunication systems, to electrical power distribution systems, to systems of highway and airline connections. Such networks are designed to distribute information, electricity, and people from place to place in an efficient and economical manner. Network systems are composed of *nodes* (representing switching centers, workstations, concentrators, facilities, cities, etc.), which are connected by a collection of *edges* (cables, pipelines, roads, or even wireless connections).

These network components (nodes, edges) are typically prone to failure. At any instant, each component is considered to be either *operating* or *failed*, and as a result the entire network system will be either operational or failed. The reliability of the system measures the probability that the system functions properly. Its value is critically dependent not only on the individual component reliabilities, but also on the manner in which they are interconnected. Important objectives are therefore to determine the reliability of a given network system (analysis) and to design cost-efficient network structures that have improved reliability (design). We focus here on certain computational aspects of network analysis.

The topology of a network is specified by its underlying *graph* $G = (V, E)$, where V denotes the set of nodes and E denotes the set of edges joining pairs of nodes. The graph is *undirected* if links between nodes are symmetric (bidirectional); in this case, the edge $e = (i, j) \in E$ indicates that nodes i and j are adjacent, joined by the undirected edge e . In a *directed* graph, there is a direction associated with each edge $e = (i, j) \in E$ (such as a one-way street); in this case, node j is considered to be adjacent from node i . Of course, an undirected graph can be modeled by replacing each undirected edge $e = (i, j)$ by a pair of oppositely directed edges (i, j) and (j, i) . Of fundamental importance is a *path* $P = [i_0, i_1, \dots, i_k]$ in a graph (undirected or directed): it is a sequence of edges $(i_0, i_1), (i_1, i_2), \dots, (i_{k-1}, i_k)$ leading from node i_0 to node i_k . A *simple* path

contains no repeated nodes. Two edges both joining node i to node j are called *parallel edges*; if there are exactly two edges (i, j) and (j, k) incident on node j , then these two edges are called *series edges*.

In general, both nodes and edges can fail. However, node failures can be simulated by edge failures [1], so we assume for simplicity that only edge failures occur. Specifically, p_e is the probability that edge e is operational at a random instant; $q_e = 1 - p_e$ is the failure probability of edge e . It is common to assume that the failures of distinct edges are statistically independent. Suppose that $S \subseteq E$ denotes the set of operating edges at a given instant, so that the edges in $E - S$ are failed. By independence, the probability of *state* S is given by

$$P(S) = \prod_{e \in S} p_e \prod_{e \in E-S} q_e \quad (1)$$

There are several ways of assessing the *reliability* $R(G, p)$ of a network having the topology specified by the graph $G = (V, E)$ and the edge reliabilities $\{p_e : e \in E\}$. In each of the following models, the system is required to maintain a certain degree of connectivity after random edge failures. If the graph is undirected, the following models are commonly used.

- [M1.] *Two-terminal reliability*: the reliability $R_{st} = R(G, p)$ is the probability that there exists a path of operating edges between two specified nodes s, t in G .
- [M2.] *All-terminal reliability*: the reliability $R_V = R(G, p)$ is the probability that paths of operating edges exist between every pair of nodes in G .
- [M3.] *K-terminal reliability*: the reliability $R_K = R(G, p)$ is the probability that paths of operating edges exist between every pair of nodes in $K \subseteq V$.

When G is directed, the above models generalize as follows.

- [M4.] *s, t connectedness*: the reliability $\vec{R}_{st} = R(G, p)$ is the probability that there exists a path of operating edges joining node s to node t in G .
- [M5.] *s reachability*: the reliability $\vec{R}_V = R(G, p)$ is the probability that paths of operating edges exist from the source node s to every node in G .

2 Computational Issues in Network Reliability

[M6.] s, K *connectedness*: the reliability $\vec{R}_K = R(G, p)$ is the probability that paths of operating edges exist from the source node s to every node in $K \subseteq V$.

In each of the above problems, the system is considered to be operating when the subgraph H of operating edges (defined by state S) possesses a certain structure; states that support operation are generically termed *pathsets*, and sets of edges whose removal leaves a state not supporting the operation are *cutsets*. *Minpaths* are minimal pathsets, and *mincuts* are minimal cutsets (see **Path Sets and Cut Sets in System Reliability Modeling**). For two-terminal reliability, nodes s and t need to be connected in the subgraph H . A minimal requirement then is that nodes s and t are connected by a simple path in H . Thus, the minpaths for this problem are the simple s - t paths in G . Similarly, the minpaths for the all-terminal reliability problem are the spanning trees [2] of G , while the minpaths for the K -terminal reliability problem are trees that span all the nodes in K . Directed analogs of minpaths are also easy to characterize: for example, minpaths for s reachability are directed spanning trees [2] of G , rooted at node s . Determining $R(G, p)$ can then be expressed as computing the probability of certain events associated with the minpaths of the system (see the section titled “Exact Computation of Reliability”).

The *unreliability* $U(G, q)$ of network G is given by $U(G, q) = 1 - R(G, 1 - q)$. It can also be defined more directly in terms of certain “dual” graph structures (mincuts). For example, the two-terminal unreliability U_{st} is equivalently the probability that the failed edges contain a set of edges whose removal disconnects s from t in G . Similarly, the all-terminal unreliability U_V is the probability that removal of the failed edges will disconnect G . Network unreliability $U(G, q)$ can also be expressed in terms of events associated with the mincuts of the system (see the section titled “Exact Computation of Reliability”).

For more complete treatments, and references for the material in subsequent sections, see [1, 3, 4].

Computational Complexity

We consider the special case of the reliability analysis problem that arises when all individual component reliabilities are equal: that is, $p_e = p$ for all edges

$e \in E$. If $|E| = m$, then the reliability can be written as a polynomial in p , $\sum_{i=0}^m F_i p^{m-i} (1-p)^i$, the *reliability polynomial*. The associated computational problem, the *functional reliability analysis problem*, takes a representation of a network (a graph with its associated edge probabilities) as input and produces the coefficients $\{F_i\}$ as output. The *rational reliability analysis problem* instead requires the evaluation of the reliability polynomial at a specific rational probability $p = a/b$.

The general term $F_i p^{m-i} (1-p)^i$ in the reliability polynomial is the probability that exactly $m-i$ edges operate and the system operates. Thus, we can interpret F_i as the number of operating system states (pathsets) having i failed edges. Equivalently the probability of network failure can be written as an *unreliability polynomial*, $\sum_{i=0}^m C_i p^{m-i} (1-p)^i$; here C_i is the number of cutsets with i edges. The problem of determining each of the coefficients F_i or C_i is a counting problem. NP and NP-complete denote classes of recognition problems, but counting problems associated with reliability can be classified as NP-hard when at least as hard as the NP-complete problems [5].

The functional reliability analysis problem can be reduced in polynomial time to the rational reliability analysis problem. Given a network, we define

c = cardinality of a minimum cardinality cutset,

C_c = number of minimum cardinality cutsets,

ℓ = cardinality of a minimum cardinality pathset,

N_ℓ = number of minimum cardinality pathsets.

Then

$$0 \leq F_i \leq \binom{m}{i} \text{ for } i = 0, 1, \dots, m$$

$$F_i = \binom{m}{i} \text{ for } i < c$$

$$F_i = \binom{m}{i} - C_c \text{ for } i = c$$

$$F_i = N_\ell \text{ for } i = m - \ell$$

$$F_i = 0 \text{ for } i > m - \ell \quad (2)$$

Table 1 Computational complexity of recognition and counting problems

	R_K	R_{st}	R_V	$\vec{R}_V$
Minimum cardinality pathset recognition	NP-hard	Polynomial	Polynomial	Polynomial
Minimum cardinality cutset recognition	Polynomial	Polynomial	Polynomial	Polynomial
Minimum cardinality pathset counting	NP-hard	Polynomial	Polynomial	Polynomial
Minimum cardinality cutset counting	NP-hard	NP-hard	NP-hard	Polynomial
General term F_i	NP-hard	NP-hard	NP-hard	NP-hard

Thus, by computing the reliability polynomial we determine important properties of the underlying graph. For example, we can determine from the reliability polynomial the size c of a minimum cardinality cutset. Thus, if the minimum cardinality cutset recognition problem is NP-hard, then computing the reliability polynomial is NP-hard. Indeed, for any network, the functional and rational reliability analysis problems are NP-hard whenever recognition or counting is NP-hard for minimum pathsets or minimum cutsets. Known complexity results for the K -terminal, two-terminal, and all-terminal problems are given in Table 1.

In light of these negative results, much research has been aimed at the analysis of structured networks. The widest class of networks known to be solvable in polynomial time involve series-parallel graphs and certain generalizations. However, the undirected two-terminal reliability problem remains NP-hard over planar networks having node degrees bounded by three and the directed two-terminal reliability analysis problems remain NP-hard over acyclic planar networks having node degrees bounded by three. Directed and undirected all-terminal reliability analysis problems are NP-hard when restricted to planar networks (see [2] for undefined graph-theoretic terms).

Exact Computation of Reliability

Now we outline exact algorithms for computing reliability measures. Most computational methods rely on a simple but important observation: there exist graph transformations that leave the values of various reliability measures unchanged, and these can often be used to simplify the network. Among these, deletion of *irrelevant* nodes and edges (those not appearing in any minpath) and standard series-parallel reductions, are the most frequently used.

When reliability-preserving transformations fail to reduce the network into a restricted class for which an efficient exact method is known, we are forced to resort to potentially exponential-time methods. In the first main class of these exact methods, we examine the possible states of the network.

A *state* of network $G = (V, E)$ is a subset $S \subseteq E$ of operational edges. The conceptually simplest exact algorithm is *complete state enumeration*. If $\mathcal{O}$ is the set of all pathsets, then

$$R(G, p) = \sum_{S \in \mathcal{O}} \prod_{e \in S} p_e \prod_{e \in E-S} (1 - p_e) \quad (3)$$

By generating all states, and determining which are operational, the reliability is “easily” (but not efficiently) computed.

Often, many states can be easily seen to be operational without listing them out one by one. Hence an immediate improvement is obtained by generating the states in a more clever way. A basic ingredient in this is the technique of *factoring*, also called *pivotal decomposition*:

$$R(G, p) = p_e R(G \cdot e, p) + (1 - p_e) R(G - e, p) \quad (4)$$

for any edge e of G . Here $G - e$ means G with edge e deleted, and $G \cdot e$ means G with edge e contracted – that is, the endpoints of e are amalgamated into a single node. Factoring carried out until the networks produced have no edges is just complete state enumeration. However, some simple observations result in improvements. When $G - e$ is failed, any sequence of contractions and deletions results in a failed network, and hence there is no need to factor $G - e$. Moreover, although we may be unable to simplify G with a reliability-preserving transformation, we may well be able to simplify $G \cdot e$ or $G - e$.

When it is possible to generate all minpaths or mincuts of the network, a second important class of algorithms is used. Suppose that the minpaths $P_1, \dots, P_h$ of G have been listed. Let E_i be the event that all edges in minpath P_i are operational. Then the reliability is just the probability that one or more of the events $\{E_i\}$ occurs. Unfortunately, the $\{E_i\}$ are not disjoint events, and hence we cannot simply sum their probabilities of occurrence. For example, representing “or” as $\vee$ and “and” as $\wedge$,

$$\Pr[E_1 \vee E_2] = \Pr[E_1] + \Pr[E_2] - \Pr[E_1 \wedge E_2] \quad (5)$$

In general $R(G, p) = \Pr[E_1 \vee E_2 \vee \dots \vee E_h]$, and hence

$$R(G, p) = \sum_{j=1}^h (-1)^{j+1} \sum_{I \subseteq \{1, \dots, h\}, |I|=j} \Pr[E_I] \quad (6)$$

where E_I is the event that all paths P_i with $i \in I$ are operational, that is, the event $\bigwedge_{i \in I} E_i$. This is a standard inclusion–exclusion expansion [1].

Many Boolean algebra methods simplify such Boolean formulas to represent all of the operational states. For example we can rewrite the probability that at least one minpath is operational as $\sum_{i=1}^h \Pr[\overline{E_1} \wedge \dots \wedge \overline{E_{i-1}} \wedge E_i]$. This expresses the reliability as the sum of h *disjoint* events; in the i th event, all of the first $i - 1$ minpaths fail, but the i th minpath operates. Simplification of the terms $\Pr[\overline{E_1} \wedge \dots \wedge \overline{E_{i-1}} \wedge E_i]$ is the key; existing methods focus on writing each as the conjunction (“and”) of the basic events of edge operation and edge failure (and hence in terms of independent events), in order to produce expressions that are as small as possible and can be calculated as simple sums of products. For algorithms and theory related to Boolean algebra methods in reliability, see [6].

Bounds and Approximations

None of the methods developed succeed in providing exact reliability metrics efficiently. Consequently two main directions have been explored, bounding and approximation. The essence of bounding is to provide practical computation of absolute upper and lower bounds on the reliability measure of interest. Investigations have pursued two quite

different lines. When practical computation is the main concern, and theoretical guarantees of efficiency are not needed, strategies based on inclusion–exclusion for states, pathsets, and cutsets have all been extensively applied. These provide the basis for widely used methods that generate the *most probable states* of the network, since a few states typically account for much of the network operation. See especially [4] for a more detailed exposition.

Theoretical guarantees of efficiency in bounding arise from three main strategies. In the first, the combinatorial or algebraic structure of the collection of operating states is used to determine, from a limited amount of information about these states, information about the rest. When information about all states can be deduced from information about only polynomially many states, efficient bounding techniques are obtained. In the second, transformations that lessen or increase the reliability measure, but at the same time simplify the structure of the network, underlie many effective algorithms. Finally, similar techniques that transform the set of pathsets or cutsets, rather than the network itself, have also proved effective. Detailed discussions of these appear in [1, 3].

Monte Carlo Methods

The intractability of exact computation and the inability of polynomial-time-bounding algorithms to provide tight bounds leads us to *simulation* techniques. The price is that the estimates obtained have a certain degree of uncertainty, but this is often well justified by superior results. The well-studied **Monte Carlo** method is most frequently used to simulate the stochastic behavior of the system.

We are interested in obtaining an estimate $\hat{R}$ for the true **network reliability** $R = R(G, p)$, in which the m edges have reliabilities $p_1, p_2, \dots, p_m$. Recall that the probability of each state $S \subseteq E$ is given by equation (1).

The *crude* method of sampling is straightforward. A sample of K states $S_1, \dots, S_K$ is chosen by generating mK independent samples u_{kj} , $k = 1, \dots, K$, $j = 1, \dots, m$ uniformly at random, and then placing e_j in S_k if and only if $u_{kj} \leq p_j$. Let $\hat{K}$ be the number of states k for which S_k is operating. Then an unbiased estimator for R is $\hat{R} = \hat{K}/K$, with variance $R(1 - R)/K$. Reduction of this variance can

be obtained by a number of standard Monte Carlo sampling techniques, such as *antithetic* and *control variates*, and *conditional*, *importance*, and *stratified sampling*.

Dagger sampling instead “spreads out” the individual edge failures in such a way that repeats of sample states is minimized. *Sequential construction* orders the edges of the graph. The edges begin as all failed, and then edges are successively “repaired”, that is, caused to operate, one by one in the specified ordering, until the system becomes operational. The reliability estimate is then a function of how long it takes for the system to become operational. This can result in better estimates than could be obtained by the crude method. *Sequential destruction* is the analogous process in which edges are removed one at a time.

Sampling using bounds is a powerful hybrid of the classical importance sampling and control variate schemes in Monte Carlo. It can, in principle, be applied to any reliability problem where the system function has associated with it a *lower bounding* function and an *upper bounding* function.

Each of the methods mentioned here admits a substantial theoretical development, and each has been extensively used in problems of practical importance.

A Small Example

Consider a network G with four nodes $V = \{1, 2, 3, 4\}$ and five edges $E = \{e_1 = \{1, 2\}, e_2 = \{1, 3\}, e_3 = \{1, 4\}, e_4 = \{2, 3\}, e_5 = \{3, 4\}\}$. Edge e_i operates with probability p_i and fails with probability $1 - p_i$. By listing the operational edges, a state $S \subseteq E$ is formed; since edges operate independently, its probability of occurrence is $\prod_{e_i \in S} p_i \prod_{e_i \in E \setminus S} (1 - p_i)$. Complete state enumeration lists all $32 = 2^5$ states, and sums the probability of occurrences over the operational states or pathsets. For the specific operation of all-terminal reliability, the single state with 5 edges and all 5 states having 4 edges are pathsets; only 8 of the 10 states with 3 edges are pathsets, and no pathset has fewer than 3 operating edges. By considering the failed edges, equivalently there is one cutset with 5 edges; 5 with 4 edges; 10 with 3 edges; and 2 with 2 edges. Minpaths are spanning trees, which are the pathsets with three edges here. There are six

mincuts: $\{e_1, e_4\}$ and $\{e_3, e_5\}$ of size two; and $\{e_1, e_2, e_3\}$, $\{e_2, e_4, e_5\}$, $\{e_2, e_3, e_4\}$, and $\{e_1, e_2, e_5\}$ of size three.

Factoring chooses an edge to delete or contract. Suppose that we factor G on edge e_1 . First we delete e_1 to form $G - e_1$, the network with $V = \{1, 2, 3, 4\}$ and four edges $E = \{e_2 = \{1, 3\}, e_3 = \{1, 4\}, e_4 = \{2, 3\}, e_5 = \{3, 4\}\}$. Edge e_4 must then operate for the network to remain connected, and hence the reliability of $G - e_1$ can be written as p_4 times the reliability of the network G_2 with $V = \{1, 3, 4\}$ and three edges $E = \{e_2 = \{1, 3\}, e_3 = \{1, 4\}, e_5 = \{3, 4\}\}$. On the other hand, contracting e_1 in G yields a network with $V = \{12, 3, 4\}$ and four edges $E = \{e_2 = \{12, 3\}, e_3 = \{12, 4\}, e_4 = \{12, 3\}, e_5 = \{3, 4\}\}$. Edges e_2 and e_4 are now in parallel and can be replaced by a single edge to produce a simpler network G_1 with the same reliability, with $V = \{12, 3, 4\}$ and three edges $E = \{e_3 = \{12, 4\}, e_5 = \{3, 4\}, e_6 = \{12, 3\}\}$, where $p_6 = p_2 + p_4 - p_2 p_4$. The reliability of G is $p_1 \text{Rel}(G_1) + (1 - p_1) p_4 \text{Rel}(G_2)$; in this case, a single factoring step produces two smaller problems, each involving networks with three edges.

To illustrate methods using minpaths, consider the network G_2 . It has three minpaths: $\{e_2, e_3\}$, $\{e_2, e_5\}$, and $\{e_3, e_5\}$. The probabilities that each occurs are $p_2 p_3$, $p_2 p_5$, and $p_3 p_5$, but these are not disjoint events. The probability that the second minpath operates when the first does not can be computed as $(1 - p_3) p_2 p_5$; the probability that the third operates when neither the first nor the second does is $(1 - p_2) p_3 p_5$. These events are disjoint and hence the reliability of G_2 is $p_2 p_3 + (1 - p_3) p_2 p_5 + (1 - p_2) p_3 p_5$, a disjoint sum-of-products form.

In practice these methods employ powerful techniques to factor so as to achieve the most simplification, or to simplify descriptions of the pathsets or cutsets so as to minimize the number of terms in the expression for the reliability.

References

- [1] Colbourn, C.J. (1987). *The Combinatorics of Network Reliability*, Oxford University Press, New York.
- [2] Chartrand, G. & Lesniak, L. (2005). *Graphs & Digraphs*, Chapman & Hall/CRC, Boca Raton.
- [3] Ball, M.O., Colbourn, C.J. & Provan, J.S. (1995). Network reliability, in *Handbook of Operations Research*:

- Network Models*, M.O. Ball, T.L. Magnanti, C.L. Monma & G.L. Nemhauser, eds, Elsevier North-Holland, Amsterdam, 673–762.
- [4] Shier, D.R. (1991). *Network Reliability and Algebraic Structures*, Oxford University Press, New York.
- [5] Ball, M.O. (1986). Computational complexity of network reliability analysis: an overview, *IEEE Transactions on Reliability* **R-35**, 230–239.
- [6] Crama, Y. & Hammer, P.L. (2007). *Boolean Functions: Theory, Algorithms, and Applications*, Cambridge University Press, Cambridge.

Related Articles

Hierarchical Markov Chain Monte Carlo (MC-MC) for Bayesian System Reliability; Parallel, Series, and Series-Parallel Systems; System Reliability: Computational Algebra Methods; System Reliability: Monte Carlo Estimation.

CHARLES J. COLBOURN AND DOUGLAS
R. SHIER

Quality Control, Computing in

Introduction

The quality movement has benefited tremendously from the development of computer software making important but difficult and/or time-consuming calculations readily available to practitioners. When these calculations were performed by hand, the default was typically to resort to simpler tools such as “rules of thumb” for sample size. **Measurement Systems analysis** (MSA) called for 10 parts, 3 operators, and 3 replicates with analysis based on the $\bar{X}$ -bar and range method. Statistical process control (SPC) charts were limited to individuals, $\bar{X}$ -bar and range, or attribute charts. The design and analysis of experiments (DOE) were often simple two-level full or fractional factorials with hand calculations of “average of highs minus average of lows”. These simplified approaches are still valuable for training purposes and may work quite well in practice, but can also produce misleading results due to oversimplification or violation of assumptions.

The following is a brief survey of important advances in **quality control** tools that are (or should be) available to practitioners using current statistical quality software tools, as of 2007:

Power and Sample Size

The power of a test ($1 - \beta$) is the probability of rejecting the null hypothesis, given that the alternative hypothesis is true. Power depends on the type of hypothesis test, and is larger with an increase in sample size (which translates to increased cost), effect size (difference to be detected), **significance level** (α), and decrease in population standard deviation (σ) [1, 2].

Power and sample size calculations should be carried out prior to performing a test of hypothesis or design of experiments. Ideally, when planning an experiment, the power results should be displayed in real time as one selects the design type and number of replicates (using standardized effect sizes). This allows the practitioner to balance the cost of an experiment with expected outcome (see **Sample-Size**

Determination in Experimental Designs; Sample-Size Determination).

The details of power calculations are outside of the scope of this article, but note that the calculations for one-sample z , one proportion, and two proportions are straightforward, with statistical power determined using the same distribution function as that used in the significance test. Power calculations for one-sample t , two-sample t , analysis of variance (ANOVA), two-level **factorial**, and χ^2 require a reference distribution function called a *noncentral distribution function* [1]. Modern computer software takes care of the challenges associated with computing these distributions.

Measurement Systems Analysis

MSA studies (also known as *gauge repeatability and reproducibility*, GRR) are often performed with 10 parts, 3 operators/appraisers, and 3 replicates using the Automotive Industry Action Group (AIAG) standard templates ([3], see **Gauge Repeatability and Reproducibility (R&R) Studies; Measurement Systems Analysis, Overview; Measurement Systems Analysis, Capability Measures for**). Unfortunately, these sample sizes typically yield wide **confidence intervals** on GRR statistics such as %GRR (%precision/total) and %precision/tolerance. Without considering the confidence intervals, practitioners run the risk of calling an unacceptable measurement system acceptable (or an acceptable measurement system unacceptable).

Computing confidence intervals for GRR statistics requires the calculation of variance components (see **Variance Components; Gauge Repeatability and Reproducibility (R&R), Variance Components in**), and the confidence intervals associated with these variance components and ratios (see [4]; **Gauge Repeatability and Reproducibility (R&R) Studies, Confidence Intervals for**). Computer software is available to perform these computations.

Statistical Process Control

The SPC chart developed by Walter Shewhart has proved to be an effective and essential tool for quality improvement. As a statistical tool, however, there are times when the **control chart** fails to give proper information about the process. Sometimes

there are false alarms leading to expensive and fruitless searches for assignable causes. On the other hand, the Shewhart chart occasionally fails to detect a significant process shift early enough to prevent rework. These problems are particularly acute in process industries.

In order for control charts to function properly, the underlying assumptions must be valid. Specifically, the observations are assumed to be normally distributed, independent, with fixed mean and constant variance. If one is using an $\bar{X}$ -bar chart, the assumption of normality is usually met because of the central limit theorem – the distribution of sample averages tends to be normal, even if the individual observations are not. If, however, the subgroup size, $n = 1$, and the observations are not normally distributed, a power transformation should be applied to the data. Suggested transformations include X^2 , $\sqrt{X}$, $\ln(X)$, $1/\sqrt{X}$, $1/X$, $1/X^2$.

A power transformation resulting in normality should be utilized. The transformed data is then used to calculate the control limits. Computer software simplifies this task by employing the Box–Cox transformation method (see **Box and Cox Transformation**). If a suitable power transformation cannot be found, an alternative technique such as the Johnson transformation may be used [5].

The requirement for samples to be statistically independent can be a much more serious problem, since there is no similar mechanism to the central limit theorem assuring reasonable results. The presence of serial autocorrelation in data will adversely affect the performance of a Shewhart chart (see **Autocorrelation Function; Autocorrelated Data**). In the process industries, autocorrelation is typically positive due to inertial elements, resulting in false alarms. If, however, the autocorrelation is negative, the Shewhart limits will be too wide, and significant process shifts will not be detected.

The key to dealing with this autocorrelation is to directly model it with **time series** analysis and then use that model to remove the autocorrelation from the data. A Shewhart control chart can then be applied to the residual prediction errors in order to detect significant process disturbances. The exponentially weighted moving average (EWMA); see **Exponentially Weighted Moving Average (EWMA)**; **Exponentially Weighted Moving Average (EWMA) Control Chart** chart can be used as a one-step-ahead

predictor, so it is a valid model to consider when dealing with autocorrelation. The EWMA is an important member of the time series models identified by Box and Jenkins as autoregressive integrated moving average (ARIMA) models (see [6, 7]; **Large Autoregressive Integrated Moving Average (ARIMA) Modeling**). It is equivalent to an ARIMA(0, 1, 1). Practically, the EWMA is a good prediction model in the majority of cases. If, however, the autocorrelation is negative or the mean is drifting too quickly, a different model will be required. Other ARIMA models should be considered using computer software.

Design and Analysis of Experiments

Computer software greatly simplifies the design and analysis of experiments, but the use of computer-generated designs is especially important in situations where standard factorial or **fractional factorial** cannot be utilized due to factor constraints. D-optimal designs are the most widely used type of computer-generated designs (see **Optimal Design**).

Many designed experiments also require the simultaneous optimization of multiple responses. Typically, an algorithm is used to maximize a desirability function to find the optimum settings of the X factors. However, as the number of factors or responses increases, conventional optimization algorithms tend to locate a local optimum rather than the global optimum.

Recent developments using genetic algorithms bring a new level of design efficiency that works well even with highly constrained factor regions [8]. Genetic algorithms are also very useful in multiple response optimization for locating the global optimum [9]. This translates into more effective experiments run at a lower cost.

References

- [1] Thomas, L. & Krebs, C.J. (1997). A review of statistical power analysis software, *Bulletin of the Ecological Society of America* **78**(2), 126–139.
- [2] Lenth, R.V. (2001). Some practical guidelines for effective sample size determination, *The American Statistician* **55**(3), 187–193.
- [3] Chrysler Corporation, Ford Motor Company, General Motors Corporation. (2002). *Measurement Systems Analysis Reference Manual*, 3rd Edition, Automotive Industry Action Group (AIAG), Detroit.

-
- [4] Burdick, R.K., Borror, C.M. & Montgomery, D.C. (2005). *Design and Analysis of Gauge R&R Studies: Making Decisions with Confidence Intervals in Random and Mixed ANOVA Models*, ASA-SIAM Series on Statistics and Applied Probability, Society for Industrial and Applied Mathematics (SIAM), Philadelphia, American Statistical Association (ASA).
 - [5] Chou, Y., Polansky, A.M. & Mason, R.L. (1998). Transforming non-normal data to normality in statistical process control, *Journal of Quality Technology* **30**(2), 133–141.
 - [6] Hunter, J.S. (1986). The exponentially weighted moving average, *Journal of Quality Technology* **18**(4), 203–210.
 - [7] Box, G.E.P. & Jenkins, G.M. (1976). *Time Series Analysis, Forecasting, and Control*, 2nd edition, Holden-Day.
 - [8] Heredia-Langner, A., Carlyle, W.M., Montgomery, D.C., Borror, C.M. & Runger, G.C. (2003). Genetic algorithms for the construction of D-optimal designs, *Journal of Quality Technology* **35**(1), 28–46.
 - [9] Ortiz, F., Jr, Simpson, J.R., Pignatiello, J.J., Jr & Heredia-Langner, A. (2004). A genetic algorithm approach to multiple-response optimization, *Journal of Quality Technology* **36**(4), 432–450.

JOHN NOGUERA

Data Mining in Quality and Reliability

Background

Few applications of data mining in quality control, quality assurance, or reliability have been published. In April 1996, in *Quality Progress*, Bert Gunter [1] urged caution in the use of data mining based on the extraordinary amount of hype and false promises it was receiving at the time. Focusing on formal experimental design, his article validated the methodologies used in data mining but contrasted them with the ability to formulate and test hypotheses using standard statistical techniques.

Since the publication of this article, it has become much easier to collect and store large amounts of data. Volumes of data may be collected from continuous operation of machines, from web logs and web transactions, and from healthcare studies involving claims data or long-term evaluations of patients. In addition to this, computing power has increased exponentially, making it possible to process millions of data points easily and quickly. Articles discussing these and other database aspects of data mining include **Processing of Survey Data, Statistics in Data and Information Quality, and Data Mining, Evaluation Techniques in**. Standard statistical techniques may not be meaningful when applied to these enormous databases because all comparisons with the mean will be significant and standard statistical measures of variability will be extremely small (*see Descriptive Statistics*). Data mining is fundamentally different from statistical analysis but in some ways it resembles some **exploratory data analysis** techniques [2]. However, it strives to identify patterns within the data using both continuous and categorical data. The results of data mining are often measured by “interestingness”. This can be defined as finding patterns in the data that are not obvious. In relating this to examining quality, this may mean finding previously unidentified anomalies in the data that may contribute to a process being out of control or previously undiscovered bias or **outliers** in the data.

In 1998, in *Technometrics*, Gerry Hahn and Roger Hoerl [3] examined the changing role of statistics (and statisticians) in business, citing enormous databases and finding proactive statistical process

control (SPC) methods for understanding and analyzing the questions and issues that arise with these data. While there has been enormous growth in the size of databases, research articles using data mining methods for quality have not followed.

This article uses a large healthcare trial in which data mining techniques provided the only way to analyze the dataset for **data quality**, to identify anomalies and nonrandomness, and to examine lack of homogeneity.

Overview of Supervised and Unsupervised Data Mining Techniques

Data mining techniques can be separated into two primary areas: those using supervised algorithms and those using unsupervised algorithms. Supervised data mining techniques most closely resemble standard statistical modeling because they are predictive with an output variable being predicted by many other variables. Unsupervised data mining techniques are associative. They are used to examine relationships between variables and may or may not have a direction specified with antecedent and consequent variables identified. Apart from traditional correlations between variables, these associations may include variables and data of any type, nominal (including dummy variables), ordinal, interval, and ratio.

Supervised Techniques

These include classification analyses and decision trees (C5.0, classification and regression trees (CaRT), and chi-squared automatic interaction detector (CHAID)), logistic regression, and neural nets. CaRT decision trees and neural nets are discussed at length in **Classification and Regression Tree Methods, Neural Networks: Construction and Evaluation and Neural Networks in Statistical Process Control**.

Decision trees using data mining algorithms differ from their statistical counterparts in that they do not have prespecified splits or branches. Using the C5.0 classification algorithm, the branch splitting is data driven and the split giving the most separation between two (or more) groups is the variable identified first. The tree may be asymmetrical in that variable splits for one branch of the tree may not appear in another branch or may appear at a different level.

This technique will combine groups of a categorical variable if they are so similar that additional branches are unnecessary. This algorithm, the most common in data mining software, uses maximum gain as its criteria for splitting into a new set of branches. It chooses the variable that has the largest expected information gain. In other words, it selects the attribute that will result in the smallest expected size of the subtrees. The definition of information gain is as follows: consider the quantity $\text{Gain}(X, T)$ defined as $\text{Gain}(X, T) = \text{Info}(T) - \text{Info}(X, T)$. This represents the difference between information needed to identify an element of T and information needed to identify an element of T after the value of attribute X has been obtained, that is, this is the gain in information due to attribute X . CaRT identifies the best classification using logistic regression techniques, and CHAID bases its algorithm on the χ^2 distribution. Most tree algorithms require that the target (dependent) variable be categorical. Such algorithms require that continuous variables be binned for use with regression. The most notable exception to this is CaRT, which handles continuous target variables directly.

Neural nets and logistic regression models are similar to each other and resemble least-squares regression in the interpretation of output. Logistic regression models provide coefficients, goodness-of-fit tests, and significance levels while neural nets list the percentage contribution of each variable to the targeted outcome variables. Neural nets are a “black box” algorithm and the output gives no indication of the form of the resulting model.

Unsupervised Techniques

If a predictive model is the desired outcome of an experiment, unsupervised data mining may be the first step in understanding the associations in the data. Primary unsupervised data mining techniques include association webs and Kohonen maps, *a priori* and generalized rule induction (GRI) methods, and cluster and anomaly analyses. Association webs are used solely with categorical data and provide the strength of the links between variables (either as a percentage or an absolute number). The darkest and widest lines connecting two variables (or categories within variables) provide evidence of stronger links and interactions. These webs are often used to construct interactions between variables prior to a supervised model. Kohonen maps provide a multivariate method

for examining associations. The Kohonen algorithm is a type of neural net that creates two (or more) variables from all the variables available and plots important characteristics of the dataset on these two (or more) dimensions. If patterns exist, the data are often clustered in different areas of the plot.

Cluster analysis is a standard multivariate statistical technique that is also a commonly used tool in data mining. It differs from its statistical counterpart by not providing tests between clusters and by allowing the inclusion of any type of variable into the analysis. Cluster analysis can be subdivided into three common data mining algorithms: two-step cluster analysis, where the algorithm will find the best number of clusters; *k*-means cluster analysis, where the number of clusters must be prespecified, and anomaly analyses, where the data points that are furthest from the mean of the cluster are identified as not fitting into any cluster easily. Anomaly analysis is of immediate use and application in quality control for processes with large amounts of data and many factors. It can identify multivariate outliers for error detection and analysis. Of all the data mining techniques outlined, anomaly detection has the easiest and immediate application in the field of quality.

Two other association techniques are *a priori* analysis and GRI. These two techniques are similar in their goal of creating rules in the data with antecedent variables and consequent variables. Often used in analyses of heterogeneous objects such as market baskets, the formation of these rules is based on the support the rule has in the data and the confidence level within the support for the rule. Support is defined as the probability of finding all the variables of interest in the dataset or the proportion of the data containing all the values of these variables. Confidence is defined as the conditional probability of finding one value of one of the variables given that the other variables are also contained. For example, in a sample of defective items, three characteristics of defects, called *A*, *B*, and *C*, exist in 30% of the sample. This is the support for the three characteristics. However, further analysis might show that given that defects *A* and *B* exist, defect *C* exists in 80% of these items. Both algorithms require that the defect of interest, in this case the consequent defect *C*, be categorical. In *a priori* analyses, the antecedents, *A* and *B* must also be categorical. In a GRI the antecedents may be of any type.

Time Series Data Mining

Many assessments of quality have a temporal component and this can be included in a data mining analysis using either supervised or unsupervised methods. In supervised models such as neural nets or logistic regression, time can be included as a variable in the model. In unsupervised models, a transaction ID may be used instead of a variable for time. For example, if a production process on multiple machines is examined over an arbitrary or fixed time interval, the machine ID may be used to gauge the quality of items produced by that machine over time. Similarly, multiple purchases over time can be tracked by customer reward card numbers to examine patterns of behavior and store efficiency.

Example: Using Data Mining for Auditing Data Quality

A multistate study was designed to compare the care given to patients by fee-for-service (FFS) providers with that given by managed care (MC) providers. The study was carried out in multiple sites within a state and over multiple states. Patients entered the study as FFS patients and data was collected on services provided for 1 year prior to being randomized to continue to receive care through a FFS provider or being switched to an MC provider for care. After **randomization**, patients were followed for an additional 24 months. The service provided and the outcomes of care were tracked through insurance claims data. The final database contained over 60 000 patients and 4 million claims. With a database of this size, traditional statistical analyses are not useful and examining patterns of care using data mining was implemented.

Analysis of Homogeneity at Baseline

Homogeneity was initially assumed for patient characteristics across treatment conditions because of strict randomization and, because treatment was standardized by government protocols, patient care within treatment conditions as well as within and between states was initially assumed to be homogenous. However, prior to the analysis of outcomes of the study, these assumptions were investigated. With the enormous number of patients in the database, any strict

statistical comparison of baseline demographics and diagnostic data would be statistically significant so several data mining techniques were used. A subset of data from three states is used for illustration.

Several Analyses to Evaluate the Quality of the Study Randomization.

1. To examine the randomization over all sites, a web and a Kohonen map were constructed.^a The web, shown in Figure 1, plots the strongest connections between baseline variables (gender, race, and diagnosis) and the randomization factor (FFS or MC). The assumption prior to examining the web was that the randomization was balanced with respect to these variables, but examining Figure 1 shows that more women and nonwhites were randomized to the MC providers. Men and whites did not show up with strong links to either MC or FFS provider, this indicates balance with respect to those variables. The Kohonen map in Figure 2 confirms that there are strong patterns in the database separating into racial and gender groups. Table 1 gives a summary of the strong links found in these analyses.
2. Another useful tool for identifying unusual patterns is an anomaly cluster analysis. Figure 3 identifies four clusters within the data and many unusual patients within each cluster. A common pattern of gender and race exists in several of

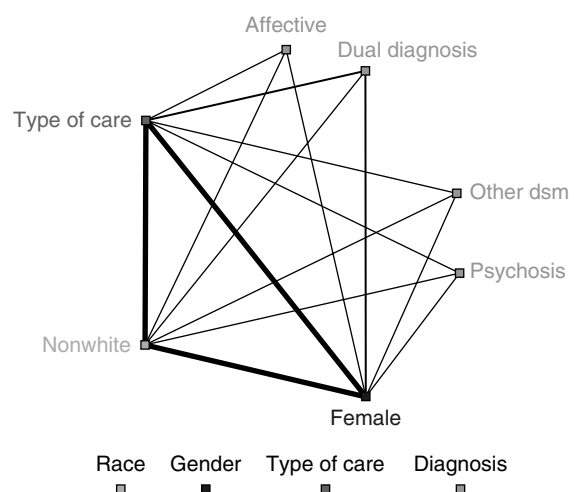


Figure 1 Web of strong links between managed care, nonwhite, and female

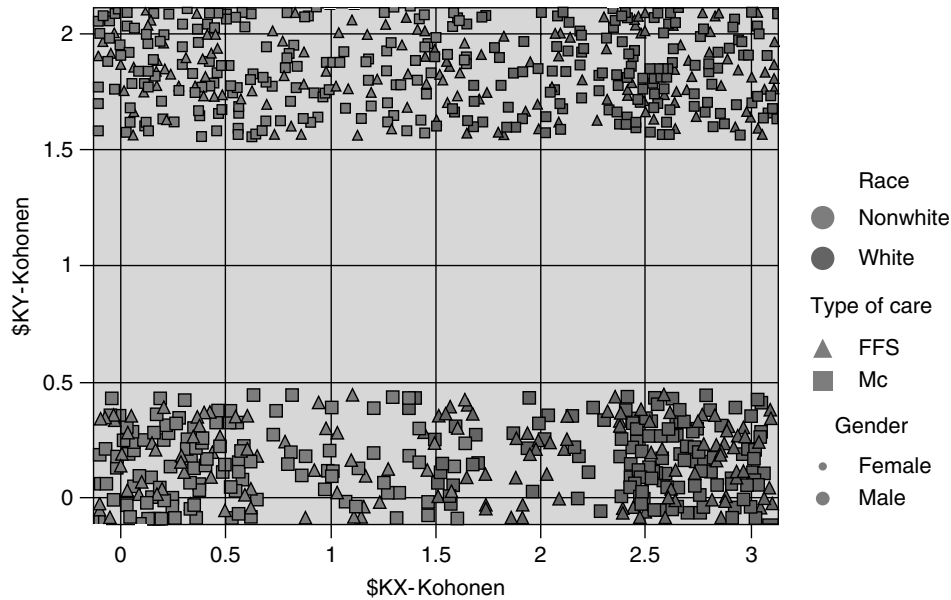


Figure 2 Kohonen map examining race and gender

Table 1 Summary of strong links between variables

Strong Links	Field 1	Field 2
Links		
14 933	Gender = "female"	Type of care = "MC"
13 517	Race = "nonwhite"	Type of care = "MC"
12 036	Gender = "female"	Race = "nonwhite"
8919	Diagnosis = "dual diagnosis"	Type of care = "MC"
7621	Gender = "female"	Diagnosis = "dual diagnosis"
6808	Race = "nonwhite"	Diagnosis = "dual diagnosis"
6595	Diagnosis = "psychosis"	Type of care = "MC"
6128	Diagnosis = "affective"	Type of care = "MC"
5563	Gender = "female"	Diagnosis = "psychosis"
5536	Race = "nonwhite"	Diagnosis = "psychosis"
5413	Gender = "female"	Diagnosis = "affective"
5071	Diagnosis = "other dsm"	Type of care = "MC"
4850	Gender = "female"	Diagnosis = "other dsm"
4689	Race = "nonwhite"	Diagnosis = "affective"
4155	Race = "nonwhite"	Diagnosis = "other dsm"

the clusters confirming the web and Kohonen maps.

- After identifying that the study randomization was biased, a state-by-state and site-by-site analysis was undertaken using standard statistical analysis^b comparing gender balance and racial balance by site and type of care. Figure 4 gives an example of the site analysis showing a plot of

females for three states where it is clear that State 3 had significantly unbalanced randomization of females into MC ($p < 0.001$).

- Solution to this imbalance included modeling the data with baseline characteristics as covariates, as opposed to pooling the study data and doing sensitivity analyses with the imbalanced site excluded. Without the initial association

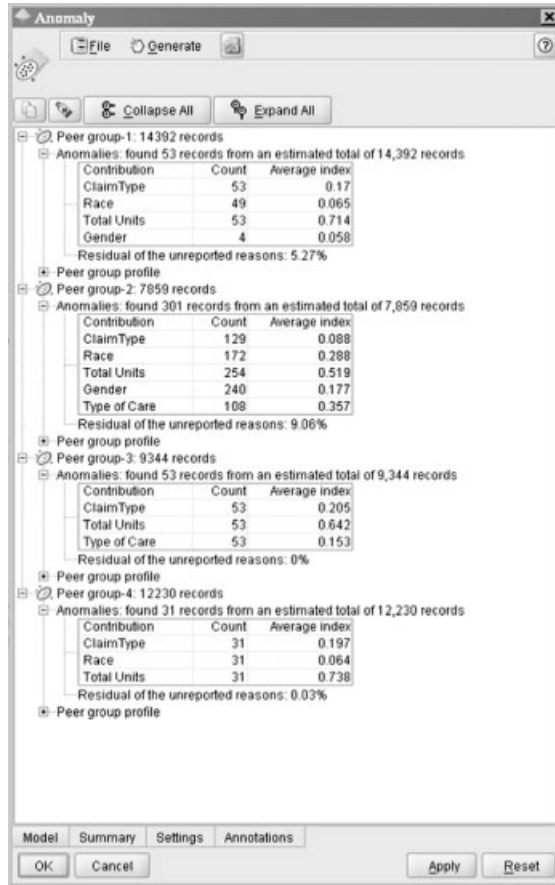


Figure 3 Anomaly cluster analysis

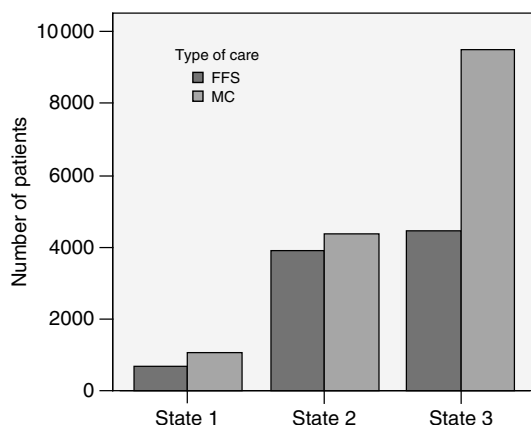


Figure 4 Females by state by type of care

analysis, this imbalance might not have been identified.

Examination of Patterns of within and between State and Site Differences

Prior to analyzing patterns in types of claims comparing FFS and MC, it was important to examine whether there were unusual patterns of service between states that could not be explained by normal variability of service. All sites followed a government-developed protocol for delivery of service that should have minimized the variability of service delivery. With over 4 million claims and 60,000 patients, using an analysis of variance approach (the standard statistical methodology for examining site differences) would have shown every comparison to be significantly different between states. Here again, a subset of three states was used.

Procedure to Evaluate the Service Heterogeneity over States.

1. A CaRT was utilized to develop a tree model to discriminate between states. This is a classification model that uses logistic regression to separate into branches. The branches are defined in the model to be binary splits of the data and each branch shows the split by the outcome variable (the state in these models). Figure 5 shows the tree with outpatient therapy imbalanced by state. While State 3 has more patients, it has over 90% of the claims for outpatient therapy.
2. After identifying that State 3 had a disproportionate number of claims for outpatient therapy, a state-by-state analysis was undertaken for each service separately using standard statistical analysis comparing type of care by type of service by state. Figure 6 gives an example of the state analysis showing a plot of outpatient therapy service claims for three states where it is clear that State 3 had significantly more claims in FFS and MC, but particularly more so in MC ($p < 0.00001$).
3. Upon further investigation of State 3, it was found that two sites within State 3 were responsible for these excess claims, and were submitting duplicate and fraudulent Medicare claims.

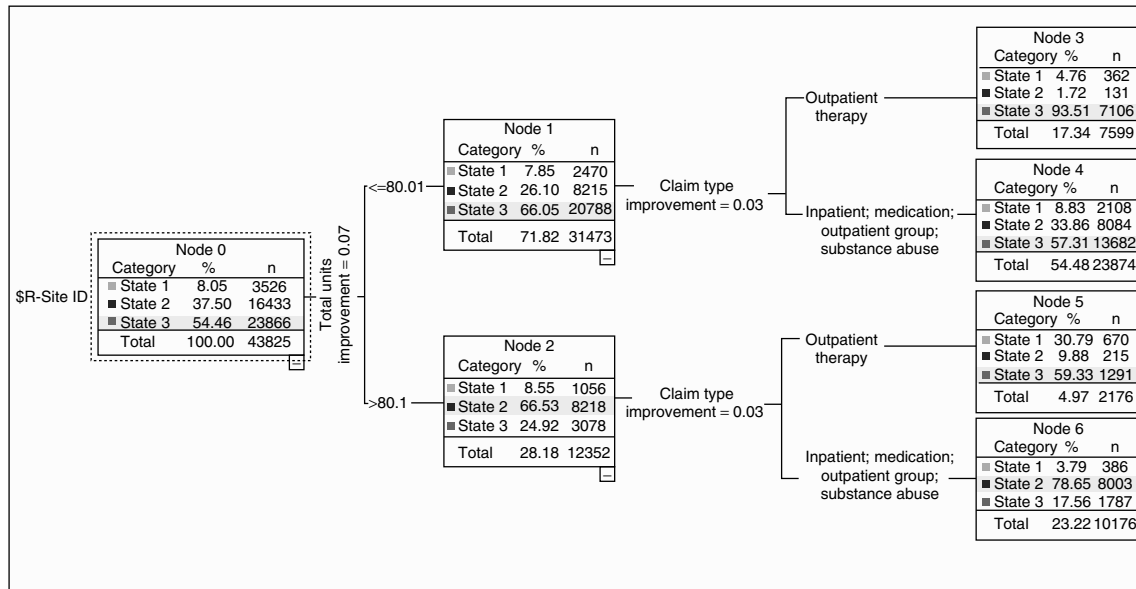


Figure 5 CaRT interactive tree showing State 3 with disproportionate claims for outpatient therapy

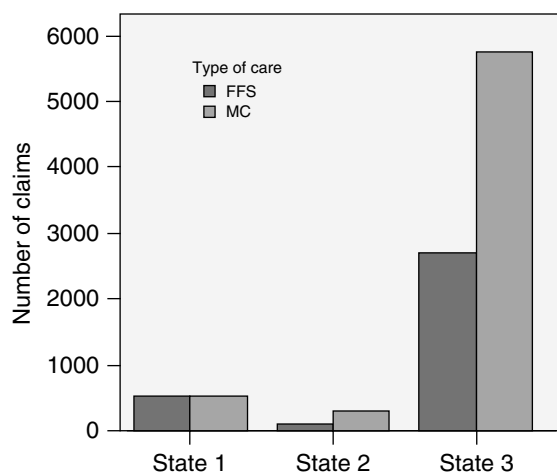


Figure 6 Excess claims filed by state by type of claim

Summary

Data mining may be the first line of analysis when a database is very large. Designed to identify unusual patterns, it can be used well for finding departures from homogeneity or identifying outliers requiring further exploration. As illustrated in the example above, the results can lead to surprising conclusions.

End Notes

- All data mining analyses were done in Clementine 10.1 from SPSS Inc., Chicago, Illinois.
- All statistical analyses were done in SPSS version 14.0 for Windows.

References

- Gunter, B. (1996). Data mining: mother lode or fools gold? *Quality Progress* **29**(4), 113.
- Hand, D.J. (1998). Data mining: statistics and more? *The American Statistician* **52**, 2.
- Hahn, G. & Hoerl, R. (1998). Key challenges for statisticians in business and industry, *Technometrics* **40**(3), 195.

Further Reading

- Adams, L. (2001). Mining the world of quality data, *Quality Magazine* (<http://www.qualitymag.com>).
- Courtheoux, R.J. (2003). Marketing data analysis and data quality management, *Journal of Targeting, Measurement and Analysis for Marketing* **11**(4), 299.
- Giudici, P. (2003). *Applied Data Mining: Statistical Methods for Business and Industry*, John Wiley & Sons, Chichester.
- Jiang, W. (2004). A joint monitoring scheme for automatically controlled processes, *IIE Transactions* **36**, 1201.
- Larose, D.T. (2006). *Data Mining Methods and Models*, Wiley-Interscience, John Wiley & Sons.

Spadola, T. (2003). Steering the course, *Best's Review* 108.

Wells, D. (2000). Quality assurances, *Enterprise Systems Journal* 16.

Whiting, R. (2005). Trend spotters: finding needles in haystacks is getting easier, *InformationWeek* 1054.

ELAINE ALLEN AND CHRIS SEAMAN

Density and Failure Rate Estimation

Introduction

In reliability and quality control, density and failure rate (or hazard) functions provide valuable information about the distribution of failure times. Shapes of density and failure rate functions may contain characteristics of interest, such as number, location, and features of modes; monotone increasing or decreasing features; and information about tail behavior. One can distinguish two basic statistical approaches for estimating density and failure rate functions, given a sample of observed failure times: parametric and nonparametric.

In the parametric approach to **estimation** and model fitting, the starting point is to assume an underlying parametric model which specifies a distribution of the observed data, typically a lifetime distribution such as Weibull, Gompertz, or gamma. The parameters are then usually identified by maximum likelihood (*see* **Maximum Likelihood**) or a Bayesian method (*see* **Bayesian Reliability Analysis**).

The parameters determine the distribution uniquely, and thus also density and failure rate functions. The main advantage of the parametric approach, which we will not discuss further in this article, is ease of inference, efficiency, and the absence of a **bias** issue, provided that the underlying parametric model assumption is correct; the latter assumption requires a leap of faith in many applied situations. In fact the main disadvantage of parametric models is that indeed the underlying distribution must be fully known: under model misspecification, parametric approaches result in inconsistent estimators, and the associated asymptotic bias leads to invalid inference.

The nonparametric or smoothing approach is considerably more flexible. Nothing more than basic smoothness assumptions on the distribution function are needed. Therefore this approach “lets the data speak for themselves”, without imposing hard-to-verify model assumptions. It is thus particularly well suited for exploratory data analysis, and is also used for final analysis in those situations where density functions or failure rates cannot be clearly specified, or when doubts remain about these specifications.

The main disadvantages of nonparametric smoothing methods are that inference, often through confidence bands, is not straightforward, and that a smoothing parameter needs to be specified. While these methods are asymptotically unbiased under weak assumptions, their application often requires one to be aware of and deal with a finite bias issue.

As was already pointed out in Rosenblatt’s seminal 1956 paper, nonparametric density estimates are finitely biased [1]. In parametric approaches, potential bias issues are mostly addressed indirectly through goodness-of-fit checking, and if such bias is detected, the model needs to be changed, often through a data transformation. One can argue that the explicit nature of the bias issue in nonparametric smoothing approaches is actually an advantage, as this bias is small, can be adjusted for, and disappears asymptotically. Although parametric models rarely fit exactly, the resulting bias is often not adequately addressed. For the multivariate case, the so-called curse of dimensionality leads to rapidly declining convergence rates of nonparametric methods as the dimensions increase. This has led to the development of various dimension reduction and semiparametric approaches [2]. In such models, smooth densities are included for low-dimensional parts of the model, while dimension reduction is achieved through a parametric component.

Nonparametric Density Estimation

We begin with basic definitions and relations. Denoting a nonnegative random variable that represents failure time or lifetime by T , the survival function is

$$\bar{F}(t) = P(T \geq t), \quad t \geq 0 \quad (1)$$

where $\bar{F}(t) = 1 - F(t)$ is the cumulative distribution function of the lifetime distribution. If this function is differentiable, we may define the **probability density function** as

$$f(t) = \lim_{\Delta \rightarrow 0} \frac{1}{\Delta} P(t \leq T < t + \Delta), \quad t \geq 0 \quad (2)$$

and the hazard or failure rate function as

$$\lambda(t) = \lim_{\Delta \rightarrow 0} \frac{1}{\Delta} P(t + \Delta > T \geq t | T \geq t) = \frac{f(t)}{\bar{F}(t)}$$

2 Density and Failure Rate Estimation

(3)

The importance of the failure rate λ in reliability is due to its interpretation as a risk function. Specifically, $\lambda(t)$ quantifies imminent risk of failure at time t , and is related with the survival function $\bar{F}$ through

$$\begin{aligned}\bar{F}(t) &= \exp \left[- \int_0^t \lambda(u) du \right], \\ \lambda(t) &= - \frac{d}{dt} \log [\bar{F}(t)]\end{aligned}\quad (4)$$

Both density and failure rate functions characterize the failure time distribution. The transformations from density to failure rate and *vice versa* are as follows [3]:

$$\begin{aligned}\lambda(t) &= \frac{f(t)}{1 - \int_0^t f(u) du}, \\ f(t) &= \lambda(t) \exp \left[- \int_0^t \lambda(u) du \right]\end{aligned}\quad (5)$$

Basic properties of densities f , failure rates λ , and the cumulative hazard rate $\Lambda(t) = \int_0^t \lambda(s) ds$ are

$$\begin{aligned}f &\geq 0, \quad \int f(u) du = 1, \\ \lambda &\geq 0, \quad \Lambda(t) \rightarrow \infty \text{ as } t \rightarrow \infty\end{aligned}\quad (6)$$

Turning to the nonparametric estimation of these functions, historically the first density estimate is the histogram which is still much in use. Other density estimates can be viewed as extensions of the histogram. Besides being of interest in assessing shape, especially symmetry, multimodality, and tails of a distribution, density estimates are also useful for a variety of specific statistical tasks, such as the smoothed **bootstrap**, the construction of **confidence intervals** of estimated **quantiles**, e.g., the median, for obtaining Fisher information, or for classification, as centers of clusters may correspond to modes in a density estimate and often multimodality in a density is of interest.

We proceed with a formal definition of the two most common nonparametric density estimates, histograms and kernel estimators, starting with a sample $X_1, \dots, X_n$ of independent and identically distributed (i.i.d.) univariate data whose distribution has density

f . To construct a *histogram*, one divides the range of the data into m bins $B_1, B_2, \dots, B_m$, which are usually chosen to be of the same size but can also vary in size, and then obtains the histogram

$$\begin{aligned}\hat{f}_{\text{hist}}(x) \\ = \frac{1}{n} \sum_{j=1}^m \frac{(\text{number of data } X_i \text{ falling into bin } B_j)}{(\text{size of bin } B_j)}\end{aligned}\quad (7)$$

The smoothing parameters are the leftmost endpoint of the first bin (which does not matter much) and the bin size (which matters). Advantages of the histogram are that it is straightforward to construct and understand, and that it is ubiquitous in statistical packages. Disadvantages are the relatively slow convergence with increasing sample size and the discontinuity of these density estimates, which is at variance with the assumed smoothness of the underlying density.

Both disadvantages are remedied by kernel density estimates. These can be conceived as generalizations of sliding histograms, or as a convolution of the empirical distribution function with a smooth kernel function,

$$\begin{aligned}\hat{f}_K(x) &= \frac{1}{nh} \sum_{i=1}^n K \left(\frac{x - X_i}{h} \right) \\ &= \int \frac{1}{h} K \left(\frac{x - u}{h} \right) dF_n(u)\end{aligned}\quad (8)$$

where $h = h(n)$ is a sequence of bandwidths or smoothing parameters, K is a kernel function, and dF_n is the empirical measure, which assigns mass $1/n$ to every observation point X_i . The kernel method for density estimation was introduced by Fix and Hodges [4], Rosenblatt [1], and Parzen [5]. The best source for basic asymptotic and finite sample properties of kernel density estimators remains the monograph of Silverman [6], along with the monograph by Scott [7] for the multivariate case. For general background on kernel methods, see [8].

Typical basic results are rates of convergence of (integrated) mean squared error (IMSE or MSE) and asymptotic distributions. Optimal rates of MSE convergence are achieved by kernel estimators and are slower than the parametric rate of n^{-1} , due to the nonparametric nature of these estimators. For example, if the underlying density is twice continuously differentiable, the sequence of bandwidths $h = h(n)$

satisfies $h \rightarrow 0$ and $nh \rightarrow \infty$ as $n \rightarrow \infty$, and the kernel satisfies $\int K(x) dx = 1$, $\int K(x)x dx = 0$, and some other basic regularity conditions, then using an IMSE-minimizing bandwidth sequence, the IMSE of the kernel density estimator satisfies

$$\begin{aligned} & \int E (\hat{f}_K(x) - f(x))^2 dx \\ &= n^{-4/5} \left[\left(\int f^{(2)}(x)^2 dx \right) \left\{ \left(\int K(u)u^2 du \right) \right. \right. \\ & \quad \left. \left. \times \left(\int K(u)^2 du \right)^2 \right\} \right]^{1/5} + o(n^{-4/5}) \end{aligned} \quad (9)$$

The IMSE rate of convergence of $n^{-4/5}$ is better than the corresponding rate $n^{-2/3}$ for histograms.

One of the main issues for nonparametric smoothing methods is the choice of the smoothing parameter. There exists a considerable literature on bandwidth selection for kernel density estimation (*cf.* the overview in Sheather [9] and also in [10]); for the related problem of kernel choice see [11]. Smoothing parameters need to be chosen such that both undersmoothing, associated with large variance and small bias, and oversmoothing, associated with large bias and small variance, are avoided and instead, the optimal trade-off between variance and bias is found.

Visual choice remains a common selection method; often done interactively by increasing a clearly undersmoothing bandwidth, until the first sign of oversmoothing appears. When comparing to a histogram with small bin width, “residuals” defined by differences between histogram values and corresponding fitted density values at bin midpoints start to show signs of lack of fit when oversmoothing sets in. Popular automatic methods are cross-validation and plug-in methods. The main obstacle for data-adaptive selection methods is that the theoretically optimal bandwidth depends on the unknown density, as seen for a special case in equation (3). The best currently known methods are plug-in methods, where these unknown quantities are estimated and substituted. Preliminary estimates of the unknown quantities include estimates of density derivatives. These preliminary estimates are also nonparametric, and their smoothing bandwidths are obtained in turn by estimating further unknown density-dependent

quantities with pre-preliminary estimators; the bandwidths of the preestimators ultimately are obtained by making parametric (usually Gaussian) assumptions on the unknown densities.

A special problem for estimating densities with compact support are boundary effects that arise when estimating at or near the endpoints. These can cause problems when a structure of interest occurs near such endpoints. One remedy is the use of special boundary adjustments such as boundary kernels as described in [12]. The basic density estimation schemes can be generalized to cover the estimation of multivariate densities, of density derivatives, and of densities obtained from incomplete or length-biased data. Other variants that have been studied include density estimation combined with data transformations, density estimation with locally varying bandwidths, density estimation under shape constraints such as unimodality, and the estimation of functionals of densities and of conditional densities. For reliability, applications of density estimation in renewal theory are of particular interest [13–15].

Alternative density estimation schemes include orthogonal series estimators, penalized maximum-likelihood estimators, and various estimation schemes based on the smoothing of histograms. Typically the histograms subjected to a subsequent smoothing step will be strongly undersmoothed and are constructed with very small bins, so that many of the bins remain empty. The midpoints of the m bins x_j and corresponding heights of the histogram bars y_j constitute a scatterplot $\{(x_j, y_j), j = 1, \dots, m\}$ to which a nonparametric regression-type smoother is applied. This smoother can be based on smoothing **splines** or B splines. An easy-to-implement and explicit smoothing scheme that also automatically adjusts for boundary effects is provided by local linear fitting [16], where one uses a nonnegative kernel function K and a bandwidth h as defined above to obtain the minimizers of

$$\sum_{j=1}^m K\left(\frac{t - x_j}{h}\right) [y_j - \{a_0 + a_1(t - x_j)\}]^2 \quad (10)$$

w.r.t. a_0, a_1 . Then the density estimate is $\hat{f}_L(t) = \hat{a}_0(t)$, the estimated intercept at each t . The problem of adjusting density estimates that do not necessarily display the properties of a density (nonnegativity and integrating to 1) have been well studied (e.g., in [17]).

Failure Rate Estimation

The failure rate $\lambda(t) dt$ represents the instantaneous chance that an equipment fails in the interval $(t, t + dt)$, given that it works at age t . The hazard or failure rate function $\lambda(t)$ then provides a trajectory of risk. Nonparametric failure rate estimation often involves the smoothing of an initial failure rate estimate. Available data may be either grouped, in which case failures are reported aggregated over time intervals, or may consist of continuously observed failure times. For continuous data, the analogy of the relationships $F(t) = \int_0^t f(x) dx$ between distribution function F and density f , and $\Lambda(t) = \int_0^t \lambda(x) dx$ between cumulative hazard rate Λ and failure rate function λ motivates kernel estimates for λ and demonstrates the closeness of failure rate and density estimation problems.

One complication that one often faces with lifetime data is that available data are incomplete, due to failure to follow-up or termination of a study. The most commonly encountered type of incompleteness is random censoring [3]. In random censoring, we assume there is a censoring random variable C_i such that observed data are $X_i = \min(T_i, C_i)$, the minimum of the lifetime T_i and censoring time C_i of the i th individual. One also observes the censoring indicator $\delta_i = 1_{\{X_i = T_i\}}$, which is one if the actual lifetime is observed and zero otherwise. Let $(X_{(i)}, \delta_{[i]})$, $i = 1, 2, \dots, n$, be the ordered sample with respect to lifetimes X_i , i.e., $X_{(1)} \leq X_{(2)} \leq \dots \leq X_{(n)}$, with $\delta_{[i]}$ the concomitant censoring indicator of $X_{(i)}$.

Hazard estimators can then be obtained by smoothing the increments of the Nelson–Aalen estimator [18]

$$\Lambda_n(t) = \sum_{i=1}^n \frac{\delta_{[i]} 1_{\{X_{(i)} \leq t\}}}{n - i + 1} \quad (11)$$

of the cumulative hazard function $\Lambda(t)$ (if there are no tied observations). Using the kernel method, we arrive at the kernel failure rate function estimator

$$\hat{\lambda}(t) = \sum_{i=1}^n \frac{1}{h} K\left(\frac{t - X_{(i)}}{h}\right) \frac{\delta_{[i]}}{n - i + 1} \quad (12)$$

if there are no tied observations. Bias and variance issues can be similarly analyzed as for density estimates.

Boundary-corrected kernel estimators and practical bandwidth choices have been considered in [12]

and bootstrap confidence intervals in [19]. For the closely related problem of **intensity function** estimation with applications in reliability, see [20, 21]. Other smoothing methods that have been considered include spline fitting with penalized likelihood [22].

For grouped data, one problem is the aggregation, which can pose problems for nonparametric failure rate estimation. The reason is that the failure rate is defined as a limit. However, nonparametric estimation of failure rates works reasonably well when based on a combination of smoothing central rates of failure and a suitable transformation [23]. Assume the data are aggregated into intervals or bins of length Δ , the number of failures is d_i in the i th such interval, and the number of items in working order and at risk of failure is n_i at the beginning and n_{i+1} at the end of the i th such interval. Then the central failure rate for the i th interval is defined as $q_{c_i} = 2d_i/(n_i + n_{i+1})$.

Denoting the midpoint of the i th interval or bin by t_i , the smoothed central failure rate $\hat{q}_{c_i}$ is obtained by smoothing the data $\{t_i, q_{c_i}\}$, for example with weighted least-squares smoothers as in equation (10) or with spline smoothers (*see Splines and other Metamodels in Reliability Analysis*). Then a recommended transformation estimate to obtain the (continuous) hazard or failure rate is [24]

$$\psi(\hat{q}_c(t)) = \frac{1}{\Delta} \log \frac{2 + \Delta \hat{q}_c(t)}{2 - \Delta \hat{q}_c(t)} \quad (13)$$

It can be shown that this smoothing and transformation approach reduces the unavoidable discretization bias.

To avoid discretization biases it is preferable to record continuous failure times if either the density or the failure rate of the lifetime distribution is of interest. Extensions of the methods discussed here are available for the case where one is interested in determining which features of products influence lifetime. This leads to semiparametric hazard regression models, such as the Cox model or the accelerated failure time model (see [3]). Such models combine a baseline hazard function as a nonparametric component with a parametric part that models the influence of the covariates.

In reliability applications shape-restricted failure rates sometimes play a role. A common assumption is increasing failure rate (IFR). Kernel estimators can be monotonized in various ways, from applying the simple pool adjacent violators algorithm (PAVA) to more

sophisticated schemes [25]. Other shape restrictions such as bathtub shape can also be implemented [21].

References

- [1] Rosenblatt, M. (1956). Remarks on some nonparametric estimates of a density function, *Annals of Mathematical Statistics* **27**, 832–837.
- [2] Liebscher, E. (2005). A semiparametric density estimator based on elliptical distributions, *Journal of Multivariate Analysis* **92**, 205–225.
- [3] Cox, D.R. & Oaks, D.R. (1990). *Analysis of Survival Analysis*, Chapman & Hall, London.
- [4] Fix, E. & Hodges, J.L. (1951). Discriminatory analysis and non-parametric estimation: consistency properties, Report No. 4, Project Number 21-49-004, USAF School of Aviation Medicine, Randolph Field, Texas.
- [5] Parzen, E. (1962). On estimation of probability density function and mode, *Annals of Mathematical Statistics* **33**, 1065–1076.
- [6] Silverman, B.W. (1986). *Density Estimation for Statistics and Data Analysis*, Chapman & Hall, London.
- [7] Scott, D.W. (1992). *Multivariate Density Estimation*, John Wiley & Sons, New York.
- [8] Wand, M.P. & Jones, M.C. (1995). *Kernel Smoothing*, Chapman & Hall, London.
- [9] Sheather, S.J. (2004). Density estimation, *Statistical Science* **19**, 588–597.
- [10] Jones, M.C., Marron, J. & Sheather, S. (1996). A brief survey of bandwidth selection for density estimation, *Journal of the American Statistical Association* **91**, 401–407.
- [11] Granovsky, B. & Müller, H.G. (1991). Optimizing kernel methods: a unifying variational principle, *International Statistical Review* **59**, 373–388.
- [12] Müller, H.G. & Wang, J.L. (1994). Hazard rate estimation under random censoring with varying kernels and bandwidths, *Biometrics* **50**, 61–76.
- [13] Watelet, L. & Winter, B.B. (1991). Nonparametric estimation of nonincreasing densities and use of data from renewal processes, *Communications in Statistics: Theory and Methods* **20**, 2073–2094.
- [14] Winter, B.B. (1989). Joint simulation and forward recurrence times in a renewal process, *Journal of Applied Probability* **26**, 404–407.
- [15] Müller, H.G., Wang, J.L., Carey, J.R., Caswell-Chen, E.P., Chen, C., Papadopoulos, N. & Yao, F. (2004). Demographic window to aging in the wild: constructing life tables and estimating survival functions from marked individuals of unknown age, *Aging Cell* **3**, 125–131.
- [16] Fan, J. & Gijbels, I. (1996). *Local Polynomial Modeling and Its Applications*, Chapman & Hall, London.
- [17] Gajek, L. (1986). On improving density estimators which are not bona fide functions, *Annals of Statistics* **14**, 1612–1618.
- [18] Nelson, W. (2000). Theory and applications of hazard plotting for censored failure data, *Technometrics* **42**, 12–25.
- [19] Cheng, M.Y., Hall, P. & Tu, D. (2006). Confidence bands for hazard rates under random censorship, *Biometrika* **93**, 357–366.
- [20] Rammlau-Hansen, H. (1983). Smoothing counting process intensities by means of kernel functions, *Annals of Statistics* **11**, 453–466.
- [21] Reboul, L. (2005). Estimation of a function under shape restrictions. Applications to reliability, *Annals of Statistics* **33**, 1330–1356.
- [22] O’Sullivan, F. (1988). Fast computation of fully automated log-density and log-hazard estimators, *Siam Journal on Scientific and Statistical Computing* **9**, 363–379.
- [23] Wang, J.L., Müller, H.G. & Capra, W.B. (1998). Analysis of oldest-old mortality: lifetables revisited, *Annals of Statistics* **26**, 126–163.
- [24] Müller, H.G., Wang, J.L. & Capra, W.B. (1997). From lifetables to hazard rates: the transformation approach, *Biometrika* **84**, 881–892.
- [25] Hall, P., Huang, L.-S., Gifford, J.A. & Gijbels, I. (2001). Nonparametric estimation of hazard rate under the constraint of monotonicity, *Journal of Computational and Graphical Statistics* **10**, 592–614.

Related Articles

Cumulative Distribution Function (CDF); Intensity Functions for Nonhomogeneous Poisson Processes; Smoothed Function Estimation for Censored Data; Splines and other Metamodels in Reliability Analysis.

HANS-GEORG MÜLLER AND JANE-LING WANG

Monte Carlo Methods, Univariate and Multivariate

Monte Carlo Method

In the world of simulation, a **Monte Carlo** method is a means of solving a deterministic problem using random numbers. The origin of the term comes from a casino in Monaco; the method requires many repetitions of a random process, which is similar to that which occurs in a casino. The term became popular with those using such methods in the 1940s and 1950s, notably Nicholas Metropolis and Stanislaw Ulam [1], among others.

The ideas underpinning the method had been recognized well in advance of this time. For example, the case of Buffon's needle is a classic example of the use of random experiments in order to estimate a deterministic quantity. Since the time of Metropolis and Ulam, the methods have been of direct interest to the statistical community; [2] indeed even in the 1980s it was recognized that such simulation methods were useful, in an age when computers were "powerful" and mathematicians were scarce [3]. Of course, computer power has continued to improve substantially over the past two decades; and such methods are increasingly common [4, 5].

In order to describe the method in more detail, we look at a simple example. Following this, some technical considerations will be highlighted, together with a discussion of how the method can be applied to a more complex situation. We then note how the method relates to that of ensembles. Finally we examine some particular methods.

A Simple Example

In order to demonstrate the idea, we solve a very simple problem using such a method.

Consider the question "What is the area of a circle inscribed within the unit square?" Graphically, the region of interest is highlighted in Figure 1. The total area from which the simulation will be carried out is the white square, with the circle shaded in black.

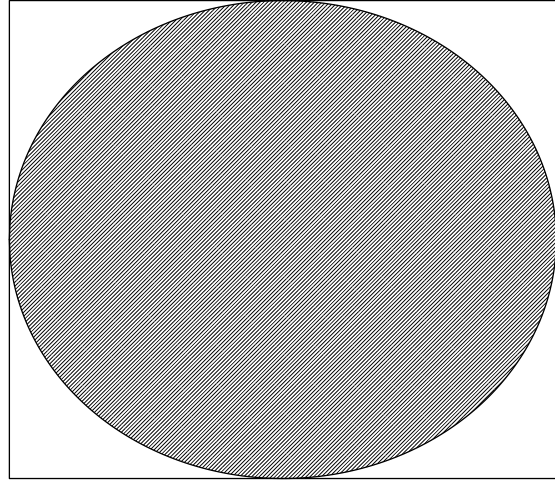


Figure 1 Pictorial representation of circle in unit square

In order to answer this, we consider an "experiment" in which there is an element of chance. The idea is that we "throw" darts uniformly at the unit square. For each of the darts, we then determine whether or not it has fallen within the shaded circle. If it has, then we count it as a success. Otherwise we think of it as a failure. The relative area of the circle to the square is then approximated by the proportion of successes to the total number of throws.

Thus, the formal algorithm, noting that the circle is centered at (0.5,0.5) and has radius 0.5 proceeds as follows;

- Set $s = 0$.
- Repeat N times
 - Generate $x \sim U[0, 1]$.
 - Generate $y \sim U[0, 1]$.
 - If $(x - 0.5)^2 + (y - 0.5)^2 < (0.5)^2$ then $s = s + 1$.

The area is then estimated as the fraction, s divided by the total number of simulations, N .

Characteristics of the Algorithm

We note that the method has the following characteristics;

2 Monte Carlo Methods, Univariate and Multivariate

- The method requires a source of random numbers (the $U[0,1]$).
- There are many repetitions of a simple process (the throwing of the darts).
- The answer that is obtained is an approximation to that which is required.

Indeed, the first of these conditions is not strictly true; it is sufficient that the source of “random” numbers be pseudorandom – the key requirement is that there is no relationship between what is being estimated and the source we are using.

We (the persons carrying out the experiment) are in charge of the number of such repetitions that we obtain. In practice, since computers are typically used to carry out this simulation process, it is “cheap” to carry out each repetition. Thus the main barrier to the size of the collection of realizations is the time available and the processor power and memory storage available to us.

Even though the answer we obtain is only an approximation to that which is required, we can quantify the size of the error. This is termed *Monte Carlo error*, and by varying the number of simulations we can make this arbitrarily small. Since this last point is important, a little more detail is provided below.

Monte Carlo Error

Where independent simulations have been obtained it is possible to quantify the size of the Monte Carlo error by the following considerations;

Consider the example above. Formally we describe the probability of the “dart” being in the unit square by a parameter, p . The throwing of each “dart” is then a Bernoulli trial; that is the result of the throw is 1 (i.e., success – landing inside the circle) with probability p and 0 (i.e., failure – landing outside the circle) with probability $1 - p$. We denote this random variable as X .

If we consider a large collection (n) of independent realizations of such random variables, we end up with a collection $\{x_1, x_2, \dots, x_n\}$ of zeroes or ones. The sum of these n realizations may be denoted S_n . It is clear that S_n is then a random variable which has a binomial distribution with parameters (n, p) . In this case, the expectation and variance of the random variable may readily be obtained as np and $np(1 - p)$ respectively.

Thus, if we consider the derived random variable $p_{\text{est}}(n)$, given by;

$$p_{\text{est}}(n) = \frac{S_n}{n} \quad (1)$$

then this has expected value p and variance $\frac{p(1-p)}{n}$. Thus, as n gets large, we are increasingly confident that $p_{\text{est}}(n)$ is close to p .

Returning to the question, we wish to estimate the relative area of the circle. By considering the setup, we note that this relative area is exactly the probability of a uniformly thrown “dart” hitting the circle. Thus, it is equal to p , and our task is to estimate p to a given degree of accuracy.

Using the central limit theorem, we note that as n increases there is approximately 95% probability of $p_{\text{est}}(n)$ being contained in a ball that stretches 2 standard deviations from p , and 99.9% probability of being in a ball that extends about 3 standard deviations from p . Thus, since we can construct $p_{\text{est}}(n)$ by “experiment” we have a method of estimating p .

By observing this fact, and noticing that we can increase n to be arbitrarily large, we note that we can obtain estimates of p that are accurate to a given accuracy.

The residual uncertainty that exists because $p_{\text{est}}(n)$ only *approximates* p may be quantified in terms of the standard deviation of the estimator considered as a random variable. This, as outlined above is $\sqrt{\frac{p(1-p)}{n}}$.

It is this residual uncertainty that is termed the *Monte Carlo error*. In any situation where a Monte Carlo method is used, it is essential that some reference be made to this quantity – either explicitly or even just to note that the relative size of the Monte Carlo error is “small”.

In general, the Monte Carlo error may be estimated using the standard deviation of the generated sequence; in the example above, this would be the standard deviation of the sequence of realizations of the Bernoulli random variables.

An Extension of the Example

The example that has been used is a particularly simple one. Indeed, it has the status of a “toy” and is used simply to illustrate the method.

However, we may readily consider situations that are of real practical interest to us.

Consider a system, S , which takes a number of inputs; we may think of the system as being described by $S(I_1, I_2, \dots, I_m; 0)$ at time 0. We may be interested in an indicator function, R , which takes the value 1 if the system is functioning and the value 0 if it has failed.

Thus, $R(S(I_1, I_2, \dots, I_m; 0)) = 1$ tells us that the system with the given inputs was functioning initially; and if $R(S(I_1, I_2, \dots, I_m; t_1)) = 1$ and the quantity at a later time is given by $R(S(I_1, I_2, \dots, I_m; t_2)) = 0$ we know that the system was functioning at t_1 but had failed by t_2 .

Now, in general, it is possible that the inputs are uncertain – that is that $I_1, I_2, \dots, I_m$ may be thought of as realizations of a multivariate random variable. It is likely that it is possible to quantify our uncertainty about the inputs in the form of a probability distribution – for example if I_1 is temperature, it may be possible to say it can be thought of as normally distributed centered at room temperature.

Regardless of the fact that the inputs are random variables, we are still concerned with some derived quantity – for example the probability that the system is operational at 100h, say. This may be obtained by integrating across the probability distribution for the multivariate inputs, and thus calculating the expectation of $R(S(I_1, I_2, \dots, I_m; 100))$. That is;

$$\begin{aligned} P(R(S(\mathbf{I}; 100))) &= 1 \\ &= E(R(S(\mathbf{I}; 100))) \\ &= \int R(S(\mathbf{I}; 100))f(\mathbf{I}) d\mathbf{I} \quad (2) \end{aligned}$$

In principle, such a calculation is possible. In practice, it is analytically intractable. Where the input space is of high dimension, it is difficult to do this using quadrature or other standard techniques [6].

Thus, a procedure exactly analogous to that used with the random dart above is employed. In this case it runs as follows;

- Repeat N times
 - Generate $\mathbf{I} \sim f(\mathbf{u})$.
 - Evolve the system to 100h.
 - If $R(S(I_1, I_2, \dots, I_m; 100)) = 1$ then $s = s + 1$

The reliability of the system at 100h may then be estimated as $\frac{s}{N}$.

Relationship with Method of Ensembles

A closely related idea is the *method of ensembles*. The notion here is to consider a large number (N) of realizations of the system, where for each of the realizations, a new value of the input parameters is simulated. The resulting collection is then termed an *ensemble* and may be thought of as providing the analyst with realizations of the system in N parallel universes.

This method is commonly used in statistical mechanics, and increasingly so in applications where financial risk is the key consideration.

The analyst then examines these multiple realizations to summarize a key quantity of interest. For example, if we are interested in whether the system has failed by 100h, we examine the ensemble, and the “best guess” at this probability is the fraction of the ensemble for which the system has failed. The method is exactly the same as the Monte Carlo strategy described above, but the way in which we express it is slightly different.

The advantage of expressing our findings in terms of an ensemble is that it is often easier for those not used to the language of probability, marginalization and expectation to think in this fashion. However, the technical considerations are exactly the same; we must concern ourselves with how many members are in the ensemble, and the associated Monte Carlo error.

Particular Examples of Direct Monte Carlo Methods

A number of specific techniques are worthy of exploration. These are briefly outlined in what follows.

Univariate Rejection Method

One method of obtaining deviates from a distribution of known form is the rejection algorithm (sometimes termed the *accept–reject method*). This is a very flexible method, which permits a wide class of probability distributions to be used in Monte Carlo simulation strategies.

4 Monte Carlo Methods, Univariate and Multivariate

The rejection method is a two stage process. Let the probability distribution from which a deviate is required be denoted $f(\cdot)$.

The first step is that a proposed deviate is generated from a proposal probability distribution. This proposal distribution, let us call it $g(\cdot)$ is chosen such that $Mg(x) > f(x)$ for some scalar M and for all x in the domain of interest. Typically the form of $g(\cdot)$ is simple, in that deviates are cheaply available.

The second step involves preferentially accepting proposed deviates from areas of $f(\cdot)$ that have high probability. Explicitly, this is done by accepting a proposed deviate x_{prop} with probability given by;

$$p_{\text{accept}} = \frac{f(x_{\text{prop}})}{Mg(x_{\text{prop}})} \quad (3)$$

If the proposed deviate is not accepted, then a new proposal is generated and the process is repeated.

Thus the algorithm is as follows;

- Generate $x_{\text{prop}} \sim g(x)$.
- Calculate $p_{\text{accept}} = f(x_{\text{prop}})/Mg(x_{\text{prop}})$.
- Generate $u \sim U[0, 1)$.
- If $u < p_{\text{accept}}$ then accept x_{prop} ;
- Else repeat the process.

The choice of $g(\cdot)$ is of critical importance in determining the efficiency of the algorithm. If $Mg(\cdot)$ does not tightly bound $f(\cdot)$, then a large number of deviates will be generated for each one that is accepted. Depending on the cost of sampling from $g(\cdot)$ the algorithm may then be impractical. It can easily be shown that the expected number of times the sampler will have to draw from $g(\cdot)$ for each successfully accepted deviate is given by the scalar, M .

A pictorial representation of the setup is shown in Figure 2.

Multivariate Rejection Method

If the density to be simulated from, $f(\mathbf{x}) = f(x_1, x_2, \dots, x_m)$ which is a multivariate density, then the multivariate version of the algorithm can be used. In order to do this, a multivariate proposal density, $g(\mathbf{x})$ is required, with the same bounding conditions.

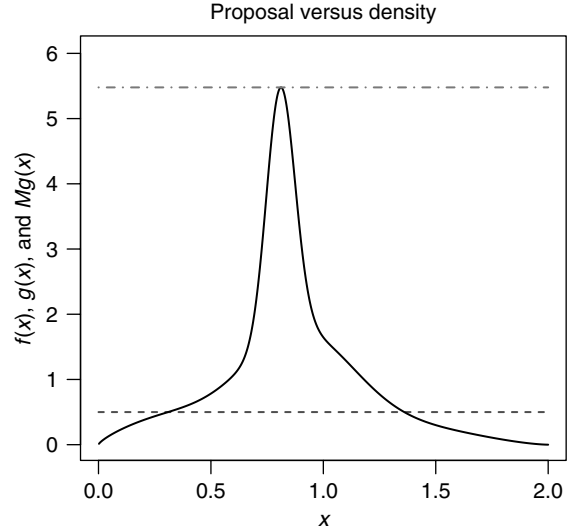


Figure 2 Typical situation where rejection is used

In this instance, the scalar, M , typically has to be much bigger, and thus the algorithm is correspondingly less efficient. Thus, in practice, fully multivariate versions of the rejection algorithm are not commonly used.

Importance Sampling

If the cost of generating deviates from the probability distribution of interest is substantial, then it makes sense that we use these as efficiently as possible. One such method, which ensures that we examine the event frequently of interest to us (and consequently reduces the variance of our estimator), is called *importance sampling*.

The idea behind importance sampling is to “adjust” the proposal from which deviates are selected, in order that the Monte Carlo algorithm evaluates the function in the region of interest more frequently. The fact that the proposal density has been adjusted has to subsequently be corrected for; this is done by reweighting the sampled values. An example makes this idea clear.

Consider the use of importance sampling in numerically approximating an integral. The basic principle is as follows:

In wanting to estimate

$$I = \int_{-\infty}^{\infty} g(x)f(x) dx \quad (4)$$

where $f(x)$ is a density function, one could sample n values of x from $f(x)$ and then approximate with

$$\hat{I} = \frac{1}{n} \sum_{i=1}^n g(x_i) \quad (5)$$

Alternatively, m values of x could be sampled from another density $h(x)$ and then I could be estimated using

$$\hat{I} = \frac{1}{m} \sum_{i=1}^m \frac{g(x_i)f(x_i)}{h(x_i)} \quad (6)$$

We have to be careful about how $h(x)$ is chosen so that the estimator is most efficient. It turns out that the most efficient form for $h(x)$, samples from areas where $g(x)$ is large, provided that $f(x)$ is not small [2].

These results extend easily to the multivariate setting. However, effective choice of $h(x)$ is more difficult in the multidimensional setting. Thus, if we have good intuition about what areas of $g(x)$ are of interest, we can exploit them using this technique. The risk, however, is that if we choose poorly we will obtain a much less efficient **estimation** strategy.

Importance sampling can be particularly useful in reliability modeling where we are concerned with rare or extreme events; for example, we may be interested in extreme loadings to a structure, which have high probability of causing damage. Thus, in order to estimate the overall reliability of the structure, we may wish to ensure we simulate many of these “high-impact events”. We then adjust the overall weighting of these according to their true probability, in order to reduce the variability of the estimated survival function in the region where it changes most.

Composition

When a density can be described as a combination of other densities, the method of composition can be used. The simplest example of this is the context in which the density is a mixture distribution.

For example, consider a situation in which we are interested in estimating the probability that a system copes with a random loading event. There may be two possible loading modes, M_1 and M_2 ; and the damage caused may be described by a random variable for each of the possible modes. Thus, for example, if M_1 occurs, the damage may be $D \sim \text{Gamma}(\alpha_1, \beta_1)$ and

if M_2 then $D \sim \text{Gamma}(\alpha_2, \beta_2)$. If we know that the probability of M_1 is p_1 then we can decompose the **pdf** for D into two parts. Specifically;

$$f(D) = p_1 f_1 + p_2 f_2 \quad (7)$$

where f_1 is $\text{Gamma}(\alpha_1, \beta_1)$ and f_2 is $\text{Gamma}(\alpha_2, \beta_2)$.

In order to simulate from such a decomposed distribution, one proceeds using the “broken stick” method as follows;

- Generate $u \sim U[0, 1)$.
- If $u < p_1$ then simulate $D \sim f_1(\cdot)$.
- Else $D \sim f_2(\cdot)$.

Thus, we simulate D from the first component with probability p_1 and from the second with probability p_2 . This can be applied to a mixture with any number of components.

It is noted that some distributions can be very well approximated by mixtures; for example, early implementations yielding deviates from a Gaussian used this method, as outlined in [3].

Box–Muller

An example of a transformation method which efficiently generates from the bivariate **normal distribution** is the Box–Muller transform. The method is a simple transformation method, which takes deviates from the uniform unit square and maps them to the bivariate standard normal. It proceeds as follows;

- Generate $u_1 \sim U[0, 1)$.
- Generate $u_2 \sim U[0, 1)$.
- Let $\theta = 2\pi u_1$.
- Let $R = \sqrt{(-2 \ln(u_2))}$.
- Return (x, y) by transformation of polar (R, θ) .

Then the pair (x, y) are independent standard normal deviates.

The advantage of this method is that it efficiently generates from the required distribution. None of the deviates are “wasted” through rejection. It is easy to implement. Such transformation methods exist for other situations.

Correlated Random Variables

Special concern is needed when the probability distributions of interest to us are correlated. This is

critical when considering the tails of distributions. Such issues arise in the context of reliability applications, and applications in **finance**, such as the probability of ruin.

As an example, consider the simple sum of standard normally distributed random variables, $Y = X_1 + X_2$. If X_1 and X_2 are independent of each other, then a Monte Carlo strategy involving Y could proceed with X_1 and X_2 each being simulated without consideration of the other. However, if there is a positive **correlation** between the random variables – that is large values of X_1 are seen more often in combination with large values of X_2 – then extreme values of Y are more likely than we would observe if we simulated them independently. Similarly, if there is a negative correlation between X_1 and X_2 , we will see extreme values of the combination less frequently than we might naively expect.

In order to ensure that we correctly account for correlation between the random variables, we could simulate X_1 from the marginal distribution and X_2 from the conditional distribution.

Other Extensions and Considerations

What has been described in this section is the Monte Carlo method in general. It should be clear that much of the work involved in implementing a Monte Carlo strategy involves careful consideration of how to simulate from the density of interest to us. This has not been discussed here as it is the subject of other articles.

In particular, the fact that we use pseudorandom numbers, that is a sequence of deviates generated by a computer which “looks” random, is a fact that has been glossed over. The detail of how such a sequence may be generated, some of the pitfalls, and some of the features of the resulting sequence one may wish to check are all the subject of other articles.

If one wishes to sample from a distribution that is not uniform, then one needs a method of doing this. For example, in considering **system reliability** above, one of the parameters was sampled from a normal distribution. For completeness, it is necessary for us to demonstrate that it is in general possible to do this. It is sufficient to note that, given the cumulative distribution function, this is always possible using the inverse transform sampling method. However, more

efficient techniques often exist, as highlighted for particular situations; for example methods for particular parametric distributions are discussed elsewhere [7–10].

Other articles [11, 12], discuss derived techniques such as **Markov chain Monte Carlo**, where the samples that are generated are no longer independent of one another. The extension has substantial advantages in high dimensional problems. However, the consideration of Monte Carlo error is not as straightforward as outlined above. Additionally, care is needed to ensure that the sampler is well behaved.

References

- [1] Metropolis, N. & Ulam, S. (1949). The Monte Carlo method, *Journal of the American Statistical Association* **44**(247), 335–341.
- [2] Kleijnen, J.P.C. (1974). *Statistical Techniques in Simulation*, M. Dekker, New York.
- [3] Ripley, B.D. (1987). *Stochastic Simulation*, John Wiley & Sons, New York.
- [4] Fishman, G.S. (1995). *Monte Carlo: Concepts, Algorithms, and Applications*, Springer-Verlag, New York.
- [5] Robert, C.P. & Casella, G. (2004). *Monte Carlo Statistical Methods*, 2nd Edition, Springer-Verlag, New York.
- [6] Concepcion Ausin, M. (2007). An introduction to quadrature and other numerical integration techniques, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons.
- [7] Rosinski, J. (2007). Simulation of Lévy processes, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons.
- [8] Marseguerra, M. & Zio, E. (2007). Simulation of life distributions, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons.
- [9] De Blasi, P. (2007). Simulation of the beta-Stacy process with application to analysis of censored data, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons.
- [10] Laud, P.W. & Pajewski N.M. (2007). Simulation of the dirichlet process, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons.
- [11] Wiper, M.P. (2007). Introduction to Markov chain Monte Carlo simulation, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons.
- [12] Mira A. (2007). Strategies to improve convergence and mixing in Markov chain Monte Carlo, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons.

CATHAL D. WALSH

Discrete-Event Simulation for Reliability Prediction

Introduction

System reliability is achieved by a combination of **component reliability**, system architecture, and maintenance. In some cases, where the structure of a system is simple and component models are well-defined then analytical solution of a systems reliability is often possible, although this often requires a large number of assumptions about specific probability models, and hence rather limits the applicability of the results. As soon as we move away from simple models into real-world situations then the mathematics gets far too difficult to solve analytically and so we need to move to an empirical approach. Discrete-event simulation (DES) is one such empirical approach. DES allows us to build real-world models and attempts to use them to predict what is likely to happen to the reliability of a system in the future. It should be kept in mind that DES produces an output based on probabilities of events occurring and so cannot be absolutely accurate but with enough experimental runs the uncertainty can be reduced sufficiently for the results to be useful. For this reason DES is widely used in reliability work and forms part of the internal workings in many pieces of reliability software.

Simulation

Simulation is primarily about the evaluation of a system design to establish its performance under various operating conditions before we commit capital or other resources. To be able to predict performance requires that we develop some kind of model of the proposed system that is independent of the real system but will anticipate its physical behavior with an acceptable degree of accuracy. All models are, by definition, simplifications and will not reproduce the behavior of the real system exactly. Simulation is about modeling, but this alone is not sufficient; it is also about conducting experiments using the model. We conduct experiments using a model to predict performance but the results are never exact. The way in which the experiments are set up, the

initial conditions, duration of “run”, and the way in which statistics are recorded can have a significant affect on the quality of results. Lastly, we have to be careful about how the results are presented and how they are interpreted. The results of a simulation are, typically, only one piece of evidence used to support the decision-making process. Sometimes the results are neither what are expected nor desired and in such cases, unfortunately, may be ignored, perhaps with “disastrous” consequences. The basic simulation process is to use a model to calculate the system state at some time t and then to move time by a small increment and recalculate the system state. Many more details about simulation in general can be found in [1].

Applying Simulation to Reliability Prediction

The root of computer simulation modeling of system reliability is **Monte Carlo simulation** [2]. Monte Carlo simulation (*see Monte Carlo Methods, Univariate and Multivariate; System Reliability: Monte Carlo Estimation*) randomly creates individual component failures, using information about each component’s failure influences, and monitors the overall system state according to predetermined rules about which components need to be operating for the system to be functional. The individual component failure distributions are easily modeled for most theoretical failure distributions and where a component failure distribution cannot be modeled theoretically, or there is insufficient information, failures and other events can be modeled by histograms or by empirical distributions [3].

An important feature of Monte Carlo simulation is that it “knows” what state the system is in at any point in time during a particular run. Accordingly, it is possible to introduce “business” and “managerial” decisions as to how to proceed in the event of a particular condition. For example, “If insufficient manpower to carry out repair, and go-no-go conditions not violated allow continued operation under restricted conditions.” This is a powerful tool for improving the veracity of the analysis, because it significantly improves the “real-life” factors that affect equipment operation.

Each Monte Carlo run will produce (internally) an answer of when the system was serviceable, and when

2 Discrete-Event Simulation for Reliability Prediction

it was in a failed state, and indeed any other pertinent conditions. Many subsequent runs will give different answers, because failure events are being created randomly. However, the many runs will show when analyzed that, say, $x\%$ of runs showed the system to be serviceable at time t . Therefore, the point estimate of reliability is $x\%$ at time t . Other parameters, such as the number of repairs being undertaken concurrently, or the spare part consumption, etc., may also be monitored and, across the many runs, estimates of the average values at any all points can be derived.

However, the answer itself is subject to variability. If the many Monte Carlo runs were repeated, a slightly different result would be obtained. To place confidence limits on the answer, the number of runs within each “simulation” must be considered. The greater the number of individual runs, the less would be the variability between several complete simulations. The formula for calculation of the variability is strictly a binomial confidence limit, but because of the extremely large number of individual runs usually undertaken, the normal approximation is appropriate [4].

Classical Monte Carlo Method

The classical method of Monte Carlo simulation is to examine the serviceability of each component at very fine time-slice intervals. If a component was serviceable at time t , the probability of it failing in time δt is small but known, and therefore, following random number generation, its serviceability at time $t + \delta t$ is deduced. Quite clearly when using very fine time slices, the serviceability of each component is calculated a very large number of times, each according to the failure distribution of the individual object. For large systems, this very quickly leads to a very slow simulation, expensive in computer time and so this approach is not recommended.

Discrete-Event Monte Carlo

A more modern approach to Monte Carlo simulation is to determine the complete time to failure of a component, rather than the failure event within small time increments, and to then to save (remember) the time when the event will occur. All simulation events which occur before that time will, by definition, show the component serviceable. The same philosophy

applies to all other events within the simulation. The operation of the system is then modeled, not by uniform time slices, but by jumping from one discrete-event to another. Provided all events have been initially calculated, or have a mechanism for being updated by the simulation itself (the usual case), then there is no change in system state between events. Accordingly, no fidelity is lost by jumping between events. Indeed, because time can now be a continuous variable rather than a finite slice, the DES offers greater fidelity.

A discrete-event Monte Carlo simulation usually must include the logic to update the list of future events as it proceeds. This will be illustrated by example. Consider a two-unit warm standby system (without repair), with $f_a(t)$ for the active component and $f_s(t)$ when in standby. The standby component adopts $f_a(t)$ when it becomes active. Initial random selection: sample $f_a(t)$ and $f_s(t)$, derive t_a and t_b being the future times when the active and standby components will fail within the particular simulation run.

1. If $t_a > t_s$, event 1 is at t_s (dormant failure of standby component), and event 2 is at t_a (failure of active component, system failure).
2. If $t_a < t_s$, event 1 is at t_a (failure of active component, standby component becomes active). This leads to a new event calculation: sample $f_a(t)$ for the new active component and derive new t_{a2} , leading to event 2 at $t_a + t_{a2}$ (second component failure, system failure).
3. If $t_a = t_s$, event 1 is at t_a (failure of both components, system failure).

Thus, in the case of 2 above, a second event calculation is made, which continues the simulation to the next event. Indeed, for this system, there are never more than three individual event calculations before system failure. Quite clearly this is far more efficient than taking fine time slices and sampling each.

More complex simulations are equally easily built-up. The logic of determining subsequent events is straightforward following each previous event. For example, in a repairable system, subsequent events following component failure would be to “initiate repair”. At the time of the failure event, the availability of “technicians” and “spare parts” could be examined within the simulation (i.e., their individual state, determined by their own events),

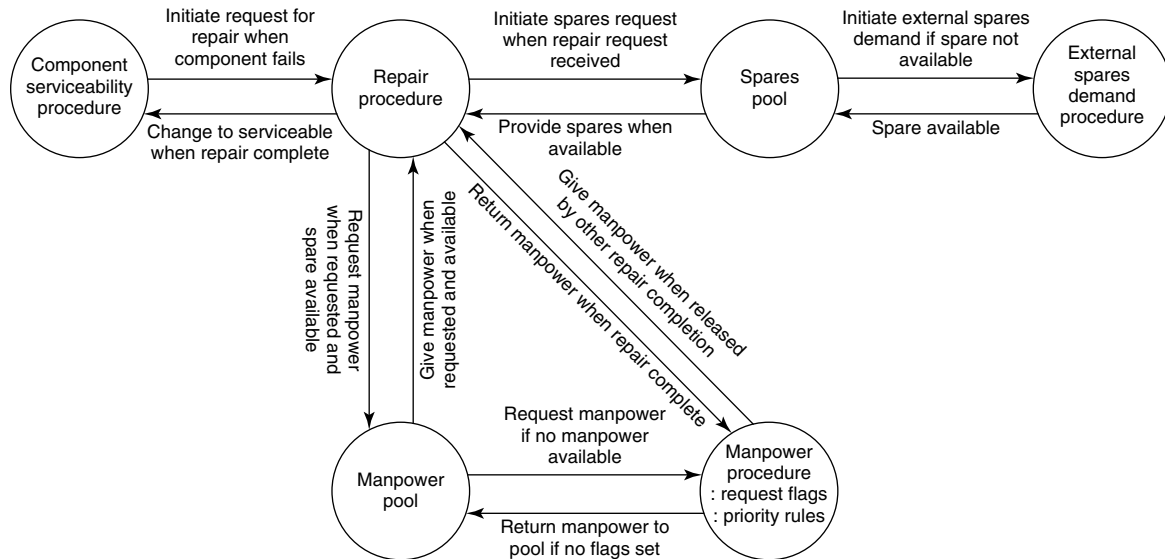


Figure 1 Discrete-event Monte Carlo simulation flow diagram

and rules governing the start and duration of repair sampled (probably also randomly sampled from a distribution, but equally easily obtained from discrete predetermined values). The simulation need only sample as far as the next logical event. New logical events may also override previously calculated event times, but such changes are easily included in the simulation model such that the list of future events is always the latest and correct list.

Thus, if the two-component warm standby system above was repairable, then a failure event would, by logic, lead to a “request” for repair. In practice, a logical flow diagram is often useful in understanding and declaring the sequence of events and decisions affecting a system’s operation. Figure 1 is a simple example of such a flow diagram. At the time of failure, the availability of technicians and spares is given by the simulation state. In the flow diagram in Figure 1, it is procedural that a spares request is initiated before a manpower request. The manpower request is only made when a spare is available. Manpower is allocated directly from the pool, if available. Otherwise the request is redirected to the manpower allocation procedure. The manpower allocation procedure is (for example) able to determine priority between competing repair tasks and maintains a list of flags, which identify waiting repair tasks. Manpower is always returned through the allocation procedure,

which would in turn return the manpower to the pool if there are no other tasks awaiting manpower. There are many variations on the flow diagram above. This particular diagram is given as way of illustration of the concept of discrete-event Monte Carlo simulation. Each procedure generates another event in response to a request. Thus, the simulation is self generating until an “end” condition is met.

Data

Supporting data is always required to create events. This applies equally to probabilistic and deterministic parameters. The more detailed the data available, the more detailed can be the computer simulation. The corollary is that with less data, it is false fidelity to apply greater detail to the model structure. The data must match the structure and *vice versa*.

In the example above, spares requirements are modeled probabilistically from the failure event. However, the level of detail with which the spare is identified should be commensurate with the level of failure probability and design information available. Where only system level information is available to any degree of accuracy, a system level “spare” should be generated by the simulation model (e.g., “engine”). Only where more detailed failure information is available, and the objectives of the simulation require

it, should more detailed breakdown be used (e.g., “inlet valve”).

Modeling Methods

Discrete-event probabilistic modeling is very closely allied to the development of a distributed database system. In a distributed database system, user inputs will generate events. The events may be arrival of data, analysis of data, output of data summaries, etc. Logical methodologies and development notation now exists to manage the specification and building of software systems in general, such as database systems. Use of such methodologies, particularly object oriented methodologies such as Booch notation [5], is highly recommended in the development of simulation models as without them, it is very easy to lose control of the event generation and flow.

Modeling Tools

A simulation model can be built using low level computer languages, such as C++. However, by far the most effective tools are dedicated, high-level tools such as Simula [6], Simscript [7], and Witness [8]. Clearly the performance of the PC will limit the size of simulation which can be run with any degree of efficiency.

DES for reliability purposes does not need to model component reliability, system architecture and maintenance using a single methodology. Rather, DES may adopt the most appropriate methodology for each. Therefore component status and maintenance are separately modeled and are inputs to a system state-space model. There are several analytical methods to derive state space. This article will briefly consider Petri nets [9], Markov diagrams [10], **fault trees** [11], reliability block diagram (RBD) [12], event trees [13], and path sets [14] as ways of representing system structure within DES.

Petri Nets. Petri nets have states (often called *places*) and transitions, and model the dynamic system state by addition and deletion of “tokens”. Petri nets provide opportunity to model every potential system state, and therefore “success”, “partial success”, “failed”, and indeed any other definable system condition of interest can be specifically identified. Therefore, it is a rich analysis methodology. A Petri

net is usually represented by a “transition matrix” and by a “token vector” prior to further computer-aided analysis, including for DES. The matrix describes the addition and deletion of “tokens” when triggered by certain input conditions, such as component failure. The total presence of “tokens”, as denoted by the “token vector”, denotes the system state. A system may remain in a given state for a period of time (say, using a simple stochastic component failure model) before addition/deletion of tokens, or may wait for external input (discrete-event) to trigger such change. The matrix cell values for stochastic models represent the probability of transfer but when a Petri net is linked to external trigger, the cells will be binary. For DES purposes, the most appropriate method is external input and, therefore, the simpler binary transfer matrix is more appropriate. When an external event triggers the Petri “token” vector, a change is calculated by repeated additive matrix/vector multiplication until a steady state is obtained in the “token” vector.

Markov Analysis. Markov analysis (*see Markov Processes; Maintenance and Markov Decision Models; Markov Renewal Processes in Reliability Modeling*) is an equally rich analysis method that also identifies all system states individually. As with Petri nets, the system states are usually represented by a matrix. A standard Markov transition matrix is also the form $k^m \times k^m$, where m is the number of components and k the number of possible component states. Each component would have an associated transition vector of length k^m . Therefore, computer representation and calculation of system state derived from Markov analysis would suffer from the same “state explosion” as from Petri nets, however, the matrix and vector size may be minimized if all system failed states are combined, as with Petri nets.

However, White developed an alternate transition matrix representation [15]. Instead of using a binary vector (matrix row) to identify the state transfer, he enumerates the states and forms a matrix of “component fault mode” and “system state”, thus immediately reducing the state matrix to $k(m-1) \times k^m$. White’s method has the advantage that no matrix multiplication is required to identify the new state. State change is identified through indexing, which is time efficient.

Fault Tree, Event Tree, and RBD. Fault trees (see **Failure Modes and Effects Analysis, Implementation of; Fault Tree Analysis for Large Systems**), event trees, and RBD may visually represent dynamic conditions, such as sequence and functional dependency, but such conditions cannot be analyzed without first being converted to, say, a Markov representation. Therefore, use of fault and event tree or RBD is often restricted to static conditions. Results from fault tree, event tree, or RBD are commonly represented as cut or path sets. These are binary vectors of (for paths) component states that, collectively, indicate system function. A fault tree captures system configurations that would lead to a specific “top” event, usually system failure. An event tree, because it potentially identifies each and every system outcome, will contain path sets or cut sets for each outcome category.

When applied to DES, the system state upon failure or repair of a component is simply an issue of updating the binary state vector, and identifying which category the new vector resides, cut or path set. This is a “find” function, usually of moderate computing efficiency $O(\ln(n))$.

Path Sets. Path and cut sets (see **Path Sets and Cut Sets in System Reliability Modeling**) collectively describe the full combination of possible component states. Therefore, dynamic systems do not increase the overall number of possibilities but, rather, conditionally change a path into a cut. If a system state is statically failed, it is not reasonable for it to be dynamically functional. Therefore, dynamic systems would not conditionally change a cut into a path. Thus, it is implicitly possible to model dynamic systems while still using path and cut sets, provided an additional data-structure is feasible to capture and identify conditions where a statically functional system is dynamically failed.

References

- [1] Law, A. & Kelton, D. (1991). *Simulation Modelling and Analysis*, McGraw-Hill.
- [2] Featherstone, A.M. (1985). Application of Monte-Carlo computer simulation techniques in reliability assessment, in *Computers in Engineering Proceedings of the International Computers in Engineering Conference*, Vol. 3, pp. 35–39.
- [3] Shanker, A. & Kelton, D. (1991). Empirical input distributions: an alternative to standard input distributions in simulation modeling, in *Winter Simulation Conference Proceedings*, Phoenix, pp. 978–985, Dec 8–11 1991.
- [4] Coit, D.W. (1997). System-reliability confidence-intervals for complex-systems with estimated component-reliability, *IEEE Transactions on Reliability* **46**(4), 487–493.
- [5] Daniels, J. (1994). Object design methods and tools, *IEE Colloquium (Digest)* (007), 3/1–3/3.
- [6] <http://www.cee.hw.ac.uk/~rjp/bookhtml/> (last accessed 20/2/2007).
- [7] <http://www.simscrip.com/> (last accessed 20/2/2007).
- [8] <http://www.sakersolutions.com/simcentre/witness.html> (last accessed 20/2/2007).
- [9] Yefremov, A.A. (2005). Failure-repair processes simulation for parallel redundant configurations using stochastic Petri nets, in *Proceedings – 9th Russian-Korean International Symposium on Science and Technology, KORUS-2005* Vol. 1, Novosibirsk, pp. 738–740, Jun 26–Jul 2 2005.
- [10] Nagrath, I.J., Misra, H.C. & Gupta, R.S. (1971). Digital simulation of power system reliability using Markov processes, *Journal of the Institution of Engineers India Electrical Engineering Division* **52**(4 pt EL 2), 65–68.
- [11] Ragheb, M., Gvillo, D. & Makowitz, H. (1986). Symbolic simulation of engineering systems on a supercomputer, *Proceedings of SPIE – The International Society for Optical Engineering* **635**, 368–374.
- [12] Wang, W., Furlong, E.R., Arno Robert, G. & Vassiliou, P. (2003). Ogden, doug. Apply simulation technique via reliability block diagram to gold book standard network, in *Conference Record of Industrial and Commercial Power Systems Technical Conference*, Orlando, pp. 38–49.
- [13] Nikolopoulos, S.D. & MacLeod, R. (1993). Experimental analysis of event set algorithms for discrete event simulation, *Microprocessing and Microprogramming* **36**(2), 71–81.
- [14] Warrington, L. Jones, J.A. (2003). Representing complex systems within discrete event simulation for reliability assessment, in *Proceedings of the Annual Reliability and Maintainability Symposium*, Tampa, p. 487, Jan 27–30 2003.
- [15] White, P. IEE colloquium (digest), in *Colloquium on Fault Tolerant Techniques*, London, No. 074. May 11 1990, p. 4.

JEFF JONES

Failure Modes and Effects Analysis, Implementation of

Introduction

Failure modes and effects analysis (FMEA) is a structured, logical, and systematic analysis of a system. To perform a FMEA an analyst examines each item in a system, considers how that item can fail, and then determines how each failure will affect the operation of the system.

FMEA provides a basis for improving a system's design by recognizing component failure modes identified in historical "lessons learned" and from component and system prototype tests. It is used to assess the safety of system components (*see Reliability, Safety, and Risk Management*), and to identify design modifications and corrective actions needed to mitigate the effects of a failure on the system. It is used in planning system maintenance activities, subsystem design, and as a framework for system failure detection and isolation capabilities. A FMEA can be performed functionally (functional FMEA), on subsystem interfaces (interface FMEA), or on the piece parts of which the system is composed (detailed FMEA). It can be applied to the processes through which a system is manufactured and maintained (process FMEA) as well as on the electrical, mechanical, and software components of the system (product FMEA).

FMEA developed as a formal methodology during the 1950s at Grumman Aircraft Corporation, where it was used to analyze the safety of flight control systems for naval aircraft [1]. During the 1970s and 1980s, various military and professional society standards were written to define the analysis methodology. Mil-Std-1629 (ships), "Procedures for Performing a Failure Mode, Effects and Criticality Analysis" [2] was published in 1974, and, through several revisions, became the basic approach for doing the analysis. The Society of Automotive Engineers Ground Vehicle Recommended Practice, J1739 [3] defines the procedure used by the major automobile companies and their suppliers. IEC 60812, "Analysis techniques for system reliability – Procedure for failure mode

and effects analysis (FMEA)" [4] is the most recent international standard specifying the technique.

The terms *FMEA* – *failure modes and effects analysis* – and *FMECA* – *failure modes, effects, and criticality analysis* – are often used almost interchangeably. Some authors take the view that FMEA is limited to an analysis of the effects of item failure modes and that FMECA also includes assessing the probability of a failure mode's occurrence. Others take the view that FMECA extends the analysis to include a ranking of the failure modes based on a combination of their probability of occurrence and the severity of their effects. Here we use FMEA as the general term and include prioritizations based on the failure mode's severity and probability of occurrence under its scope.

Example: Analysis of a Domestic Hot Water Heater

To illustrate the basic FMEA technique we will begin with an analysis of the gas stop valve for a domestic hot water heater whose schematic is shown in Figure 1.

Functionally, the hot water heater takes cold water and gas as inputs and produces hot water as an output; flue gases and heat leakage are also produced as waste outputs. The water temperature is regulated by the controller opening and closing the main gas valve (labeled *stop valve*) to keep the water temperature within the preset limits of 140–180°F. The pilot light is always on and the gas valve operates the main burner in full-on/full-off modes. The controller is operated by the temperature measuring and comparing device. The check valve in the water inlet pipe prevents reverse flow due to overpressure in the hot water system, and the relief valve opens if the system pressure exceeds 100 psi (pounds-force per square inch).

Table 1 shows the results of the FMEA on the stop valve. Observe that the analyst has identified all the ways in which the stop valve can fail and determines the effect of each failure locally, on the stop valve itself, and on the entire system. In most cases the analyst also identifies the possible causes of failure and attempts to either eliminate them from the design or, if that is not possible, mitigate their effects. For example, the stop valve could fail to open or close properly due to corrosion. This, in turn, could be caused by electrolysis between different

2 FMEA, Implementation of

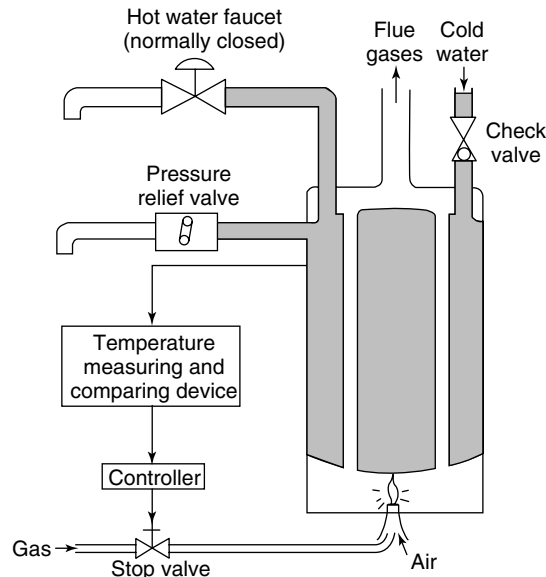


Figure 1 Schematic for a gas hot water heater

metals, or by a reaction of the valve material with the gas. Careful consideration of the valve materials would eliminate this failure mode. (At least it would eliminate this cause of the failure mode; there could be other causes.)

An analysis of the pilot light would reveal the hazard of gas leakage if it fails (it is not lit) and the gas is on. This failure mode cannot readily be eliminated but it can be mitigated by adding an interlock to prevent the gas valve from opening if

the pilot light is not on. Many gas hot water heaters have this type of interlock. If the interlock fails, the burner will not be able to heat the water even if the pilot light is on.

FMEA Methodology

The FMEA methodology is based on a hierarchical, inductive approach to analysis. The basic procedure consists of:

1. Identify all item failure modes;
2. Determine the effect of the failure for each failure mode both locally and on the overall system being analyzed, and if necessary, for each subsystem of which the item is a part;
3. Classify the failure by its effects on the system operation and mission;
4. Determine the failure's probability of occurrence (if possible);
5. Identify how the failure mode can be detected (This is especially important for fault tolerant configurations.);
6. Identify any design changes to eliminate the failure mode, or if that is not possible, mitigate or compensate for its effects.

The results of the analysis are captured on analysis worksheets, which provide a description of the failure modes and their consequences traceable to diagrams or other design documentation. Generally the worksheets include:

- identification of the component being analyzed;

Table 1 FMEA analysis of the stop valve for a hot water heater

Component	Failure mode	Effect	
		Local	System
Stop valve	1. Fails closed	Burner off	No hot water
	2. Fails open	Burner will not shut off	Overheats, release valve releases pressure, may get scalded
	3. Does not open fully	Burner not fully on	Water heats slowly or does not reach desired temperature
	4. Does not respond to controller – stays open	(same as 2)	
	5. Does not respond to controller – stays closed	(same as 1)	
	6. Leaks through valve	Burner will not shut off, burns at low level	Water overheats (possibly)
	7. Leaks around valve	Gas leaks into room	Possible fire or gas asphyxiation

- its purpose or function;
- the component failure mode being analyzed;
- the cause of the failure and how the failure is detected;
- the local, subsystem, and system-level effects of the failure mode;
- the severity classification and probability of occurrence of the failure mode.

A FMEA normally analyzes each item failure as if it were the only failure within the system. However, when the failure is undetectable or latent or the item is redundant, the analysis may be extended to determine the effects of another failure, which in combination with the first failure could result in an undesirable condition.

When analyzing failure effects, the analyst must also be concerned about possible failure cascades and common cause failures where a single event can lead to multiple failures. Such failures often result from the physical placement of the components rather than their operational functions. For example, a failure of a disk turbine in which the disk disintegrates and throws off pieces of broken metal could disable several independent hydraulic systems in an aircraft if they are all routed near the turbine.

The FMEA Process

The FMEA process tracks the system development from concept to detailed implementation. It starts with a functional fault analysis, progresses to an interface analysis, and concludes with a detailed analysis. Each analysis represents a more detailed iteration from the previous analyses. Data needed for the FMEA include the system operating modes and functions, required performance levels, environmental considerations, and safety or regulatory requirements. Libraries with descriptions of failure modes and consequences are also useful. Such libraries provide direction as to the level of detail for the analysis and help ensure consistency in terminology, types of failure modes considered, and so on, among all the analysts (including future analysts) contributing to the project. Ground rules specifying assumptions, limitations, boundary conditions, and what constitutes a failure should also be established before beginning the analysis [5]. The ground rules for the analysis of the hot water heater included the assumption of a continuous supply of gas and water, the water heater is for home use, and it is operated between 140–180 °F.

Functional Fault Analysis

A functional fault analysis is performed on the conceptual design to verify that provisions to compensate for component failures are both necessary and sufficient. The relief valve in the hot water heater is an example of such a compensating provision; it is intended to relieve excessive pressure that potentially could rupture the tank if the water overheats. A functional FMEA of the hot water heater would show that it is needed and that it provides the required protection. The relief valve also illustrates the concern about undetected failures – how does one know that the valve is operating properly? Thus, in addition to showing that the valve is needed, the functional FMEA might also result in a requirement for periodic inspections of the relief valve or a requirement for a monitor to detect a failure of the valve.

Typical functional failure modes are of the form: “function fails to perform”, or “function is performed at the wrong time”. In the hot water heater example, the function of the stop valve is to control the flow of gas in full-on/full-off mode. Thus its functional failure modes are a failure to control the flow of the gas; for example, the gas is on when it should be off, or the gas is not full-on or full-off.

Ideally, a functional fault analysis focuses on the functions that an item, or group of items, performs rather than the characteristics of the specific components used in their implementation. In practice, the types of failure modes considered for a function may depend on how the function will be implemented. When a functional analysis is applied to manufacturing processes, typical failure mode categories include manufacturing and assembly operations, receiving inspection, and testing. Process failure modes are described by process characteristics such as part misorientation, part hole off-center, and so on. When the analysis is applied to software typical functional failure modes are, failure to execute, execution at an incorrect time (early or late), or incorrect result [6].

A major benefit of a functional FMEA is that the functional failure modes can be identified in the conceptual design phase before the detailed design has been developed. Typical results of the analysis identify functional failure modes that need to be eliminated or mitigated by changing the functional system design or the requirements [5].

Interface Fault Analysis

An interface fault analysis focuses on determining the characteristics of failures in the interconnections between subsystem modules. The analysis begins by defining the failure modes for the specific type of interface between subsystem elements. A typical electrical failure mode would be “input shorted to ground”; a typical mechanical failure mode would be “hydraulic pressure low”. Software interface failure modes focus on failures affecting the interfaces between disparate software elements or software and hardware elements such as incorrect timing, or errors in the values presented at the interface. Process errors that could result in misalignment or improper connections are important failure modes for a process interface FMEA.

For the hot water heater, typical failure modes that would need to be considered at the temperature/controller interface would include “no signal” and “incorrect temperature transmitted”.

The advantage of a separate interface fault analysis is that it can be performed as soon as the subsystem inputs, outputs, and their interconnections are defined. This is usually well before detailed module designs are available. Typical results of this analysis are interface failure modes that need to be eliminated or mitigated by interface design changes.

Detailed Fault Analysis

A detailed fault analysis is used to verify that the design complies with the system requirements for: (a) failures that can cause the loss of system functions; (b) single point failures; and (c) fault detection and isolation capabilities. A detailed fault analysis on the system hardware (sometimes called a *piece part fault analysis* since it is done on the “piece parts” that comprise the system) is what has traditionally been done in a FMEA. Table 1 is an example of such an analysis. To perform the analysis, a set of component failure modes and their corresponding occurrence ratios are especially useful. For example, the failure modes normally considered for a capacitor, are “open”, “short”, and “leaking”. Failure mode ratios allow the item failure probability to be apportioned among its failure modes to give the failure mode probability of occurrence. Such ratios are best obtained from field data that is representative of the particular item application but generic references

such as [7] can be used for guidance if such data is not available. Failure mode ratios for a particular component type often vary widely depending on the operating environment, manufacturer, application, and other factors.

A detailed software fault analysis is intended to assess the effect of failures on the software functionality or on the variables used in the software. Unlike hardware, the software errors analyzed by a FMEA are not due to “potential failures”; rather, they already exist and will be activated under the “right” set of conditions [5, 6]. Detailed procedures for performing a FMEA on software are given in [7, 8]. Process related failure modes are specific to the manufacturing or maintenance process. For example a wax coating may be applied too thinly to provide the corrosion protection it is intended to provide.

One problem associated with a detailed fault analysis is that it cannot be initiated until the design has matured to the point that detailed schematics and parts lists are available. This means that any major errors found by the analysis are likely to be very expensive to fix. Conversely, errors, even major errors in the design concept, are relatively easy to fix in the early design phase when the functional and interface fault analyses are done.

Identify Failure Consequences

The consequences of a failure mode analyzed in a FMEA are its effects, a classification of the severity of the failure mode based on its system-level effects, and the probability of the failure mode occurrence. The analysis is conducted for all phases of the system operation including normal operating modes, contingency modes, and test modes. After the effects of an item failure mode have been determined corrective actions or compensating provisions must be identified within each applicable operating mode.

Assessment of the failure mode effects must identify the system conditions or operational modes that manifest the anomalous behavior. For example, failures in the landing gear of an airplane that would cause “loss of landing gear extension” will have very different effects if the failure occurs on the ground than if it occurs while attempting to land. Likewise, the discovery mechanism for detecting the failure may be different.

Failure modes are usually classified by an assessment of the significance of the end-effect on the

system operation and mission. The FMEA ground rules developed in the planning stage should provide a ranking and classification system, and appropriate criteria for assessing the severity of failures for the product being analyzed. In many cases a four-level classification system with severity classifications ranging from catastrophic to insignificant is sufficient [2, 4]. When items are redundant and there is no warning that a redundant item has failed, the severity should be assessed as if all of the redundant items have failed.

The Failure Cause Model

A failure cause model is used to combine the causes of the failure mode and evaluate the failure mode probability of occurrence [5, 9]. The failure cause model often takes the form of a fault tree (*see Fault Trees; Fault Tree Analysis for Large Systems*) whose top event is the failure mode of interest. The fault tree can include operational causes of failure, manufacturing process errors, and design errors. Thus, the failure cause model for a failure mode unifies the various types of FMEA for product, process, software, and so on.

To the extent possible, statistical data should be used to determine the failure mode probability of occurrence. If sufficient data is not available it can be estimated from field data, laboratory or simulation studies, experience with similar systems, tables of generic component failure rates, and so on.

Corrective Action Recommendations

Corrective actions are needed for undetectable faults and for faults having significant consequences – for example, unsafe conditions, mission or safety-critical single point failures, adverse effects on operating capability, or high maintenance costs. Corrective actions may not be needed if the risks for the specific consequence(s) of a failure are acceptable based on a low enough probability of occurrence. Corrective actions generally take the form of changes in requirements, design, processes, procedures, or materials to eliminate the design deficiency.

Once a corrective action is implemented, the affected fault analyses must be revised to reflect the new baseline configuration. When design changes to correct a deficiency are not possible or feasible, compensating provisions, often in the form of design

provisions or designated operator actions, must be identified to circumvent or mitigate the effect of the failure when it occurs.

Prioritization

A FMEA on even a moderately sized system usually results in the identification of hundreds of potential failure modes whose effects may range from insignificant to catastrophic. Hence, it is necessary to prioritize them for corrective actions. Most prioritization techniques are based on an assessment of the severity of the failure effect and the probability of the failure mode occurrence. Generally, failure modes within a severity class having probabilities of occurrence above some threshold must be corrected, but the threshold may vary, depending on the severity category. For example, corrective actions may be required for any catastrophic failure mode with a probability of occurrence of more than 10^{-9} , but a failure mode resulting in an operating limitation may be acceptable if its frequency of occurrence is less than 10^{-4} .

The automobile FMEA [3] specifies the use of a risk priority number (RPN) to rank failure modes. This methodology ranks the severity (S), probability of occurrence (O), and likelihood of detection (D) on a scale from 1 to 10 and multiplies them to obtain the RPN. High RPNs are assumed to be more serious than lower RPNs. Use of the RPN to rank failure modes requires caution and good judgment [4].

Automation

When done manually, FMEAs tend to be tedious and time-consuming tasks. Product designers usually prefer to focus on the tasks a system must do and how to design the system to do those tasks well rather than on what can go wrong. Analyzing what will happen when something fails requires viewing the system in a different way and propagating a failure mode from the component to system level is often difficult.

A number of computer-based tools have been developed to aid the analysis process. The simplest do little more than provide a consistent format for the analysis. Others have features that help improve the quality of the analysis such as providing lists of appropriate component failure modes and their apportionments or making completeness and consistency checks. Some specialized programs may

even simulate the system to help the designer evaluate how various component failure modes will affect the system operation and output [10].

Summary

FMEA has evolved from an *ad hoc* technique, dependent on the designer's "experience", to a formal and accepted analysis technique applicable at all phases of the system development process. Failure modes can be described functionally and their effects analyzed early in the design phase. Interface and detailed FMEAs can be applied as the design detail develops. Fault trees can be used to combine the contributions to the failure mode due to design errors, material properties, manufacturing errors, service mistakes, and even user errors to assess the probability with which the failure mode occurs. Failure modes can be eliminated by removing their causes or at least have their probabilities of failure reduced to acceptable levels.

References

- [1] Coutinho, J.S. (1964). Failure-effect analysis, *Transactions of the New York Academy of Sciences* **26**, 564–585.
- [2] Procedures For Performing A Failure Mode Effects and Criticality Analysis, US Mil-Std-1629 (ships), November 1, 1974; US Mil-Std-1629A, November 24, 1980; US Mil-Std-1629A/Notice 2, November 28, 1984.
- [3] Potential Failure Mode and Effects Analysis In Design (Design FMEA) and Potential Failure Mode and Effects Analysis In Manufacturing and Assembly Processes (Process FMEA) Reference Manual, Society of Automotive Engineers, Surface Vehicle Recommended Practice, J1739, August 2002.
- [4] International Electrotechnical Commission, IEC 60812: Analysis Techniques for System Reliability – Procedure for Failure Mode and Effects Analysis (FMEA), January 2006.
- [5] Recommended Failure Mode and Effects Analysis (FMEA) practices for non-automobile applications, Society of Automotive Engineers, Aerospace Recommended Practice, ARP 5580, July 2001.
- [6] Goddard, P.L. (2000). Software FMEA techniques, in *Proceedings Annual Reliability and Maintainability Symposium*, Los Angeles, pp. 118–123.
- [7] FMD-97. (1997). *Failure Mode/Mechanism Distributions*, Reliability Analysis Information Center.
- [8] Ozarin, N. (2004). Failure modes and effects analysis during design of computer software, in *Proceedings Annual Reliability and Maintainability Symposium*, Los Angeles, pp. 201–205.
- [9] Krasich, M. (2000). Use of fault tree analysis for evaluation of system-reliability improvements in design phase, in *Proceedings Annual Reliability and Maintainability Symposium*, Los Angeles, pp. 1–7.
- [10] Price, C.J., Pugh, D.R., Wilson, M.S. & Snooke, N. (1995). The flame system: automating electrical Failure Mode & Effects Analysis (FMEA), in *Proceedings Annual Reliability and Maintainability Symposium*, Washington, DC, pp. 90–95.

Related Articles

Fault Trees; Fault Tree Analysis for Large Systems; Reliability, Safety, and Risk Management; System Reliability: Computational Algebra Methods.

JOHN B. BOWLES

Fault Tree Analysis for Large Systems

Introduction

Fault tree analysis (*see* **Fault Trees**) is commonly used for the quantification of system failure probability or frequency. The analysis, which for all but the simplest of system structures requires a software, is traditionally performed using the two-stage method of kinetic tree theory developed by Vesely [1]. Stage one of the computer implementation manipulates the fault tree structure using the rules of Boolean algebra to determine what are known as *minimal cut sets* (*see* **System Reliability: Computational Algebra Methods**). A *cut set* is a list of component-failed states which result in system failure. The more useful concept is that of a *minimal cut set*, which is a list of necessary and sufficient component failure events that result in system failure. In a logic equation expressing the system failure causes, the events within a minimal cut set are combined with an AND ($\cdot$) as they are all required to occur. Since any of the N_c minimal cut sets will result in a system failure, the top event, T , (system failure mode) can be expressed as:

$$T = C_1 + C_2 + \cdots + C_{N_c} \quad (1)$$

where C_i are the minimal cut sets and “+” represents OR.

The process of generating the minimal cut sets is very simple and can be performed top down or bottom up on the fault tree [2]. Cut sets are formed by sequentially substituting out gate events in the causes of the system failure, by events appearing lower down in the tree structure. This continues until only basic events remain. The cut sets are then minimized. During this activity, the software has to do many comparisons between the different events within each cut set and then between the cut sets themselves applying the laws of idempotence (equations 2 and 3) and absorption (equation 4). The two idempotence laws ensure that repeated events do not appear within the component failure combinations (equation 2) nor do repeated failure combinations occur (equation 3). Equation (4) deletes any nonminimal failure combinations and

reduces the logic expression to the minimal cut sets. Although the process is simple, since the implementation requires many comparisons, it is time consuming.

Idempotent

$$A \cdot A = A \quad (2)$$

$$A + A = A \quad (3)$$

Absorption

$$A + A \cdot B = A \quad (4)$$

A second difficulty which has to be overcome is the way the information is stored within the computer memory during processing. One approach is to form a two-dimensional array (table) where the rows provide a list of the cut sets and the entries in the row (columns) contain the component failures which appear in the cut sets. This approach requires the number of columns to be set at the maximum number of components which occur in any cut set (prior to any minimization). Since many, usually most, of the minimal cut sets will not need this amount of storage space, it becomes inefficient and may result in insufficient memory to solve the problem. The housekeeping routines required to overcome this again increases the processing time.

Stage two of the analysis is then to use the minimal cut sets, together with the component failure characteristics to predict the failure probability or the failure frequency of the system. Having obtained the minimal cut sets, we can express the top event logic equation as the disjunction (OR) of the minimal cut sets (equation 1). The system failure probability, Q_{sys} is then the probability of this disjunction given by the inclusion–exclusion expansion:

$$Q_{sys} = P(T) = P(C_1 + C_2 + \cdots + C_{N_c}) \quad (5)$$

$$\begin{aligned} Q_{sys} = & \sum_{i=1}^{N_c} P(C_i) - \sum_{i=2}^{N_c} \sum_{j=1}^{i-1} P(C_i \cap C_j) \\ & + \sum_{i=3}^{N_c} \sum_{j=2}^{i-1} \sum_{k=1}^{j-1} P(C_i \cap C_j \cap C_k) - \cdots \\ & \cdots + (-1)^{N_c+1} P(C_1 \cap C_2 \cdots \cap C_{N_c}) \quad (6) \end{aligned}$$

2 Fault Tree Analysis for Large Systems

This process too is computationally intensive. Consider a moderate sized fault tree which delivers 100 000 minimal cut sets. The number of elements in first summation term of equation (6) is 10^5 , in the second term $\approx 5 \times 10^9$ and in the third term $\approx 1.7 \times 10^{14}$ and so on for the 10^5 terms in the equation. Even with modern, fast, digital computers, this is an enormous number of calculations and will take a considerable time to complete. In practice, approximations are required for the quantification.

Binary Decision Diagrams

A method that overcomes the deficiencies identified in the section titled “Introduction” is the binary decision diagram (BDD) method. Figure 1 shows the form of the BDD [3], which like the fault tree, represents the causes of failure of the system in terms of the failure of its basic events or components. In this case, the BDD shows how failures of basic events A, B, and C combine to cause system malfunction.

The entry point on the BDD is the root vertex. The intermediate nodes all correspond to components in the system and two paths leave each node: a “1” branch representing component failure and a “0” branch for component functioning. To construct the diagram, an “ordering” must be assigned to the basic events, which specifies the order in which each component is considered when constructing the diagram. A path on the diagram starts at the root vertex and progresses along a specific output branch from each node which determines the status of that particular component. As soon as the component conditions specified on any path establish the system status, then the path is terminated with the appropriate terminal-0 or terminal-1 node for system functioning

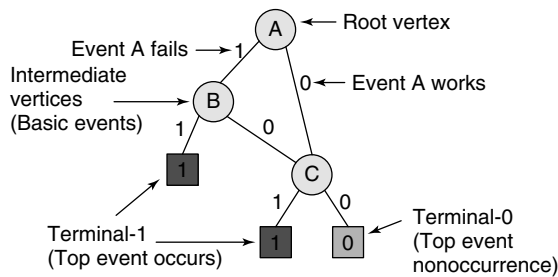


Figure 1 Binary decision diagram structure

or system failure respectively. The concise nature of these diagrams is appealing from the mathematical perspective as they possess characteristics which result in efficient and accurate analysis. From the engineering perspective, they are not documented and do not develop the failure logic in any structured way. As such, it is usually best to develop the system failure logic as a fault tree and then convert to a BDD for analysis.

Binary Decision Diagram Construction

The most efficient fault tree conversion process manipulates equations using a special ite (if–then–else) form. Consider the basic event represented by a node X.

IF the event X occurs

THEN we pass down the “1-branch” and consider the Boolean function, f1, represented by the BDD connected to this branch.

ELSE (the event has not occurred), we pass down the “0-branch” and consider the Boolean function, f2, represented by the BDD connected to this branch.

This can be written in a very concise way:

$$\text{ite}(X, f1, f2) \quad (7)$$

With the ite notation allowing us to express each node on the BDD, mathematical rules can be established which are applied in a bottom-up fashion to the fault tree and deliver the BDD form. Initially each basic event in the fault tree is expressed in the ite notation. So for any basic event X we have:

$$X = \text{ite}(X, 1, 0) \quad (8)$$

When we encounter a gate of logic type $\oplus$, (where $\oplus$ = AND or OR), which has two inputs G and H already expressed in ite form:

$$G = \text{ite}(X, g1, g2) \text{ and } H = \text{ite}(Y, h1, h2) \quad (9)$$

Then the following rules [4] can be applied to get an ite expression for the gate event. It depends on where the basic events X and Y appear in the ordering. Either $X < Y$ (X appears before Y) or $X = Y$ (X

and Y are the same variable). Then:

$$G \oplus H = \text{ite}(X, g1 \oplus H, g2 \oplus H) \quad (10)$$

$$\text{Or}G \oplus H = \text{ite}(X, g1 \oplus h1, g2 \oplus h2) \quad (11)$$

Once an ite expression is formed for any gate it is simplified as much as possible using the Boolean identities:

$$G + 1 = 1 \quad (12a)$$

$$G + 0 = G \quad (12b)$$

$$G \cdot 1 = G \quad (12c)$$

$$G \cdot 0 = 0 \quad (12d)$$

As an example of the construction process, consider the fault tree illustrated in Figure 2. We will use the rules expressed in equations (10) and (11) and an ordering of the basic events $A < B < C$ to convert this to a BDD. First all basic events are put into ite form:

$$\text{ite}(A, 1, 0) \quad \text{ite}(B, 1, 0) \quad \text{ite}(C, 1, 0)$$

Considering now GATE1:

$$\begin{aligned} \text{GATE1} &= A + C \\ &= \text{ite}(A, 1, 0) + \text{ite}(C, 1, 0) \end{aligned}$$

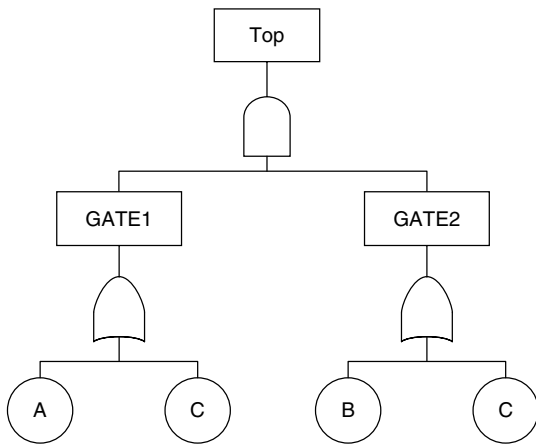


Figure 2 Example fault tree structure

$$= \text{ite}(A, 1 + \text{ite}(C, 1, 0), 0 + \text{ite}(C, 1, 0))$$

using equation 10

$$= \text{ite}(A, 1, \text{ite}(C, 1, 0))$$

using equation 12a, b (13)

In exactly the same way GATE2 can be expressed as:

$$\begin{aligned} \text{GATE2} &= B + C \\ &= \text{ite}(B, 1, \text{ite}(C, 1, 0)) \end{aligned} \quad (14)$$

Having now established expressions for GATE1 and GATE2, the process can now consider the top event:

$$\begin{aligned} \text{TOP} &= \text{GATE1} \cdot \text{GATE2} \\ &= \text{ite}(A, 1, \text{ite}(C, 1, 0)) \cdot \text{ite}(B, 1, \text{ite}(C, 1, 0)) \\ &\text{using equation (10), gives:} \\ &= \text{ite}(A, 1 \cdot \text{ite}(B, 1, \text{ite}(C, 1, 0))), \\ &\quad \text{ite}(C, 1, 0) \cdot \text{ite}(B, 1, \text{ite}(C, 1, 0))) \end{aligned} \quad (15)$$

applying equation (12c) to the 1-branch and equation (10) to the term on the 0-branch of the above expression gives:

$$\begin{aligned} &= \text{ite}(A, \text{ite}(B, 1, \text{ite}(C, 1, 0))), \text{ite}(B, 1 \cdot \text{ite}(C, 1, 0), \\ &\quad \text{ite}(C, 1, 0) \cdot \text{ite}(C, 1, 0))) \end{aligned}$$

Simplifying

$$\begin{aligned} &= \text{ite}(A, \text{ite}(B, 1, \text{ite}(C, 1, 0))), \text{ite}(B, \text{ite}(C, 1, 0), \\ &\quad \text{ite}(C, 1, 0))) \\ &= \text{ite}(A, \text{ite}(B, 1, \text{ite}(C, 1, 0))), \text{ite}(C, 1, 0)) \end{aligned} \quad (16)$$

Equation (16) represents the ite expression for the final BDD. It can be seen that the root vertex is A. On its “1-branch” it has the expression $\text{ite}(B, 1, \text{ite}(C, 1, 0))$, i.e. it is connected to the vertex B which has 1 and $\text{ite}(C, 1, 0)$ on its 1 and 0 branches respectively. The “0-branch” of A is connected to vertex C. This produces the BDD illustrated in Figure 1.

The ite structure lends itself very well to computer implementation as each node on the BDD can be represented as an ordered triple: the variable name, a pointer to the node on the 1-branch, and a pointer to the node on the 0-branch. The whole BDD is a list of its nodes.

System Failure Probability

It is with the system quantification that the BDD scores heavily over the conventional fault tree analysis method. Due to the way that the BDD is formed, it has the characteristic that all paths through the BDD are disjoint. Consider any intermediate node X . A path will pass this node on either the 1-branch meaning that the event has occurred or the 0-branch meaning that it has not. These conditions are mutually exclusive and so all failure causes progressing below the 1-branch therefore contain the component-failed state and will be mutually exclusive to any path through the 0-branch on which the component is considered to function. The branching at all other variables results in all paths being disjoint.

Since the BDD expresses the system structure in a disjoint form, the system failure probability is then simply the sum of the probability of all paths from the root node to a terminal-1 node (accounting for the failure/success of all components included in the path):

$$Q_{\text{sys}} = \sum_{\text{all paths}} P(\text{path } i \text{ to terminal } - 1) \quad (17)$$

For the BDD shown in Figure 3, the disjoint paths to a terminal-1 node are:

1. $A \cdot C \cdot B$
 2. $A \cdot \bar{C} \cdot B \cdot D$
 3. $\bar{A} \cdot B \cdot D$
- (18)

Summing the probabilities of these paths gives:

$$\begin{aligned} Q_{\text{sys}} &= q_a q_b q_c + q_a (1 - q_c) q_b q_d \\ &\quad + (1 - q_a) q_b q_d \\ &= q_a q_b q_c + q_b q_d - q_a q_b q_c q_d \end{aligned} \quad (19)$$

The points to note about this process are:

1. the top event probability obtained is exact [5] (approximations are not needed on any size of problem) and
2. the minimal cut sets were not required as an intermediate step in the calculations [6].

This gives the BDD approach advantages in terms of both accuracy and efficiency over the conventional fault tree analysis method.

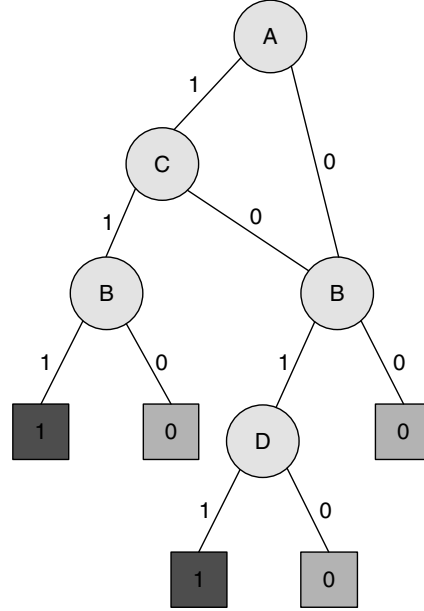


Figure 3 Example BDD

The procedure to evaluate the system failure frequency from the BDD is given in references [3] and [5].

Qualitative System Analysis

As indicated in the previous sections the minimal cut sets are not needed to quantify the system reliability performance. They do however provide useful information in their own right indicating how the system fails and can be determined from the BDD structure.

Each path leading to a terminal-1 node contains the component conditions required for system failure. If the working component states are ignored and only the failed component states are listed, then this corresponds to a cut set. For the BDD shown in Figure 1, there is a path to a terminal 1 passing through nodes A and B on their 1-branches. So AB is a cut set. There are two paths passing through the BDD to the second terminal-1 node. The first passes through the 1-branch of A, the 0-branch of B, and the 1-branch of C, giving cut set AC. The second passes through the 0-branch of A and the 1-branch of C, giving cut set C. Component failure combinations leading to system failure are:

1. AB
2. AC
3. C

Since we can remove component A from combination 2 and the system remains in the failed state, we can see that this is not a minimal cut set. Removing this from the list gives the minimal cut sets {AB}, {C}.

The above approach demonstrates the way minimal cut sets can be obtained in a system assessment; however, for a full minimal cut set analysis, the efficient process would be to construct a zero-suppressed BDD, which encodes only minimal cut sets. Details can be found in reference [7].

Variable Ordering Schemes

The sequence specified to determine the order in which the component events are considered for inclusion in the BDD structure plays a vital role in determining the size, and therefore the efficiency of the BDD produced to represent the logic function. Consider the fault tree given in Figure 4.

With the two ordering schemes of $A < B < C < D$ and $D < C < B < A$, the resulting BDDs are illustrated in Figure 5(a,b) respectively.

The BDD in Figure 5(a) has four intermediate nodes and produces three cut sets (all minimal). The BDD in Figure 5(b) is not such an efficient representation producing five cut sets (two nonminimal). For a problem of this size, the efficiency of the representation of the system failure logic is not important. However, as the size of the system grows, this can be

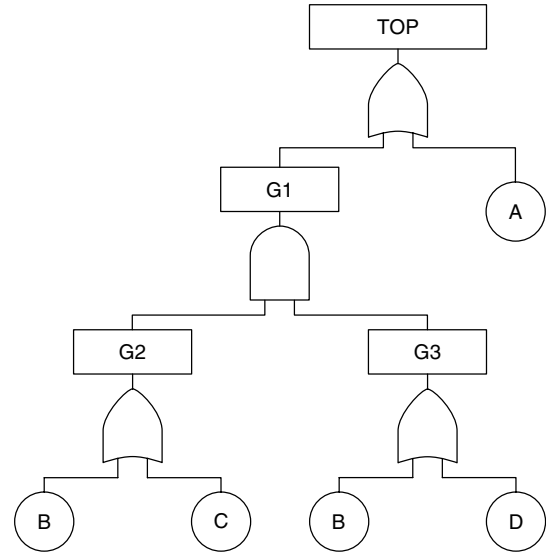


Figure 4 Fault tree to investigate ordering schemes

critical if a BDD is to be constructed. There is no universally accepted way in which the variable ordering can be specified; a number of approaches are possible. The problem then becomes: given the characteristics of the original fault tree, select which strategy should be used to specify the variable ordering. **Neural networks** have been used for this selection process (references [8–10]).

Potential strategies to order the fault tree basic events are:

- “Neighborhood” methods where the fault tree is traversed in a specified way and basic events listed as they are encountered. One of the most

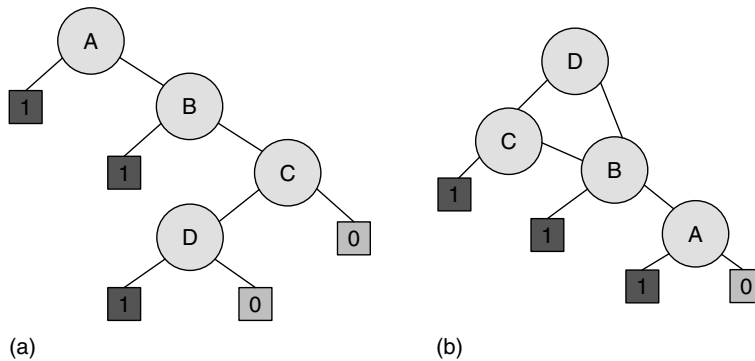


Figure 5 BDDs for different basic event ordering schemes

common ordering methods falls in this category and is known as *top down, left right* where the tree is traversed from top to bottom and on each level encountered from left to right listing the variables as they are encountered for the first time.

- “Structural importance” methods allow nodes to be selected from anywhere in the tree structure. Nodes are allocated a “weighting” which indicates their contribution to the top event. Highest “weightings” are ordered first.

Summary

This article has described the basic properties of the BDD technique which provides a fast, efficient, and accurate method by which fault trees can be analyzed. The implementation of this method has increasing benefits with the size of the problem and offers significant advantages for fault trees constructed to represent the failure causes of large complex systems. In using the BDD method, the system variables, basic events in the fault tree, need to be placed in an ordering. Care should be taken with the process to obtain a concise BDD structure.

References

- [1] Vesely, W.E. (1970). A time dependent methodology for fault tree evaluation, *Nuclear Design and Engineering* **13**, 337–360.
- [2] Andrews, J.D. & Moss, T.R. (2002). *Reliability and Risk Assessment*, Professional Engineering Publications.
- [3] Andrews, J.D. (2006). Fault tree analysis using binary decision diagrams, in *Tutorial Notes, Annual Reliability and Maintainability Symposium*, Newport Beach, January 23–26.
- [4] Rauzy, A. (1993). New approaches for fault tree analysis, *Reliability Engineering and System Safety* **03**(40), 203–211.
- [5] Sinnamon, R.M. & Andrews, J.D. (1997). Improved accuracy in quantitative fault tree analysis, *Quality and Reliability Engineering International* **13**, 285–292.
- [6] Sinnamon, R.M. & Andrews, J.D. (1997). Improved efficiency in qualitative fault tree analysis, *Quality and Reliability Engineering International* **13**, 293–298.
- [7] Minato, S. (1993). Zero-suppressed BDDs for set manipulation in combinatorial problems, in *Proceedings of the 30th ACM/IEEE design Automation Conference, DAC'93*, Dallas. pp. 272–277.
- [8] Bartlett, L.M. & Andrews, J.D. (1999). Efficient basic event ordering schemes for fault tree analysis, *Quality and Reliability Engineering International* **15**(2), 95–103.
- [9] Bartlett, L.M. & Andrews, J.D. (2001). An ordering heuristic to develop the binary decision diagram based on structural importance, *Reliability Engineering and System Safety* **72**, 31–38.
- [10] Bartlett, L.M. & Andrews, J.D. (2002). Selecting an ordering heuristic for the fault tree to binary decision diagram conversion process using neural networks, *IEEE Transactions on Reliability* **51**(3), 344–349.

JOHN D. ANDREWS

Integrating Computer and Physical Experiment Data

Calibration and Prediction for Scalar Outputs

This article provides motivation for a statistical approach to *calibrate* a physics simulation model to experimental data and to *predict* a future experiment given information from calculations and data already in hand. These notions will be made more specific as this discussion proceeds. The statistical framework for the calibration and prediction methodology outlined below is Bayesian [1–3]. To begin, assume that experimental data y , as a function of *system* inputs $\mathbf{x}$, follows a basic measurement error model,

$$y(\mathbf{x}) = \zeta(\mathbf{x}) + \epsilon(\mathbf{x}) \quad (1)$$

Here, $y(\mathbf{x})$ represents the data actually observed, $\zeta(\mathbf{x})$ represents the *true, unknown* response of the physical system, and $\epsilon(\mathbf{x})$ represents observation error. The system inputs $\mathbf{x}$ are known to the scientist with certainty. They represent such quantities as physical constants (e.g., ambient temperature, nominal density) or variables specific to the geometry of the physical system under study, such as spatial coordinates or time. The concepts involved in what follows will be illustrated for the simple physical setting shown in Figure 1. The *observed* datum $y(x)$ is a measurement of the time required for an object to impact the ground after being dropped at rest from height x . In this setting there is one system input (height), but in general multiple system inputs are allowed. The system response $\zeta(x)$ represents the true underlying physical relationship between height x and impact time. This function is generally not observable, due to measurement error and other components of uncertainty inherent in collecting data, represented broadly here by the error term $\epsilon(x)$. The function $\zeta(x)$ is an object of statistical inference.

For complicated physical systems, actual experimental data may only exist for a very small number of $\mathbf{x}$ values. Replicated data may not be available. In this situation, inference about complex functions $\zeta(\mathbf{x})$ is difficult or impossible if based on the observed data alone. For this reason, scientists build intricate simulation models that incorporate advanced physics

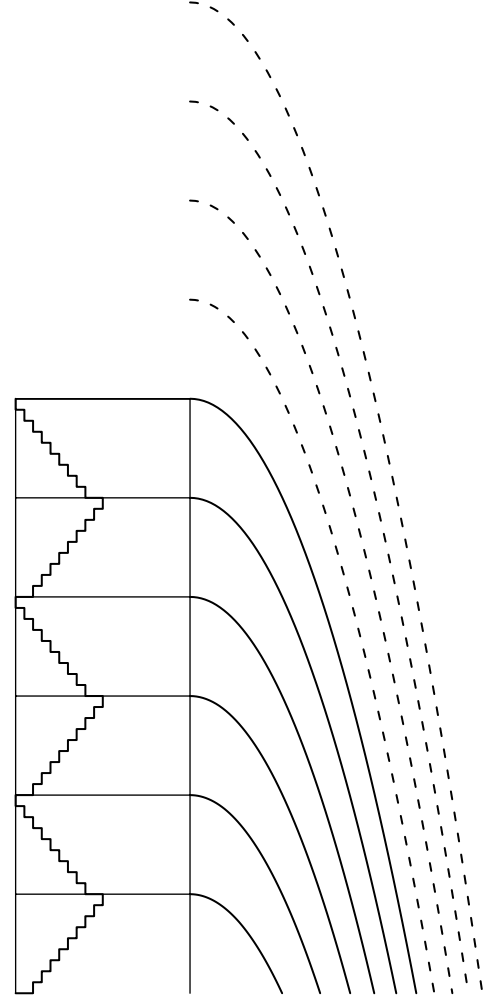


Figure 1 The time taken for an object to drop from each of six floors of a tower (solid lines) is recorded. Predictions of unobserved times are desired for drops from higher floors (dashed lines)

with the goal of effectively modeling $\zeta(\mathbf{x})$. These simulation models produce output as a function of the system inputs $\mathbf{x}$, but they often include additional parameters $\mathbf{t}$ that are commonly used in practice to tune the simulation model to experimental data. These additional parameters are often referred to as *knobs*. They may represent actual physical concepts in the mathematical description underlying the simulation model, or they may represent abstractions employed to account for physical concepts that cannot be or are not explicitly modeled, or they may

2 Integrating Computer and Physical Experiment Data

represent parameters in an empirical model serving as the representation of a physical system. Whatever these parameters represent, in contrast to the system inputs $\mathbf{x}$, they possess uncertainty in their values. *Tuning* refers to any methodology designed to find optimal point estimates of the knobs $\mathbf{t}$, with optimal being defined in terms of minimizing a criterion function such as the sum absolute or squared differences between experimental data and calculations.

For a simulation model that can be made to perfectly represent the physical system through appropriate choice of knob settings $\mathbf{t} = \boldsymbol{\theta}$,

$$\zeta(\mathbf{x}) = \eta(\mathbf{x}, \boldsymbol{\theta}) \quad (2)$$

where $\eta(\mathbf{x}, \boldsymbol{\theta})$ denotes the simulator output at system inputs $\mathbf{x}$ and the true value of the knobs $\boldsymbol{\theta}$. In application, the true value $\boldsymbol{\theta}$ is unknown and must be estimated, traditionally using tuning criteria such as those suggested previously. In the example introduced above, the simulation model implements a differential equation that provides the relationship between the position of the object and time:

$$\frac{d^2s}{d\tau^2} = -1 - \theta \frac{ds}{d\tau} \quad (3)$$

where s and τ represent position and time respectively. The initial conditions are given by

$$s(0) = x \text{ and } \left. \frac{ds}{d\tau} \right|_{\tau=0} = 0 \quad (4)$$

The constant term represents acceleration due to gravity. The height x is assigned to the initial position of the object, and the model output $\eta(x, \theta)$ is the root of the equation $s(\tau) = 0$. The knob θ is called the *drag coefficient*, and it represents the effect of air resistance on the falling object. The value of this parameter is clearly uncertain, as it depends on environmental conditions in effect when the object is dropped.

In many applications, the simulation model does not contain all of the relevant physics to explain the physical system being represented. This can be due to lack of knowledge about applicable physics for certain components of the system or inability to implement advanced physics due to computational complexity. A corollary to this situation is a simulation model that contains all of the (known) relevant physics but for which simplifications are required (such as collapsing over dimensions or choice of mesh) for computational tractability.

Many applications involve inadequate physics *and* computational simplification. In this reality, a model for the physical system can be specified as follows,

$$\zeta(\mathbf{x}) = \eta(\mathbf{x}, \boldsymbol{\theta}) + \delta(\mathbf{x}) \quad (5)$$

Here, $\delta(\mathbf{x})$ is defined as the difference $\zeta(\mathbf{x}) - \eta(\mathbf{x}, \boldsymbol{\theta})$ between the true physical system response $\zeta(\mathbf{x})$ and the simulation model evaluated at the best value for the knobs $\boldsymbol{\theta}$, $\eta(\mathbf{x}, \boldsymbol{\theta})$. This term is often referred to as *discrepancy* or *model inadequacy*. Nonzero values for $\delta(\mathbf{x})$ indicate that for system inputs $\mathbf{x}$, the simulation model evaluated at optimal settings for the knobs $\boldsymbol{\theta}$ is an incomplete representation of the physical system. For example, suppose the simulation model solves the differential equation $s''(\tau) = -1$, subject to the initial conditions $s(0) = x$ and $s'(0) = 0$. The impact time $\eta(x, 0)$ is the root of the equation $s(\tau) = 0$. This model is appropriate for releasing objects in a vacuum but inadequate under any circumstances where air resistance affects the velocity of the falling object. If the latter applies, $\delta(x)$ will be an increasing function of x as the effect of drag on impact time will increase with drop height.

Figure 2 illustrates the concept of model inadequacy for the situation where the simulator produces output under the assumption of a vacuum ($\theta = 0$) while data has been generated under the assumption of a positive drag coefficient ($\theta = 0.5$) with a small amount of observation error. In the notation of the previous development, the physical system response is $\zeta(x) = \eta(x, 0.5)$ and the discrepancy function is $\delta(x) = \eta(x, 0.5) - \eta(x, 0)$, that is, the difference between the system response and the simulation model $\eta(x, 0)$. In this example, the simulation model is so simple that there are no calibration parameters. All statistical modeling of these data is concentrated on the discrepancy function. It is important to note that statistical modeling does a reasonable job of predicting the system response within the range of available data, but does quite poorly when extrapolating to drop heights beyond the upper range of the data, as indicated by the rapidly increasing spread in the **quantiles** of the predicted system response. This highlights the difficulty of using extrapolations with an inadequate simulator. This behavior is a consequence of the fact that statistical modeling of discrepancy in this example does not incorporate any knowledge of physics to assist predictive capability in data-poor regimes. Statistical discrepancy modeling can be used to identify the possible presence of

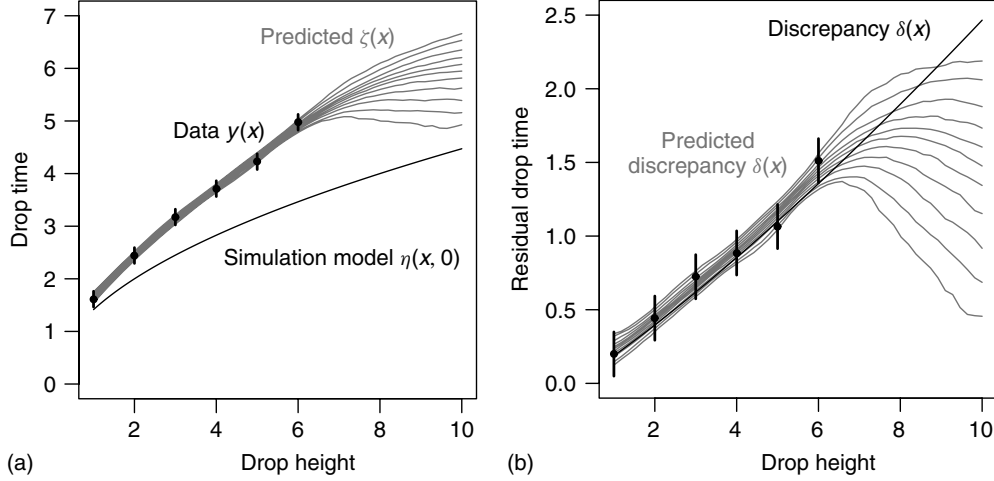


Figure 2 (a) The black trace is the assumed simulation model (vacuum release) and the black dots are data (drag coefficient, 0.5) with two standard deviation error bars. A statistical model was fit to these data ($y(x)$), resulting in 500 predictions of the physical system response $\zeta(x)$. The 11 gray traces are the 0.05, 0.1 (0.1) 0.9, 0.95 quantiles of these predicted curves. (b) The black trace is the discrepancy function $\delta(x)$ and the black dots are residual data ($y(x) - \eta(x, 0)$) with two standard deviation error bars. The 11 gray traces represent the above specified quantiles of 500 discrepancy curve predictions

inadequate physics in simulation models under the following circumstances:

- nonzero discrepancy predictions in the range of the data,
- poor extrapolation of system response predictions.

Of course, absence of these indicators does not imply that the simulation model contains completely adequate physics.

In many applications, there will be no single value of the knobs θ that best matches the data. Many values of θ will give similar results. Multiple experimental datasets will generally be associated with different best estimates of θ . Changing experimental conditions could influence the best θ , such as the degree of laminar or turbulent flow on the best drag coefficient. Different values of θ may be appropriate across changing sets of conditions in dynamic simulations; for example, best material strength parameters varying as a function of strain rate and temperature. For these reasons, the statistical approach to inference about θ assumes that initial uncertainty about the value of θ is represented as a probability distribution. This prior assessment of θ is updated using experimental data and simulation runs to yield a refined assessment of uncertainty in θ . This process

is referred to as *calibration* of the simulation model to experimental data. It is distinct from tuning in that the end product is a probability distribution for θ , while tuning produces a single best estimate for θ . The calibration approach to gaining information about θ recognizes that typically there will be residual uncertainty in knowledge about θ even after updating this knowledge with data and calculations. This remaining uncertainty should be accounted for in any follow-on analysis (such as prediction) involving the knob parameters.

A common inferential problem at this stage is to *predict* the value of experimental data at system inputs x , which are generally composed of input conditions at which no experiment has previously been conducted. That is, the goal is to provide a best estimate with uncertainty bounds for the value of data that would be generated by a future experiment. Several prediction problems are of interest: prediction of physical system response $\zeta(x)$ using the *calibrated simulator*, where the refined uncertainty in the knob parameters from the calibration analysis is essentially propagated through the simulator to predict $\zeta(x)$; prediction of discrepancy $\delta(x)$, where differences between data and simulation model are examined; prediction of $\zeta(x)$ using the discrepancy-adjusted calibrated simulator; and

prediction of experimental data $y(\mathbf{x})$ based on the entire model, which incorporates uncertainties from observation error in addition to the above sources (calibrated simulation and discrepancy). Predictions of discrepancy with uncertainty bars not containing zero can be used as one diagnostic tool for making an assessment of inadequate physics in the simulation model, as discussed previously.

Figure 3 shows results obtained by enhancing the physics of the drop-time simulation model for vacuum environments of Figure 2 to incorporate air resistance. The data set shown in Figure 2 is used to calibrate this improved simulator. Initial uncertainty in the drag coefficient θ was assumed to be uniform on the unit interval $[0, 1]$. The experimental data constrains this uncertainty substantially so that all probability in the calibrated θ is tightly concentrated around 0.5, the value used in generating the data for this example. It is also evident that the reduction in uncertainty of θ due to calibration leads to much more informative (less variable) predictions of the system response. The discrepancy is probabilistically zero in this analysis, so that the calibrated simulator is an effective predictor of the system response. It is evident that extrapolation of the calibrated simulator beyond the upper range of drop heights in the observed data leads to well-behaved predictions of the system response. This is to be expected when the simulation model incorporates all or most of the relevant physics.

Calibration and Prediction for Functional Outputs

The calibration and prediction methodology described above can be extended to functional computer model (and experimental data) outputs. Let $\boldsymbol{\eta}(\mathbf{x}, t)$ denote a vector of computer model output calculated at input specification $(\mathbf{x}, t)$. The output vector can be functional (e.g., velocity as a function of time) or a multivariate collection of selected quantities (e.g., important features extracted from curves). It is assumed that the computer model is expensive to evaluate, so that a limited number of model runs is feasible and a statistical model must be adopted to approximate the computer model output at unsampled inputs based on the information available from the limited runs (*see Computer Experiments*). In particular, the code is run at m different combinations of the inputs $(\mathbf{x}, t)$ as specified by an experimental design (*see Assessment of Experimental Designs*).

Statistical modeling of the computer model output focuses on dimension reduction. An orthogonal basis representation for the code output is specified:

$$\boldsymbol{\eta}(\mathbf{x}, t) = \mathbf{k}_1 w_1(\mathbf{x}, t) + \cdots + \mathbf{k}_{p_\eta} w_{p_\eta}(\mathbf{x}, t) = \mathbf{K} \mathbf{w}(\mathbf{x}, t) \quad (6)$$

The basis vectors $\mathbf{k}_1, \dots, \mathbf{k}_{p_\eta}$ are fixed, while the coefficients $w_i(\mathbf{x}, t)$ are assigned independent, mean-zero Gaussian process (GP) models [4, 5]. A common

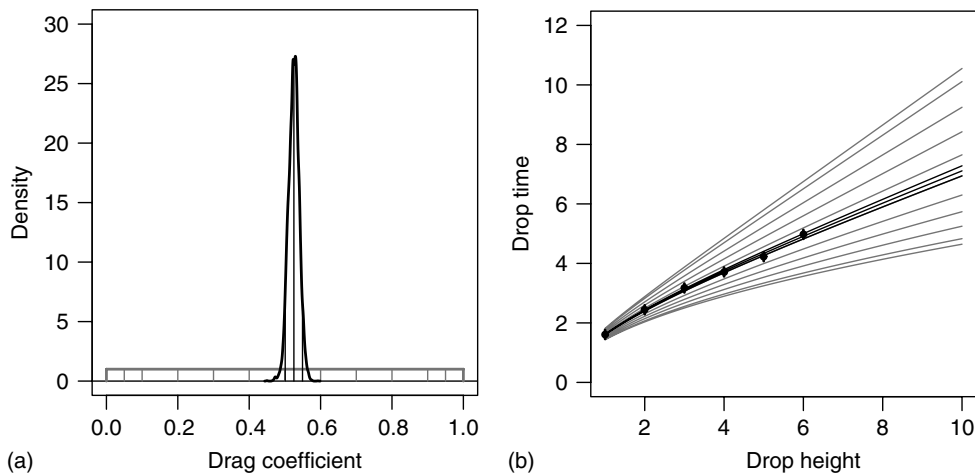


Figure 3 (a) Comparison of the initial uncertainty distribution of the drag coefficient θ (gray; 0.05, 0.1 (0.1) 0.9, 0.95 quantiles indicated) with the calibrated distribution (black; 0.05, 0.5, and 0.95 quantiles indicated). (b) Simulation model output for θ set to the indicated quantiles of its initial distribution (gray) and its calibrated distribution (black)

and efficient approach to obtaining an orthogonal set of basis vectors involves calculating the principal components of standardized computer model output via the singular value decomposition. Often fewer than five components are required to explain over 99% of the variation in computer model output, resulting in considerable dimension reduction and hence more efficient statistical analysis.

The statistical model for discrepancy $\delta(\mathbf{x})$ assumes a basis representation as well,

$$\delta(\mathbf{x}) = \mathbf{d}_1 v_1(\mathbf{x}) + \cdots + \mathbf{d}_{p_\delta} v_{p_\delta}(\mathbf{x}) = \mathbf{D} \mathbf{v}(\mathbf{x}) \quad (7)$$

For functional data, the basis vectors $\mathbf{d}_1, \dots, \mathbf{d}_{p_\delta}$ are often chosen to be local in extent (such as a kernel basis) to accommodate nonstationarity that often arises in the discrepancy process. The coefficients $v_i(\mathbf{x})$ are assumed to be independent, zero-mean GP models. Depending on the application, the coefficients may be partitioned into groups having common GP models. The discrepancy and computer model coefficients are assumed *a priori* to be independent.

Each field experiment may produce different amounts of functional data. The code and discrepancy

bases must be interpolated onto the grid of locations available for each of the field experiments. The statistical model for field experiment i is thus given by

$$\mathbf{y}(\mathbf{x}_i) = \mathbf{K}_i \mathbf{w}(\mathbf{x}_i, \boldsymbol{\theta}) + \mathbf{D}_i \mathbf{v}(\mathbf{x}_i) + \boldsymbol{\varepsilon}(\mathbf{x}_i) \quad (8)$$

where $\boldsymbol{\varepsilon}(\cdot)$ represents the observational error model.

These developments are illustrated with computer simulations of a tantalum flyer plate experiment. These experiments are conducted to explore the behavior of materials subjected to high strain rates. Flyer plate experiments involve forcing a plane shock wave through stationary test samples of material and measuring the free surface velocity at the back side of the target as a function of time. This experiment can be modeled as a function of 10 parameters governing the equation of state (EOS), material strength, and spall strength of the target material [3]. Material strength is incorporated through the Preston–Tonks–Wallace (PTW) model of plastic deformation [6]. Seven PTW model parameters are calibrated.

The simulator was run at 128 unique combinations of the 10 calibration parameters as specified by an orthogonal array (OA)-based Latin hypercube design (*see Orthogonal Arrays; Latin Hypercube*

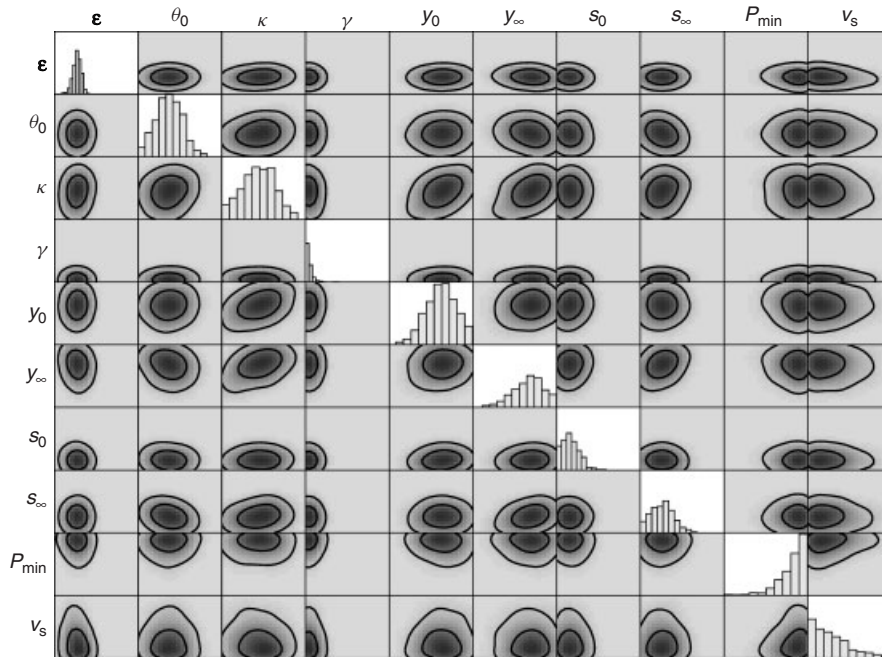


Figure 4 Marginal and bivariate posterior density estimates of 10 calibration parameters. Black contours represent 90% and 50% bivariate posterior probability regions

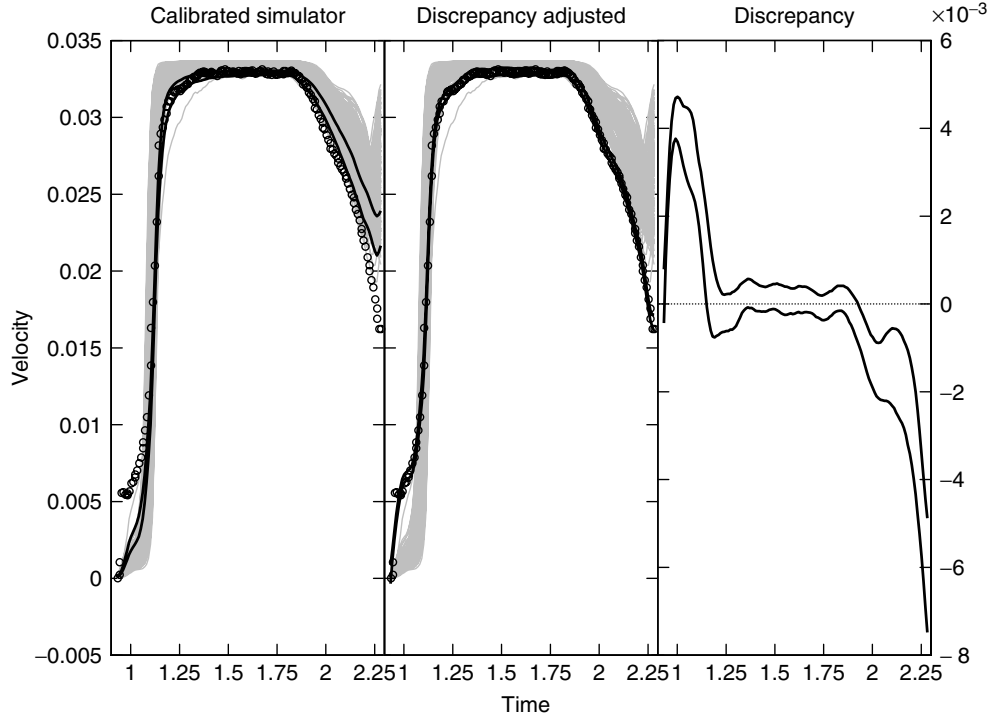


Figure 5 Assessment of calibrated code predictions and discrepancy. Calculated velocity profiles (gray curves) and velocity measurements (black dots). The black curves give 5%/95% pointwise probability intervals for the calibrated code, system response, and discrepancy predictions (a, b, and c)

Designs) generated from a two-level, strength 3 OA design [7]. For each run in the experimental design, a velocity profile (velocity as a function of time) is calculated. A basis consisting of the first three empirical orthogonal functions computed from the 128 velocity profiles was adopted to represent the simulation outputs. A basis consisting of 31 localized Gaussian kernels was adopted to accommodate considerable nonstationarity in the discrepancy process.

A Bayesian hierarchical model fit of the PTW model to independent stress–strain data provided the prior distribution for all seven material strength parameters, coupled with uniform priors on the initial ranges of the remaining parameters. Figure 4 shows projections of the posterior distributions for all 10 calibration parameters resulting from incorporating the flyer plate experimental data as an additional constraint. Note, in particular, the considerable amount of knowledge gained about the first parameter (EOS perturbation ϵ): it began with a uniform distribution

on $[-5\%, 5\%]$ and calibrates to 95% of its probability on $[-3.34\%, -1.4\%]$.

Figure 5 provides an assessment of computer model performance as measured by the ability of a properly calibrated code to predict experimental data. Substantial discrepancy is present in code predictions of the initial loading phase and throughout the release phase of the flyer plate experiment. The 10 calibration parameters are probabilistically tuned to be consistent with experiment during part of the loading phase and throughout the peak velocity phase.

This article has focused on calibration of uncertain computer model parameters to available experimental data, assessment of calibrated model predictions, and the associated evaluation of model discrepancy. However, the statistical model adopted for these objectives can be used for other analyses as well. In particular, practitioners are often interested in some form of global sensitivity analysis, which involves quantitative assessments of input parameter variations on model output variation [8].

References

- [1] Kennedy, M. & O'Hagan, A. (2001). Bayesian calibration of computer models (with discussion), *Journal of the Royal Statistical Society. Series B* **63**, 425–464.
- [2] Higdon, D., Kennedy, M., Cavendish, J., Cafo, J. & Ryne, R.D. (2004). Combining field observations and simulations for calibration and prediction, *SIAM Journal on Scientific Computing* **26**, 448–466.
- [3] Williams, B., Higdon, D., Gattiker, J., Moore, L., McKay, M. & Keller-McNulty, S. (2006). Combining experimental data and computer simulations, with an application to flyer plate experiments, *Bayesian Analysis* **1**(4), 765–792.
- [4] Sacks, J., Welch, W.J., Mitchell, T.J. & Wynn, H.P. (1989). Design and analysis of computer experiments (with discussion), *Statistical Science* **4**, 409–423.
- [5] Santner, T.J., Williams, B.J. & Notz, W.I. (2003). *Design and Analysis of Computer Experiments*, Springer, New York.
- [6] Preston, D.L., Tonks, D.L. & Wallace, D.C. (2003). Model of plastic deformation for extreme loading conditions, *Journal of Applied Physics* **93**, 211–220.
- [7] Tang, B. (1993). Orthogonal array-based Latin hypercubes, *Journal of the American Statistical Association* **88**, 1392–1397.
- [8] Oakley, J.E. & O'Hagan, A. (2004). Probabilistic sensitivity analysis of complex models: a Bayesian approach, *Journal of the Royal Statistical Society. Series B* **66**, 751–769.

BRIAN J. WILLIAMS AND DAVID M. HIGDON

Markov Chain Monte Carlo, Introduction

Introduction to Monte Carlo Simulation

High-dimensional integrals and simulation problems are seen in areas such as statistics, reliability, statistical physics, and so on. For example, in Bayesian statistics, given a prior distribution $P(\theta)$ and a likelihood function $l(\theta|\text{data})$, then the posterior mean of θ is calculated *via* Bayes theorem as

$$E[\theta|\text{data}] = \frac{\int \theta P(\theta) l(\theta|\text{data}) d\theta}{\int P(\theta) l(\theta|\text{data}) d\theta} \quad (1)$$

In many cases, the integrals in the numerator and denominator cannot be evaluated analytically and various alternatives may be considered. One possibility is to use numerical integration techniques (*see Quadrature and Numerical Integration*), and an alternative is to use analytic approximations of the integrands such as the Laplace approximation, (*see e.g. [1] for a general review and Laplace Approximations in Bayesian Lifetime Analysis in the context of lifetime data analysis*). However, the first method suffers strongly from the curse of dimensionality and approximation methods are only appropriate under restricted conditions. A more general approach is to use simulation.

Suppose then that for a variable X with distribution π that we wish to estimate the expected value of some functional $f(X)$. Then the **Monte Carlo** method, see [2], is to generate a sample of values $x_1, \dots, x_N$ from the density π and estimate the expectation using the sample mean

$$E_\pi[f(X)] \approx \frac{\sum_{i=1}^N f(x_i) \pi(x_i)}{\sum_{i=1}^N \pi(x_i)} \quad (2)$$

It is usually difficult or impossible to directly sample from the density π so that certain alternatives such as importance sampling or rejection sampling may be used. A full review of the Monte Carlo approach is

given in **Monte Carlo Methods**. See also **System Reliability: Monte Carlo Estimation**.

In a large number of cases, however, exact Monte Carlo samples may not be easily constructed. Therefore, an alternative strategy is to try to generate an approximate Monte Carlo sample. One approach is to construct a discrete-time stochastic process (**Markov chain**) such that the distribution of the states of the process converges over time to the distribution of interest π when a sample can be taken and used in the same way as a standard Monte Carlo sample. This is the basis of the Markov chain Monte Carlo (MCMC) algorithms.

Markov Chains

Before introducing the different MCMC algorithms, it is first useful to review the properties of Markov chains. A full review is given in, for example [3].

A Markov chain, $\{X_t\}$, is defined to be a sequence of variables $X_1, X_2, \dots$ such that the distribution of X_t , given the previous values $X_0, \dots, X_{t-1}$ only depends on X_{t-1} , so that

$$P(X_t \in A | X_1, \dots, X_{t-1}) = P(X_t \in A | X_{t-1}) \quad (3)$$

for any set A , where, $P(\cdot|\cdot)$ represents a conditional probability. A Markov chain is further said to be time homogeneous if

$$P(X_{t+k} \in A | X_t) = P(X_k \in A | X_0) \quad (4)$$

for any k . A time homogeneous Markov chain is thereby determined by the distribution or value of the initial state X_0 and by the transition kernel

$$P(x, y) = P(X_{t+1} = y | X_t = x) \quad (5)$$

For most problems of interest to us in the context of MCMC, the Markov chain will take values in some subset of $\mathbb{R}^m$ for some possibly large m . However, in order to illustrate the most important concepts, it is convenient to assume that the state space is countable so that the possible states of the chain can be represented as the nonnegative integers $\{0, 1, 2, \dots\}$. We can then denote the t -step transition probabilities as

$$p_{ij}(t) = P(X_t = j | X_0 = i) \quad (6)$$

2 Markov Chain Monte Carlo, Introduction

Then our interest concerns the conditions under which the t -step transition probabilities converge to a stationary distribution π so that

$$p_{ij}(t) \longrightarrow \pi(j) \quad \text{for all } i, j, \quad \text{as } t \longrightarrow \infty \quad (7)$$

so that the distribution of the states of the chain becomes approximately independent of the initial state or distribution as t increases.

A first necessary condition is that the Markov chain is irreducible. A Markov chain is said to be irreducible if every state can be reached from any initial state, that is if $p_{ij}(t) > 0$ for some t and all i, j . Secondly, all states of the Markov chain must be positive recurrent. A state, i , is said to be recurrent, if the Markov chain is certain to return to i , so if we define $\tau_i = \min\{t > 0 : X_t = i | X_0 = i\}$ to be the time taken to return to state i , then $P(\tau_i < \infty) = 1$. A recurrent state i is said to be positive recurrent if its expected recurrence time is finite, that is if $E[\tau_i] < \infty$. For an irreducible Markov chain, it can be demonstrated that if any state is positive recurrent, then all states are positive recurrent. Finally, all states of the chain must be aperiodic. The period of state i is defined to be $d(i)$ where,

$$d(i) = \text{greatest common divider } \{t : p_{ii}(t) > 0\} \quad (8)$$

For an irreducible Markov chain, it can be shown that all states have the same period. The chain is said to be aperiodic if this period is one.

It can then be shown that for an irreducible, positive recurrent, aperiodic Markov chain, a unique stationary distribution π exists so that for all i, j ,

$$\pi(j) = \lim_{t \rightarrow \infty} p_{ij}(t) \quad (9)$$

Moreover, suppose that X is a random variable with distribution π and that $f(X)$ is a function such that $E_\pi[|f(X)|] < \infty$. Then it follows that if $\{X_t\}$ is a Markov chain with stationary distribution π , then

$$\frac{1}{T} \sum_{t=1}^T f(X_t) \longrightarrow E_\pi[f(X)] \quad (10)$$

as $t \rightarrow \infty$.

A sufficient condition for the existence of a unique stationary distribution is that of reversibility. A Markov chain with transition probabilities $p_{ij} =$

$P(X_{t+1} = j | X_t = i)$ is said to be (time) reversible, if there exists a probability density π that satisfies the detailed balance equation so that for any i, j , then

$$p_{ij}\pi(i) = p_{ji}\pi(j) \quad (11)$$

For a time-reversible Markov chain, then it can be immediately demonstrated that π is then a stationary distribution of the Markov chain because, in this case,

$$\sum_i p_{ij}\pi(i) = \sum_i p_{ji}\pi(j) = \pi(j) \quad (12)$$

The previous ideas can be generalized to the case of a continuous state space, although the conditions are slightly more technical, for details, see [4, 20]. In this case, given a transition kernel $P(x, y)$ such that

$$\int_y P(x, y) dy = 1 \quad (13)$$

then the stationary distribution π satisfies

$$\pi(y) = \int_x \pi(x) P(x, y) dx \quad (14)$$

The objective of the MCMC approach is therefore to construct a Markov chain with a given stationary distribution π . A general algorithm for doing this is discussed in the following section.

The Metropolis–Hastings Algorithm

The Metropolis–Hastings algorithm of [5, 6] is a general method of constructing a Markov chain with a given equilibrium distribution $X \sim \pi$. A full description is given in [7]. The algorithm generates a Markov chain in which each state only depends on the previous state as follows.

1. Given the current value, $X_t = x$, generate a candidate value, y , from a proposal density $q(y|x)$.
2. Calculate the acceptance probability

$$\alpha(x, y) = \min \left\{ 1, \frac{\pi(y)q(x|y)}{\pi(x)q(y|x)} \right\} \quad (15)$$

3. With probability $\alpha(x, y)$ define $X_{t+1} = y$ and otherwise reject the proposed value and set $X_{t+1} = x$.
4. Repeat until convergence is judged and a sample of the desired size is obtained.

Firstly, it is clear that the acceptance probability $\alpha(x, y)$ used in the Metropolis–Hastings algorithm only depends on the ratio $\pi(y)/\pi(x)$, which is not dependent upon the integration constant of the density π . Secondly, it is straightforward to demonstrate that π is a stationary distribution of the Markov chain as follows. The transition kernel for the density of a move from x to y is given by

$$P(x, y) = \alpha(x, y)q(y|x) + \left(1 - \int \alpha(x, y)q(y|x) dy\right) \delta_x \quad (16)$$

where δ_x is the Dirac mass at x . Now it is easy to show that

$$\pi(x)q(y|x)\alpha(x, y) = \pi(y)q(x|y)\alpha(y, x) \quad (17)$$

and that

$$\left(1 - \int \alpha(x, y)q(y|x) dy\right) \delta_x = \left(1 - \int \alpha(y, x)q(x|y) dx\right) \delta_y \quad (18)$$

and therefore, we have the detailed balance equation

$$\pi(x)P(x, y) = \pi(y)P(y, x) \quad (19)$$

as in equation (11) which implies that π is the stationary distribution of the chain. Note also that if the proposal density $q(y|x) = \pi(y)$, then the Metropolis–Hastings algorithm simply reduces to standard Monte Carlo sampling.

One might expect that the Metropolis–Hastings algorithm would, in general, be more efficient if the acceptance probability $\alpha(x, y)$ was high. However, this is not the case in general. In [8] it is recommended that for high-dimensional models, the acceptance rate for random walk type Metropolis–Hastings algorithms (see below) should be around 25%, whereas in models of dimension one or two, the acceptance rate should be around 50%. However, general results are not available and the efficiency of a Metropolis–Hastings algorithm will generally be heavily dependent on the proposal density $q(y|x)$. Various possibilities are considered in the following section.

Different Types of MCMC Samplers

Theoretically, the Metropolis–Hastings algorithm can be implemented with almost any proposal density $q(\cdot|\cdot)$. However, the performance of the algorithm can be influenced by the form of this density. Usually, the proposal density is chosen so that it is easy to sample from.

One possibility is to use an independence sampler, see e.g. [9] that is $q(y|x) = q(y)$ independent of x . This will often work very well if the density q is similar to π , although with somewhat heavier tails, in a similar way, to the rejection sampler in standard Monte Carlo sampling.

Another alternative is the Metropolis [5] sampler which has the property that $q(x|y) = q(y|x)$ and has the advantage of simplifying the acceptance probability in equation (15) to $\alpha(x, y) = \pi(y)/\pi(x)$. A special case is the random walk Metropolis algorithm, which assumes that $q(y|x) = q(|y - x|)$.

In many cases, X will be a high-dimensional random variable and it may be much more convenient to implement the Metropolis–Hastings algorithm in blocks $X = (X_1, \dots, X_k)$. The simplest algorithm of this type is the so-called Gibbs sampler (see [10, 11]). In order to illustrate this approach, suppose initially that $X = (X_1, X_2)$ has joint density $\pi(x_1, x_2)$ and that the marginal densities are $\pi_{X_1}(x_1)$ and $\pi_{X_2}(x_2)$, then the Gibbs sampler proceeds as follows:

1. Set an initial value $X_1 = x_{10}$.
2. For $t = 1, 2, \dots$ generate:
 - (a) $X_{2t} \sim \pi_{X_2|X_1}(\cdot|X_1 = x_{1t-1})$
 - (b) $X_{1t} \sim \pi_{X_1|X_2}(\cdot|X_2 = x_{2t})$
 where $\pi_{X_1|X_2}(\cdot|x_2)$ and $\pi_{X_2|X_1}(\cdot|x_1)$ are the associated conditional distributions so that $\pi_{X_1|X_2}(x_1|x_2) = \frac{\pi(x_1, x_2)}{\pi_{X_2}(x_2)}$ and $\pi_{X_2|X_1}(x_2|x_1) = \frac{\pi(x_1, x_2)}{\pi_{X_1}(x_1)}$.

This algorithm can be easily extended to the case of a higher-dimensional X by successively sampling from the full conditional distributions $\pi_{X_i|X_{-i}}(x_i|x_{-i})$ where $X_{-i} = (X_1, \dots, X_{i-1}, X_{i+1}, \dots, X_k)$ of each component at each step.

It is clear that the Gibbs algorithm generates a Markov chain and it can be shown that it corresponds to a composition of Metropolis–Hastings steps with acceptance probabilities equal to 1 and stationary distribution π . Full descriptions of the Gibbs sampler are given in [12, 20].

A number of methods have been developed for the implementation of Gibbs sampling algorithms. Firstly, the adaptive rejection sampler for log concave densities was introduced in [13] and extended to arbitrary densities in [14]. Generic software, e.g. WinBUGS, for implementation of Gibbs sampling based on these methods is available (see [15]), and software for assessing convergence of Gibbs samplers, e.g. CODA, has also been developed (see [16]).

When full Gibbs sampling is impossible, as some of the conditional distributions cannot be sampled directly, it is possible to use Metropolis–Hastings updates in these steps which produces a hybrid, Metropolis–within–Gibbs sampler, see e.g. [9].

Note finally that as well as the various options discussed here, many alternative MCMC samplers have been developed. Of particular importance, we may cite the slice sampler [17] which is a simple approach for sampling relatively low-dimensional distributions and the reversible jump sampler [18], which has been developed to sample over state spaces of various dimensions and the perfect sampler [19] which, for a restricted range of problems, is designed to avoid the problems of assessing MCMC convergence by generating an exact sample from the target distribution π . See e.g. [20] for more details.

Apart from the perfect sampler, although theoretically, the Metropolis–Hastings algorithm guarantees convergence to the stationary distribution. In practice, this may be very slow and the performance of the algorithm will depend on both the initial values chosen to start the chain and the form of the proposal distribution $q(\cdot|\cdot)$. Different convergence monitoring techniques are discussed in the following section.

Monitoring the Markov Chain

The objective of this section is to present the basic techniques for monitoring the output of a Markov chain for convergence. Full reviews of this area are given in, e.g. [21, 23]. Firstly, it is important to assess when the sampled values X_t have approximately converged to the stationary distribution π . This will depend largely on how well the MCMC algorithm is able to explore the state space and also on the levels of correlation between the X_t 's. Secondly, it is important to assess the convergence of MCMC averages $\frac{1}{T} \sum_{t=1}^T f(X_t)$, as in equation (10), and finally we need to be able to assess how close a given sample is to being independent and identically distributed.

When assessing the performance of an MCMC algorithm, one possibility is to consider running the algorithm various times independently with a variety of widely dispersed starting values. Then, for example, one might assess convergence to equilibrium by examining the time at which the sample averages of functionals $f(X)$ of interest have converged for all the chains. Formal diagnostics for doing this are introduced in, e.g. [22].

An alternative is to use a single run of a Markov chain. In this case, it is always useful to produce simple graphs of the generated values X_t *versus* time, which can be used to assess how quickly or slowly the Markov chain moves or any deviations from stationarity. Furthermore, convergence may be assessed using graphs of sample averages as in equation (10). More formal approaches are discussed in, e.g. [23].

Figure 1 shows the states generated using the Metropolis sampler described in the section titled “Example” when this is run for 5000 iterations starting from the same initial value with standard deviations of 50 and 0.01, respectively.

In the first case, the chain mixes quite badly; only 3% of the proposed values are accepted and the chain spends relatively long periods in the same state before jumping to another value. In the second case, almost all of the proposed values (around 99%) are accepted, but it is clear that the chain has not explored the full sample space. We can improve the mixing of the chain by changing the standard deviation of the proposal distribution. It is often useful, in general, to use the first part of a Metropolis–Hastings run to tune the sampler (for example, by varying the variance of the proposal distribution) to achieve reasonable mixing and movement through the state space so that equilibrium can be reached.

Even when convergence has been assessed, it is important to take into account that for any Markov chain–based sampler, the output is correlated. Although this does not affect the estimation of the mean, a standard Monte Carlo estimate of the variance of X ,

$$\frac{1}{T} \sum_{i=1}^T (x_i - \bar{x})^2 \quad (20)$$

will underestimate the true variance because of the positive correlation between the sampled data. One method of taking this into account is to calculate the

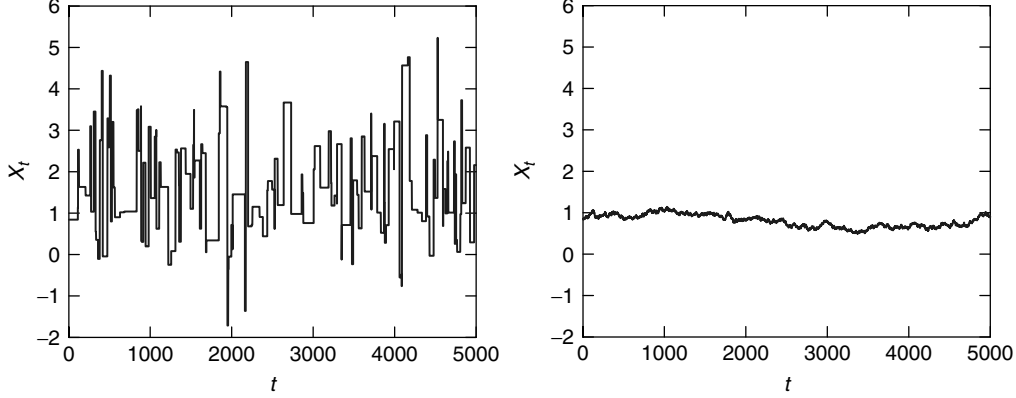


Figure 1 Metropolis samplers showing poor mixing properties

effective sample size T/κ , where $\kappa = 1 + 2 \sum_{i=1}^{\infty} \rho_i$, where ρ_i is the autocorrelation (see **Autocorrelation Function**) of lag i (see [24]). For example, in the case discussed earlier of running the Metropolis sampler with $\sigma = 50$, for 5000 iterations, the effective sample size is approximately 125.

A full review of the different strategies for improving the convergence and mixing of MCMC algorithms is given in **Convergence and Mixing in Markov Chain Monte Carlo**.

Example

We wish to sample a distribution of the form

$$\pi(x) \propto \phi\left(\frac{x + 0.94}{5/\sqrt{14}}\right) \bar{\Phi}\left(\frac{5 - x}{5}\right)^6 \quad (21)$$

where $\phi(x) = \exp(-\frac{1}{2}x^2)$ is the standard normal density function and $\bar{\Phi}(x) = \int_x^{\infty} \phi(u) du$ is the normal survivor function. This distribution is the posterior distribution derived from observing a sample of 20 normally distributed lifetimes with mean x and standard deviation 5, $Z_i|x \sim N(x, 25)$, where 14 values less than 5, say $z_1, \dots, z_{14}$ are observed completely and have mean -0.94 and the remaining 6 values are truncated so that it is only known that they take values greater than 5 and a uniform prior distribution for X is assumed.

We shall consider three different MCMC samplers: an independence sampler $q(y|x) = N(0.85, 5)$, a Metropolis sampler $q(y|x) = N(x, 5)$ with initial values $x_0 = 0.85$ in both cases, and thirdly a Gibbs.

The Gibbs sampler is implemented by defining and introducing truncated normal latent variables $Z_i|x \sim TN(x, 5)$ such that $Z_i > 5$, for $i = 15, \dots, 20$ to represent the true values of the truncated normal data. Note that it is straightforward to sample from a truncated normal distribution *via* rejection sampling (see e.g. [25]; **Monte Carlo Methods**). Given the full sample $z_1, \dots, z_{14}, Z_{15} = z_{15}, Z_{20} = z_{20}$, the conditional density of X is normal

$$X|z_1, \dots, z_{14}, z_{15}, \dots, z_{20} \sim N\left(\frac{0.94 \times 14 + \sum_{i=15}^{20} z_i}{20}, \frac{5}{\sqrt{20}}\right) \quad (22)$$

This leads to the Gibbs sampling algorithm:

1. Fix $x_0 = 0.85$.
2. For $t = 1, 2, \dots$
 - (a) Generate $z_{ti} \sim TN(x_{t-1}, 5)$ for $i = 15, \dots, 20$.
 - (b) Calculate $\mu_t = (0.94 \times 14 + \sum_{i=15}^{20} z_{ti})/20$.
 - (c) Generate $x_t \sim N(\mu_t, 5/\sqrt{20})$.

Figure 2 gives the convergence of the means $\frac{1}{T} \sum_{i=1}^T x_i$ for each of the three algorithms. In this case, the mean for the Gibbs algorithm has converged after about 5000 iterations, whereas the Metropolis and independence samplers perform somewhat worse. In Figure 3 a histogram of the data generated from

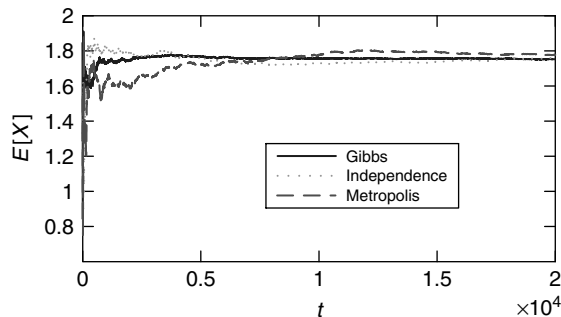


Figure 2 Convergence of the sample estimates of $E[X]$

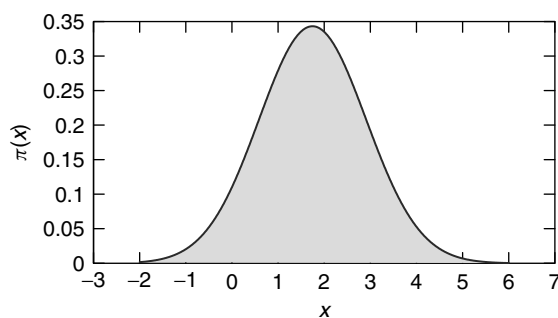


Figure 3 Histogram of the Gibbs-sampled data and fitted density estimate

the Gibbs algorithm along with an estimate of the density function of X is shown. The modal value is very close to 2 which was the true value of x used to generate these data.

References

- [1] Tierney, L. & Kadane, J. (1986). Accurate approximations for posterior moments and marginal densities, *Journal of the American Statistical Association* **81**, 82–86.
- [2] Metropolis, N. & Ulam, S. (1949). The Monte Carlo method, *Journal of the American Statistical Association* **44**, 335–341.
- [3] Norris, J. (1998). *Markov Chains*, University Press, Cambridge.
- [4] Tierney, L. (1996). Introduction to general state-space Markov Chain theory, in *Markov Chain Monte Carlo in Practice*, W. Gilks, S. Richardson & D.J. Spiegelhalter, eds, Chapman & Hall, London, pp. 59–74.
- [5] Metropolis, N., Rosenbluth, A., Rosenbluth, M., Teller, A. & Teller, E. (1953). Equations of state calculations by fast computing machines, *The Journal of Chemical Physics* **21**, 1087–1092.
- [6] Hastings, W. (1970). Monte Carlo sampling methods using Markov chains and their application, *Biometrika* **57**, 97–109.
- [7] Chib, S. & Greenberg, G. (1995). Understanding the Metropolis Hastings algorithm, *The American Statistician* **49**, 327–335.
- [8] Gelman, A. Gilks, W. R. & Roberts, G.O. (1997). Weak convergence and optimal scaling of random walk Metropolis algorithms, *Annals of Applied Probability* **7**, 110–120.
- [9] Tierney, L. (1994). Markov chains for exploring posterior distributions (with discussion), *Annals of Statistics* **22**, 1701–1762.
- [10] Geman, S. & Geman, D. (1984). Stochastic relaxation, Gibbs distributions and the Bayesian restoration of images, *IEEE Transactions on Pattern Analysis and Machine Intelligence* **6**, 721–741.
- [11] Gelfand, A. & Smith, A. (1990). Sampling based approaches to calculating marginal densities, *Journal of the American Statistical Association* **85**, 398–409.
- [12] Casella, G. & George, E. (1992). Explaining the Gibbs sampler, *The American Statistician* **46**, 167–174.
- [13] Gilks, W. & Wild, P. (1992). Adaptive rejection sampling for Gibbs sampling, *Applied Statistics* **41**, 337–348.
- [14] Gilks, W., Best, N. & Tan, K. (1995). Adaptive rejection metropolis sampling within Gibbs sampling, *Applied Statistics* **44**, 455–472.
- [15] Spiegelhalter, D., Thomas, A., Best, N. & Gilks, W. (1995). BUGS: Bayesian inference using Gibbs sampling, Technical Report, MRC Biostatistics Unit, University of Cambridge.
- [16] Best, N., Cowles, M. & Vines, K. (1995). CODA: convergence diagnosis and output analysis software for Gibbs sampling output, version 0.30, Technical Report, MRC Biostatistics Unit, University of Cambridge.
- [17] Neal, R. (2003). Slice sampling (with discussion), *Annals of Statistics* **31**, 705–767.
- [18] Green, P. (1995). Reversible jump MCMC computation and Bayesian model determination, *Biometrika* **82**, 711–732.
- [19] Propp, J. & Wilson, D. (1996). Exact sampling with coupled Markov chains and applications to statistical mechanics, *Random Structures and Algorithms* **9**, 223–252.
- [20] Robert, C. & Casella, G. (2004). *Monte Carlo Statistical Methods*, 2nd Edition, Springer, New York.
- [21] Brooks, S. & Roberts, G. (1998). Assessing convergence of Markov chain Monte Carlo algorithms, *Statistics and Computing* **8**, 319–335.
- [22] Gelman, A. & Rubin, D. (1992). Inference from iterative simulation using multiple sequences (with discussion), *Statistical Science* **7**, 457–511.
- [23] Mengersen, K., Robert, C. & Guihenneuc Jouyau, C. (1999). In *MCMC Convergence Diagnostics: A "Review"*, Bayesian Statistics, Vol. 6, J.M. Bernardo, J.O. Berger, A.P. Davis & A.F.M. Smith, eds, Oxford University Press, pp. 415–440.

- [24] Liu, J. & Chen, R. (1995). Blind deconvolution via sequential imputations, *Journal of the American Statistical Association* **90**, 567–576.
- [25] Geweke, J. (1991). Efficient simulation from the multivariate normal and student-t distributions subject to linear constraints, in *Computing Science and Statistics: Proceedings of the Twenty-Third Symposium on the Interface*, E.M. Keramidas, ed., Interface Foundation of North America, Fairfax, pp. 571–578.

(MCMC) for Bayesian System Reliability; Monte Carlo Methods; System Reliability: Monte Carlo Estimation.

MICHAEL P. WIPER

Related Articles

Convergence and Mixing in Markov Chain Monte Carlo; Hierarchical Markov Chain Monte Carlo

Laplace Approximations in Bayesian Lifetime Analysis

Introduction

The French mathematician and astronomer Pierre-Simon Marquis de Laplace (1749–1827) receives the most credit from statisticians for his seminal work on the central limit theorem, which in turn fueled the development of statistical asymptotic theory. Laplace of course was one the original “Bayesians”, as evidenced by his mathematical system of inductive reasoning and the resulting *rule of succession*. Perhaps not surprisingly, the roots of Laplace approximation can also be traced to a Bayesian problem, as reported in [1], where, in 1774, Laplace solved a quandary Thomas Bayes could not. He was able to construct the posterior distribution of a binomial probability arising from a uniform prior, which led to an approximation of the incomplete beta-function. Further modifications on this approximation method resulted in Laplace’s Fourier-type inversion formula for point probabilities associated with the generating function of a sum of independent discrete random variables [1]. It is noteworthy that this approximation method was also influential in Laplace’s development of the central limit theorem and was later generalized greatly to the behavior of sums of independent variables.

In the rest of this article, we will first briefly introduce the paradigm of Bayesian inference. Please see **Bayesian Reliability Analysis, Empirical Bayes Estimation of Reliability**, and the references therein for more detailed introduction and more comprehensive treatment of this vast area. We will then provide some details of the mechanics and the properties of Laplace approximation followed by a few examples of its use in **Bayesian reliability** literature, as a modest attempt to demonstrate the method’s utility in this area.

A Brief Introduction to the Bayesian Paradigm

Observed data are used to test scientific hypotheses, which are often stated in the form of probability distributions with unknown parameters. In

the perhaps more familiar frequentist (also known as *classical*) statistics approaches, these parameters are treated as fixed yet unknown constants to be estimated from the observed data. In contrast, the Bayesian framework treats these parameters as random variables and allows *a priori* beliefs about them to be expressed *via* what are called *prior distributions*. As relevant data becomes available, these priors are updated and the modified knowledge is expressed *via* what are called *posterior distributions*. These posterior distributions provide the basis for Bayesian inference. We will now briefly summarize the procedure involved in data-based updating procedure that leads to these posterior distributions mainly to set up the notation that will be used in the rest of this article.

As stated above, the Bayesian approach is based on the assumption that the unobserved parameter vector Θ is random and has a known distribution $\pi(\Theta)$, called the *prior distribution*. A Bayesian statistical model can be dissected into the following components [2]:

1. an observed random variable or random vector $\mathbf{X} \in \mathbf{X}$, where $\mathbf{X}$ is the data and $\mathbf{X}$ is the sample space;
2. an unobserved random variable or vector $\underline{\Theta} \in \Omega$, where Ω is the parameter space;
3. the conditional density $f(\mathbf{X}|\underline{\theta})$ of $\mathbf{X}$ given $\underline{\theta}$, also called the *likelihood*; and
4. the marginal density $\pi(\underline{\theta})$ of $\underline{\theta}$, also called the *prior*.

As stated above, most Bayesian inferential procedures depend on the so-called *posterior distribution* of $\underline{\Theta}$, which is defined as the conditional distribution of $\underline{\Theta}$ given $\mathbf{X}$. This posterior distribution can be obtained using the components provided above *via*

$$f(\underline{\theta}|\mathbf{X}) = \frac{f(\mathbf{X}|\underline{\theta})\pi(\underline{\theta})}{\int f(\mathbf{X}|\underline{\theta})\pi(\underline{\theta}) d\underline{\theta}} \quad (1)$$

This posterior distribution represents the updated knowledge about $\underline{\Theta}$ based on the observed data, $\mathbf{X}$.

Note that the functional form of the posterior is highly dependent on the form of the prior distribution, which represents our beliefs about the parameters before data collection. When little or no information is available about the parameters, it could be difficult to find a suitable prior. In such cases, what is called a *noninformative prior* can be used. Such priors

2 Laplace Approximations in Bayesian Lifetime Analysis

assign equal probability to each possible value of the parameter and they are often improper, i.e. they have infinite mass.

In principle, the posterior distributions can always be obtained; however, if the problem at hand is complex enough—and many realistic applications are—or when the problem formulation does not involve a conjugate prior-likelihood pair, then the necessary integrals may not be available in a closed form and thus, they may have to be approximated or numerically estimated. The Laplace's approximation is a desirable analytical method which may be employed for this purpose. It is often more accurate than the normal approximation, yet it is not as computationally intensive as purely numerical methods such as Markov chain Monte Carlo (*see* [3]; **Hierarchical Markov Chain Monte Carlo (MCMC) for Bayesian System Reliability; Markov Chain Monte Carlo, Introduction**).

Laplace's Approximation Method

Laplace's method attempts to approximate the density to be integrated by a **normal distribution** centered at the posterior mode. The variance of this normal distribution is given by the inverse of the second derivative of the log-posterior density evaluated at its mode. As long as the posterior distribution is unimodal, or it is at least dominated by a single mode, this approximation will produce reasonable results [4], provided the dimensionality of the parameter vector is not too high compared to the available sample size.

Laplace's approach to approximating integrals has wide applicability as it can be employed with all functions that can be expressed in the following form:

$$I = \int f(\underline{\theta}) \exp\{-nh(\underline{\theta})\} d\underline{\theta} \quad (2)$$

In equation (2), f is assumed to be a smooth positive function of $\underline{\theta}$, which is a p -dimensional parameter vector; and h is a smooth function of $\underline{\theta}$ where $-h$ has a maximum at $\hat{\underline{\theta}}$. The intuition behind the Laplace approximation is as follows [1]: Suppose that $\exp\{-nh(\underline{\theta})\}$ has a very strong peak at $\hat{\underline{\theta}}$ and decreases very rapidly as one moves away from $\hat{\underline{\theta}}$ on $\theta \in \Omega$ as $n \rightarrow \infty$. Then as $n \rightarrow \infty$, the major contribution to I can be attributed to the behavior of the integrand in a small neighborhood around $\hat{\underline{\theta}}$. The approximation of equation (2) is achieved by

expanding f and h using Taylor series about $\hat{\underline{\theta}}$. In the univariate case, the presentation given in [5], can be summarized as follows [6]:

Using a Taylor series expansion on f and h ,

$$\begin{aligned} I &\approx \int \left[f(\hat{\theta}) + \frac{f'(\hat{\theta})}{1!}(\theta - \hat{\theta}) \right. \\ &\quad \left. + \frac{f''(\hat{\theta})}{2!}(\theta - \hat{\theta})^2 \right] \times \exp \left\{ -nh(\hat{\theta}) \right. \\ &\quad \left. - \frac{nh'(\hat{\theta})}{1!}(\theta - \hat{\theta}) + \frac{nh''(\hat{\theta})}{2!}(\theta - \hat{\theta})^2 \right\} d\theta \\ &= e^{-nh(\hat{\theta})} \int \left[f(\hat{\theta}) + \frac{f'(\hat{\theta})}{1!}(\theta - \hat{\theta}) + \frac{f''(\hat{\theta})}{2!} \right. \\ &\quad \left. \times (\theta - \hat{\theta})^2 \right] \times \exp \left\{ -\frac{(\theta - \hat{\theta})^2}{2(nh''(\hat{\theta}))^{-1}} \right\} d\theta \quad (3) \end{aligned}$$

since $h'(\hat{\theta}) = 0$ from the definition of $\hat{\theta}$. Given that the exponential term inside the integral is proportional to an $N\left(\hat{\theta}, \frac{1}{nh''(\hat{\theta})}\right)$ density for θ , this implies that the $f'(\hat{\theta})(\theta - \hat{\theta})$ term will vanish since $E(\theta - \hat{\theta}) = 0$. Therefore the integral can be approximated as

$$\hat{I} \approx e^{-nh(\hat{\theta})} f(\hat{\theta}) \left(\frac{2\pi}{n} \right)^{1/2} \sigma \left(1 + \frac{f''(\hat{\theta})}{2f(\hat{\theta})} \sigma^2 \right) \quad (4)$$

where $\sigma = (h''(\hat{\theta}))^{-1/2} = O(\frac{1}{n})$.

For the cases where $\underline{\theta}$ is p dimensional, the approximate integral can be written as

$$\hat{I} \approx (2\pi)^{p/2} \{\det(nD^2h(\hat{\underline{\theta}}))\}^{-1/2} f(\hat{\underline{\theta}}) \exp\{-nh(\hat{\underline{\theta}})\} \quad (5)$$

where $\hat{\underline{\theta}}$ maximizes $-h(\underline{\theta})$ and where $D^2h(\hat{\underline{\theta}})$ is the Hessian matrix of h evaluated at $\hat{\underline{\theta}}$.

The seminal work in [4] provides a modification to this general approach, which is tailored for obtaining the marginal posterior densities necessary for Bayesian analysis. The authors note that their approach is similar to the saddle-point method [7] and proceeds as follows:

Consider the partition $\underline{\theta} = (\theta_1, \dots, \theta_p) = (\theta_1, \underline{\theta}_2)$, where θ_1 is the first component of $\underline{\theta}$ and $\underline{\theta}_2$ is a $(p-1) \times 1$ vector composed of the remaining components. Let $\hat{\underline{\theta}} = (\theta_1, \hat{\underline{\theta}}_2)$ represent the posterior

mode, i.e. $\hat{\theta}$ maximizes πe^ℓ , where π is the prior density and ℓ is the log-likelihood function. Now define Σ , a $p \times p$ matrix, as the negative inverse of the Hessian of $\frac{\ell + \ln \pi}{n}$ at $\hat{\theta}$. Note that for a given θ_1 , $\hat{\theta}_2^* = \hat{\theta}_2^*(\theta_1)$ is a $(p-1) \times 1$ vector and it maximizes $h(\cdot) = \pi(\theta_1, \cdot) e^{\ell(\theta_1, \cdot)}$. In other words, the function $h(\cdot)$ is equivalent to πe^ℓ with θ_1 held fixed. Similarly, we can define $\Sigma^* = \Sigma^*(\theta_1)$, where Σ^* is a $(p-1) \times (p-1)$ matrix and represents the negative inverse of the Hessian of $\frac{\ln h(\cdot)}{n}$. The following expression can be used to obtain the marginal posterior of θ_1

$$\pi_1(\theta_1 | \mathbf{X}^{(n)}) = \frac{\int \pi(\theta_1, \theta_2) e^{\ell(\theta_1, \theta_2)} d\theta_2}{\int \pi(\theta) e^{\ell(\theta)} d\theta} \quad (6)$$

and applying Laplace's method to its numerator and denominator leads to the following approximation:

$$\hat{\pi}_1(\theta_1 | \mathbf{X}^{(n)}) = \left(\frac{\det \Sigma^*(\theta_1)}{2\pi n \det \Sigma} \right) \frac{\pi(\theta_1, \hat{\theta}_2^*) e^{\ell(\theta_1, \hat{\theta}_2^*)} d\theta_2}{\pi(\hat{\theta}) e^{\ell(\hat{\theta})} d\theta} \quad (7)$$

It has been shown [4] that the error associated with this approximation is

$$\pi_1(\theta_1 | \mathbf{X}^{(n)}) = \hat{\pi}_1(\theta_1 | \mathbf{X}^{(n)}) \left(1 + O_{\theta_1} \left(\frac{1}{n} \right) \right) \quad (8)$$

where $O_{\theta_1}(\frac{1}{n})$ is of order $O(\frac{1}{n})$ and depends on θ_1 . Given the true parameter vector $\underline{\theta}_0$, under certain regularity conditions [8] it can be shown that for θ_1 in some fixed neighborhood of $\theta_{0,1}$, the first component of $\underline{\theta}_0$, the error term $O_{\theta_1}(\frac{1}{n})$ is uniformly of order $O(\frac{1}{n})$. In contrast, the absolute error of the usual normal approximation is only of order $O(\frac{1}{\sqrt{n}})$. Recall that the relative error of all Edgeworth-type approximations, which include the normal approximation, usually tends to zero at a rate of $1/\sqrt{n}$ on neighborhoods that shrink toward $\theta_{0,1}$ [4]. The error in equation (8) is larger mainly due to the fact that the dimensions of the two integrals in the numerator and the denominator are different. The error in the functional form of $\pi_1(\theta_1 | \mathbf{X}^{(n)})$ is in fact on the order of $O(\frac{1}{\sqrt{n}})$ in $1/\sqrt{n}$ neighborhoods of $\hat{\theta}_1$. It is the constant of integration which is associated with the larger error.

As an illustration, consider a case where a new product is tested until it fails. Suppose 120 items

are put on test and the variable of interest is posterior odds, i.e. the ratio of the number that failed over the number that did not. In this scenario, let Y , the number of items that fail, have a binomial distribution, i.e. $Y \sim \text{Binomial}(120, \theta)$. Suppose prior experience indicates that a Beta(1, 5) is an appropriate prior for θ . Then the marginal distribution of the posterior odds, $\lambda = \frac{\theta}{1-\theta}$ can be easily obtained using the Laplace approximation. In this case, the exact distribution is also analytically tractable and can be obtained by simply generating realizations from the posterior distribution of θ and calculating the value of λ for each realization. The following plot displays the exact distribution as well as the marginal posterior obtained *via* the Laplace approximation. As expected the agreement between the two curves is quite good (Figure 1).

Some Example Applications from Bayesian Reliability Literature

In this section, we provide a few application examples from the reliability literature where Laplace approximation was utilized. The examples below are not meant to be comprehensive in any way and they were chosen without any particular criteria other than their familiarity to the author.

Laplace approximations have been used in Bayesian residual analysis of regression models for lifetime data with and without censoring [9]. In this work, posterior probabilities of observations being **outliers** have been obtained utilizing an approach

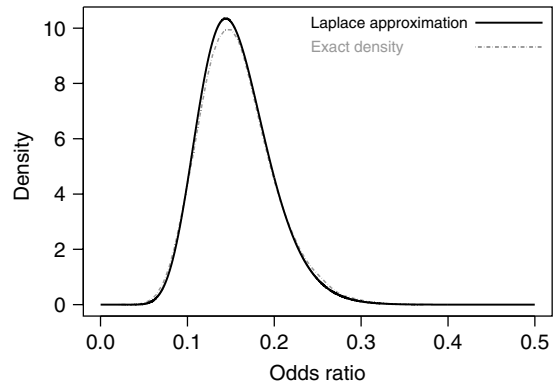


Figure 1 The exact distribution and the corresponding laplace approximation of the marginal distribution of posterior odds ratio

similar to the one that led to Cox–Snell residuals [10]. The posterior distributions of the realized errors, which represent the uncertainty about the functions of the parameters, are obtained *via* Laplace approximations, which are then used to estimate realized errors *via* intervals.

Laplace approximations have also been employed in estimating the survival function of the Pareto distribution of the second kind in the context of failure-censored data [11]. In this work, the authors compared their approach to an alternative technique proposed by Lindley [12] and to the maximum-likelihood method. They concluded that for small samples, the approach in [12] was slightly superior, whereas for large samples, all three methods gave comparable results.

Laplace approximations have been popular especially in the context of Bayesian accelerated testing applications (*see Accelerated Life Models; Accelerated Life Tests; Bayesian Design; Accelerated Life Tests: Designs Comparison within a Bayesian Framework*). The link functions used in this setting to represent the dependence of the survival distribution parameters on the acceleration variables often lead to intractable integrals for the posteriors, providing a fertile ground for the use of approximation or numerical methods. For example, Laplace’s method along with Jeffreys’ and reverse-reference priors was used in constructing the marginal posterior distributions of the acceleration and the model parameters associated with the accelerated inverse Gaussian distribution in a cumulative damage setting [13]. Similarly the approach was employed in modeling accelerated life test data under a stress–response relationship with a threshold stress, where the lifetimes are assumed to be exponential [14].

Discussion

Having introduced the Laplace approximation method, it may be useful to summarize some of its advantages and disadvantages [6]. Laplace’s method has been desirable partially due to the fact that it is not very complex and is relatively fast computationally, since it does not require an iterative algorithm. Further, since the approach requires numerical differentiation, it is often easier and numerically more stable than techniques that require numerical integration. Additionally Laplace’s method

does not rely on random numbers, i.e. it is deterministic, and hence the consistency of the results can be checked by other users provided the same prior, the same model, and the same data set are used.

Laplace’s approximation also has several limitations. As mentioned previously, in order to achieve a valid approximation, the approach requires a posterior distribution which is at least dominated by a single mode. Further, the accuracy of the approximation depends on the parameterization used (e.g. θ vs $\log \theta$), which could be problematic since choosing the correct one may be difficult. On this issue, the use of a variance-stabilizing transformation is recommended in order to enhance the performance of the Laplace approximation method [15–17]. Another possible limitation of the Laplace method is that the size of the data set used must be fairly “large” – where “large” is likely to be problem specific. One of the major drawbacks of the method however, is the fact that the associated asymptotics are in n , which implies that in order to improve the accuracy of the approximation, one needs to collect more data. Finally, if $\underline{\theta}$ is moderate-to-high dimensional, it will not only be very difficult to perform the required computations associated with the Hessian matrices, but the accuracy of the approximation provided by Laplace’s method will often be insufficient as well. It has been suggested that the Laplace approximation is reliable when the dimension of the parameter vector, $p = O(n^{1/3})$, and some modifications have been proposed [18] when the intention is to use the method to approximate high-dimensional integrals.

Since the Tierney and Kadane [4] paper, which largely enabled the use of Laplace approximation in both Bayesian and non-Bayesian applications, other attractive and largely computational alternatives, most notably Markov chain Monte Carlo, have emerged and have been made accessible to a large audience *via* user friendly software such as WinBugs (<http://www.mrc-bsu.cam.ac.uk/bugs/>). Despite the availability of such alternatives, Laplace’s approximation remains a viable option especially for problems that are not excessively complex or very high dimensional.

References

- [1] Strawderman, R.L. (2000). Higher order asymptotic approximation: Laplace, saddlepoint, and related methods, *Journal of the American Statistical Association* **95**, 1358–1364.

- [2] Arnold, S.F. (1990). *Mathematical Statistics*, Prentice-Hall, Englewood Cliffs.
- [3] Gilks, W.R., Richardson, S. & Spiegelhalter, D.J. (1996). *Markov Chain Monte Carlo*, Chapman-Hall, London.
- [4] Tierney, L. & Kadane, J.B. (1986). Accurate approximations for posterior moments and marginal densities, *Journal of the American Statistical Association* **81**, 82–86.
- [5] Kass, R.E., Tierney, L. & Kadane, J.B. (1990). The validity of posterior expansions based on Laplace's method, in *Essays in Honor of George A. Barnard*, S. Geisser, J.S. Hodges, J.S. Press & A. Zellner, eds, North-Holland, Amsterdam.
- [6] Carlin, B.P. & Luis, T.A. (1996). *Bayes and Empirical Bayes Methods for Data Analysis*, Chapman & Hall, New York.
- [7] Daniels, H.E. (1956). Saddle point approximations in statistics, *Annals of Mathematical Statistics* **25**, 631–650.
- [8] Walker, A.M. (1969). On the asymptotic behavior of posterior distributions, *Journal of the Royal Statistical Society. Series B* **31**, 80–88.
- [9] Chalonder, K. (1991). Bayesian residual analysis in the presence of censoring, *Biometrika* **78**, 637–644.
- [10] Cox, D.R. & Snell, E.J. (1968). A general definition of residuals (with discussion), *Journal of the Royal Statistical Society. Series B* **30**, 248–275.
- [11] Howlader, H.A. & Hossain, A.M. (2002). Bayesian survival estimation of Pareto distribution of the second kind based on failure censored data, *Computational Statistics and Data Analysis* **38**, 301–314.
- [12] Lindley, D.V. (1980). Approximate Bayesian methods (with discussion), *Trabajos de Estadística e Investigación Operativa* **31**, 232–245.
- [13] Onar, A. & Padgett, W.J. (2001). Some Bayesian procedures with accelerated inverse Gaussian distribution, *Journal of Applied Statistical Science* **10**, 192–209.
- [14] Tojeiro, C.A.V., Louzada-Neto, F. & Bolfarine, H. (2004). A Bayesian analysis for accelerated lifetime tests under an exponential power law model with threshold stress, *Journal of Applied Statistics* **31**, 685–691.
- [15] Slate, E.H. (1994). Parameterizations for natural exponential families with quadratic variance functions, *Journal of the American Statistical Association* **89**, 1471–1482.
- [16] Achcar, J.A. (1989). Reparameterizations and Laplace approximations for posterior moments in binomial models, *Revista de Matematica Estatística* **7**, 73–86.
- [17] Achcar, J.A. & Smith, A.F.M. (1990). Aspects of reparameterization in approximate Bayesian inference, in *Essays in Honor of George A. Barnard*, S. Geisser, J.S. Hodges, J.S. Press & A. Zellner, eds, North-Holland, Amsterdam.
- [18] Shun, Z. & McCullah, P. (1995). Laplace approximation of high dimensional integrals, *Journal of the American Statistical Association* **57**, 749–760.

Related Articles

Accelerated Life Models; Accelerated Life Tests; Bayesian Design; Accelerated Life Tests: Designs Comparison within a Bayesian Framework; Bayesian Reliability Analysis; Empirical Bayes Estimation of Reliability; Hierarchical Markov Chain Monte Carlo (MCMC) for Bayesian System Reliability; Markov Chain Monte Carlo, Introduction.

ARZU ONAR

Neural Networks in Statistical Process Control

Introduction

The resurgence of research interests in neural computing came about in the early 1980s. Enormous interest exists in studying and applying neural computing in almost all fields of science and technology (engineering, business, and medical science, to name a few). The application of **neural networks** in industrial, production engineering or related domains started gaining popularity in the latter 1980s. The momentum of the growing research interest peaked around the mid-1990s and is still sustaining a high interest level as of 2006. The tremendous interest and wide applicability of neural computing theories and technology motivate researchers to explore how this class of modeling tools can be effectively applied to statistical process control (SPC).

The fundamental issue at the core of classical statistical quality control has remained relatively unchanged since the introduction of **control charts** by Walter Shewhart in the 1920s (*see Statistical Quality Control*). The basic principle for a control scheme is to signal an out-of-control state when the behavior of critical quality variables in a process is considered statistically abnormal, but also to leave the process intact when no evidence of statistical abnormality is present. However, as a result of modern computer and data acquisition technology, the manifestation of an out-of-control state in a process may take various forms, e.g., characteristically patterned or highly autocorrelated. Under the circumstances, the data collected often render the traditional statistical control charts ineffective or inapplicable. From the process improvement point of view, practitioners desire to have the knowledge about a faulty process beyond just “a violation of a certain set of control limits, e.g., 3σ ” (*see Control Charts and Process Capability; Quality Management, Overview; Six Sigma*). The aforementioned concerns have been recognized and the need to identify nonrandom or unnatural patterns was well documented long ago, e.g., in the widely referenced *Statistical Quality Control Handbook* [1]. They also have drawn keen attention from the statistical quality control research community in the last couple

of decades (*see Sampling Inspection of Products and Statistical Process Control; Control Charts, Overview; Control Charts for the Mean; Exponentially Weighted Moving Average (EWMA) Control Chart; Cumulative Sum (CUSUM) Chart; Multivariate Control Charts Overview; Hotelling's T^2 Chart; Multivariate Exponentially Weighted Moving Average (MEWMA) Control Chart; Multivariate Control Charts, Interpretation of; Control Charts for Short Production Runs; Control Charts for Batch Processes; Control Charts, Selection of; Statistical Process Control in Clinical Medicine*).

Two primary categories of problems are the detection of shifts and the recognition of unnatural patterns from the process monitored. The former is concerned with process mean shift, which may occur when the process variable is independently and identically distributed (i.i.d.), or under some underlying structure such as **autocorrelation** (*see Autocorrelated Data*). The latter focuses on patterns exhibited within the typical control limits, which cannot be recognized using conventional decision criteria. The neural network-based schemes are flexible and adaptable enough to deal with both categories of problems. The expanded capabilities to distinguish between mean shift and structural change or patterns are of great value for the purpose of process control. The article focuses on various models appropriate in the SPC context and the general modeling approach employing various neural network algorithms.

Process Models

Modeling Unnatural Patterns and Shifts

The realization of a process variable over time, X_t , can be considered comprising of the process mean and two noise components, i.e., one from the random cause and the other is a special disturbance from assignable causes, as expressed in equation (1).

$$X_t = \mu + x_t + d_t \quad (1)$$

where

- X_t : quality measurement at time t ,
- μ : process mean when the process is in control,
- σ : process standard deviation when the process is in control,

2 Neural Networks in Statistical Process Control

- x_t : random normal variate at time t , $x_t \sim N(0, k\sigma)$, $0 < k \leq 1$,
 d_t : special disturbance at time t .

Note that d_t may take in various forms to model unnatural patterns such as trends, cycles, systematic variables, mixtures, stratification, and sudden shift in level, and so on [2]. A shift in process mean can be expressed as a type of unnatural pattern.

Modeling Shifts under Autocorrelation

Let Z_t represent the generic form of an autoregressive and moving average (ARMA) model (see **Large Autoregressive Integrated Moving Average (ARIMA) Modeling**). Suppose the process under consideration is one that has an additive, deterministic step shift of magnitude, δ , at some point of time, ψ . Then the process can be represented by equation (2).

$$X_t = Z_t + \delta h_t \quad (2)$$

where

$$h_t = \begin{cases} 0 & \text{if } t < \psi \\ 1 & \text{if } t \geq \psi \end{cases} \quad (3)$$

Modeling Shifts in Multivariate Processes

Suppose there are p process variables which are jointly distributed following a p -variate **normal distribution** with known mean vector μ , and known **covariance** matrix Σ . Equation (4) represents a multivariate process.

$$X_t = \mu + x_t + \Delta_t \quad (4)$$

where

- X_t : a vector of p quality measurements at time t ,
 μ : a vector of process means when the process is in control,
 x_t : a vector of random normal variate at time t , x_t follows a p -variate normal distribution with mean vector of zeros and known covariance matrix Σ ,
 Δ_t : step shift vector at time t .

Modeling Shifts in Multivariate Autocorrelated Processes

A multivariate autocorrelated process of p variables can be expressed as a vector autoregressive

model. Equation (5) represents a Vector Autoregressive model of order p , or VAR(p).

$$Y_t - \mu_t = \Phi_1(Y_{t-1} - \mu_{t-1}) + \Phi_2(Y_{t-2} - \mu_{t-2}) + \cdots + \Phi_p(Y_{t-p} - \mu_{t-p}) + \varepsilon_t \quad (5)$$

where μ_t is the vector of mean values at time t , ε_t is an independent multivariate normal random vector with the mean vector of zeros and covariance matrix Σ , and Φ_i ($i = 1, 2, \dots, p$) is a matrix of autocorrelation parameters.

Let $Z_t = Y_t - \mu_t$. Equation (6) describes a process that has an additive, deterministic step shift vector, Δ , beginning at some point of time, ψ .

$$X_t = Z_t + \Delta h_t \quad (6)$$

where

$$h_t = \begin{cases} 0 & \text{if } t < \psi \\ I & \text{if } t \geq \psi \end{cases} \quad (7)$$

The above models represent some of the widely characterized processes. Other forms of process models are possible.

Neural Network Modeling

Algorithms

Some of the neural network algorithms that have been adopted include two-layer, multilayer perceptrons with the widely used and proven learning algorithm backpropagation [2, 3], radial basis function (RBF) [4], Boltzmann machines [5, 6], adaptive resonance theory (ART) [7, 8] and its variants ARTMAP and fuzzy ARTMAP [9], and learning vector quantization (LVQ) network [10]. The data types handled can be coded binary or raw analog data series. The network output is in the form of pattern class, predicted magnitude of the shift, or value of the observation in the series. The network training regimes include supervised, e.g., backpropagation, and unsupervised, e.g., ART. Supervised training requires the desired or target output to be presented during the learning, while unsupervised training does not. In terms of the functionality of the system developed, it can be a general-purpose or special-purpose system.

Beyond doubt, the backpropagation algorithm is the most widely adopted in building neural networks

for a wide range of applications including both pattern recognition and shift detection. The standard backpropagation uses a generalized Delta rule [3] that updates network connection weights according to the gradient descent method using equation (8).

$$\Delta w_{ij} = \eta \delta_i o_j \quad (8)$$

where o_j is the actual output of the j th unit, η is the learning constant and δ_i is defined as follows:

$$\delta_i = (t_i - o_i) f'(net_i) \quad (9)$$

where $f'(net_i)$ is the derivative of the activation function f for the i th unit evaluated at the net total input net_i and t_i is the desired output of the i th unit. Two of the most widely used activation functions are the sigmoid (equation 10) and the hyperbolic tangent (equation 11).

$$f(z) = (1 + e^{-z})^{-1} \quad (10)$$

$$f(z) = \frac{e^z - e^{-z}}{e^z + e^{-z}} \quad (11)$$

For more discussions about neural networks, *see* [11, 12] and **Neural Networks: Construction and Evaluation**. For a review of applying neural networks to SPC, *see* [13].

The Modeling Process

The modeling process typically consists of two stages, namely, network development and performance evaluation. In the network development stage, one has to determine the following key parameters: the window size of input data, the number of hidden layers and hidden nodes, the number of output nodes, the size of the training file, the size of the testing file, the learning rule, the **transfer function**, the training termination or convergence criterion, and so on. These decisions have direct impact on the capabilities of the trained network. Depending on detecting shifts or recognizing patterns, the requirement of the window size varies. The window size should be as small as possible, say five or lesser, and yet sufficient to capture the nature of the data series. When the underlying structure of the process data is more complex, such as autocorrelation, or the patterns of interest are more complicated, such as concurrent presence of multiple patterns or features, the window size has to be sufficiently large, say 50 or more. *See* [14–16]

for some guiding principles, e.g., avoid ambiguity, adequate representation, and balanced representation, in designing the training data sets. The goal at the end of this stage is to develop a network that has not only learned all the training examples well but can also be applied to new, unseen data series with reasonable generalization capabilities. Root mean squared (RMS) errors are frequently used as the measure to determine the stability of learning or if a satisfactory level of training has been reached.

Performance Evaluation

To evaluate the capabilities of the trained network, one needs to measure the probabilities of type I and type II errors. In the case of shift detection, the primary concern is how quickly the trained network can detect the shift when it actually occurs without generating too many false alarms. Similar to the conventional SPC measures, average run length (ARL) is a useful performance measure (*see* **Average Run Lengths and Operating Characteristic Curves**). Given a prespecified in-control ARL, the trained network is “tuned” to a corresponding level of discriminating **power**. For instance, this level will be determined by choosing an appropriate output activation threshold value in a backpropagation neural network.

In the case of pattern recognition, one is concerned with not only the correct classification rates but also the timeliness of the recognition. Run length measures similar to the conventional ARL should be considered. However, depending on the number of pattern classes to be recognized, the possible outcomes for each classification are not dichotomous. Therefore, the run length must be based only on the correct classifications. Target-pattern-based measures and ARL indices, which combine both the correct classification rate and the ARL, should be used [2].

An Illustrated Example

Consider a situation whereby two vendors order the same type of goods from the same supplier. The order quantities placed by each vendor form a **time series** with autocorrelation. Owing to the common demand population and the common source of supply, the two vendors’ order quantities are also cross-correlated, thus, the order quantities follow a VAR(1)

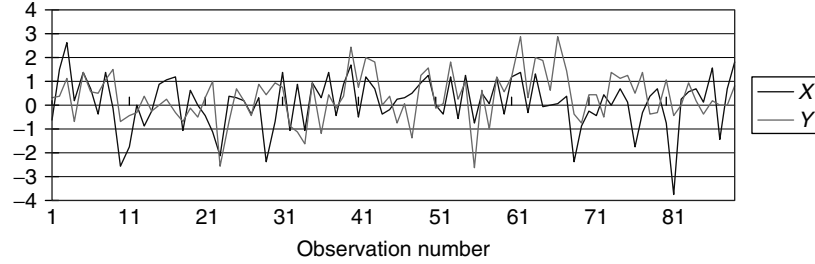


Figure 1 VAR(1): The standardized order quantities

model. At some point in time, one of the vendors embarked on a promotional drive to stimulate the demand. Over a 3-month period, order quantities were collected and shown as in Figure 1. The data are preprocessed as moving windows of data (see Figure 2) and presented to a backpropagation neural network previously trained for similar data. The network has a configuration of 100 input nodes and 5 output nodes, with no hidden layers. The 100 input nodes receive 50 pairs of (X, Y) observations. The five output nodes represent shift in X , shift in Y , autocorrelation in X , autocorrelation in Y , and cross-correlation between X and Y , respectively.

Figure 3 presents the neural network output chart for the first two output nodes, i.e., shift in X and shift in Y . The chart clearly indicates the progression of the two time series. The upward shift in order quantities from vendor Y can be clearly observed. In fact, the mean increases by an amount equivalent to 0.5 times the standard deviation of the existing ordering process from observation no. 60 onward. With proper tuning of in-control ARL, the cutoff threshold is set at 0.190768. The neural network-based control scheme detected the shift and identified the source at observation no. 65 (or in window no. 16).

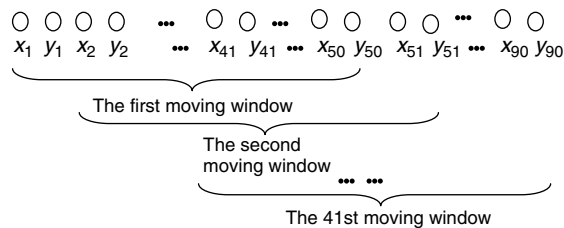


Figure 2 Moving windows for neural network input

Strengths

Being flexible, adaptable, and capable of learning are the three major strengths. A neural network can be trained for a wide range of data. Unlike traditional statistical approaches, to apply a trained network does not require assumptions about the data distribution, the time series model, or the prediction errors of the new data. The approach is relieved of the requirement of the validity of typical statistical assumptions. Neural networks are more flexible and can be applied to a wider range of data.

The capabilities of a network are closely dependent on the data upon which it was trained. By carefully designing and selecting the training data set, a network can be trained to adapt to any particular data set or a set of features. For example, one may adapt a network to focus more on small to moderate shifts by providing more training data of such features. Neural networks are more adaptable and ideal for developing special-purpose systems.

By presenting the training data to a neural network, the network learns to capture the features or form clusters from the data. In addition to learning from the given desired outputs, so-called supervised learning, neural networks can also learn to form unknown clusters or pattern classes. The progression of learning can be in stages at different times and cumulative. In other words, what is learned in the past can become part of the memory and the basis upon which new learning may take place.

Weaknesses

A lack of analytical robustness and the need for intensive data are the two major weaknesses. The production of the data sets used in training and/or testing

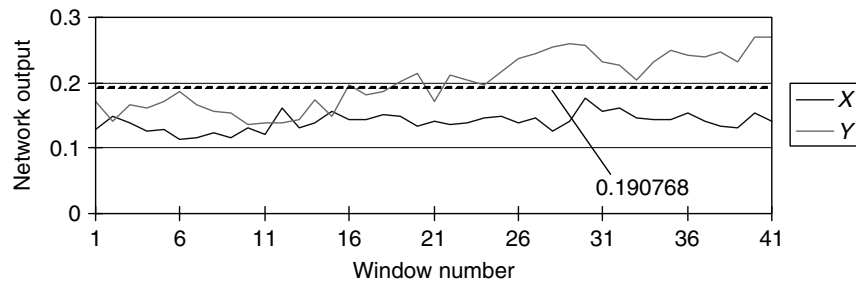


Figure 3 The neural network output chart

involves some degree of randomness. The process of iterative training itself also involves randomness. Despite the efforts in producing pseudorandom data and in training for stability, the resultant trained network is not an outcome of some analytical modeling. As a result, the outcome is often data-dependent or training-dependent making it difficult to offer a robust analysis or explanation of the results. In terms of the amount of data required, the neural network approach suffers from the scalability problem, e.g., when dealing with processes of high dimension.

Summary

The pattern classification capability of neural networks enables one to develop pattern recognizers while the adaptive learning capability allows networks to be “equipped” for different types of patterns or different purposes for process control. As the nature of modern industry or business transactions evolves and the mode of data acquisition advances, the need for dealing with complex and enormous amounts of data becomes more pressing. As demonstrated in this article, it is promising to explore further the utility of neural networks in dealing with data collected from autocorrelated processes in a multivariate setting. The advantages of the neural network approach over the traditional statistical approaches in such applications will depend highly on the ability to interpret the neural network outputs and the diagnosis of the assignable causes associated with each of the variables. These two domains must be integrated to develop a highly automated and intelligent neural network-based SPC system. Much work remains to be done.

References

- [1] Western Electric Company. (1956). *Statistical Quality Control Handbook*, Western Electric Company, New York.
- [2] Hwang, H. & Hubele, N. (1993). Back-propagation pattern recognizers for $\bar{X}$ control charts: methodology and performance, *Computers and Industrial Engineering* **24**, 219–235.
- [3] Rumelhart, D., Hinton, G. & Williams, R. (1986). Learning internal representation by error propagation, in *Parallel Distributed Processing: Explorations in the Microstructure of Cognition, Vol. 1: Foundations*, D. Rumelhart & J. McClelland, eds, MIT Press, Cambridge, pp. 318–362.
- [4] Moody, J. & Darken, C. (1989). Fast learning in networks of locally-tuned processing units, *Neural Computing* **1**, 281–294.
- [5] Hinton, G. & Sejnowski, T. (1986). Learning and relearning in Boltzmann machines, in *Parallel Distributed Processing: Explorations in the Microstructure of Cognition, Vol. 1: Foundations*, D. Rumelhart & J. McClelland, eds, MIT Press, Cambridge, pp. 282–317.
- [6] Hwang, H. & Hubele, N. (1992). Boltzmann machines that learn to recognize patterns on control charts, *Statistics and Computing* **2**, 191–202.
- [7] Carpenter, G. & Grossberg, S. (1987). A massively parallel architecture for a self-organizing neural pattern recognition machine, *Computer Vision, Graphics, and Image Processing* **37**, 54–115.
- [8] Hwang, H. & Chong, C. (1995). Detecting process non-randomness through a fast and cumulative learning ART-based pattern recognizer, *International Journal of Production Research* **33**, 1817–1833.
- [9] Hwang, H. (1997). A neural network approach to identifying cyclic behavior on control charts: a comparative study, *International Journal of Systems Science* **28**, 99–112.
- [10] Kohonen, T. (1989). *Self-Organization and Associative Memory*, Springer-Verlag, Berlin.

6 Neural Networks in Statistical Process Control

- [11] Fu, L. (1994). *Neural Networks in Computer Intelligence*, McGraw-Hill, New York.
- [12] Hertz, J., Krogh, A. & Palmer, R. (1991). *Introduction to the Theory of Neural Computation*, Addison-Wesley, Redwood.
- [13] Zorriassatine, F. & Tannock, J. (1998). A review of neural networks for statistical process control, *Journal of Intelligent Manufacturing* **9**, 209–224.
- [14] Hwarng, H. & Hubele, N. (1993). $\bar{X}$ control chart pattern identification through efficient off-line neural network training, *IIE Transactions* **25**, 27–40.
- [15] Hwarng, H. (1995). Proper and effective training of a pattern recognizer for cyclic data, *IIE Transactions* **27**, 746–756.
- [16] Hwarng, H. (2004). Detecting process mean shift in the presence of autocorrelation: a neural network based monitoring scheme, *International Journal of Production Research* **42**, 573–595.

H. BRIAN HWARNG

Bayesian Robustness: Theory and Computation

Introduction

Robustness is an important issue in Bayesian analysis, being concerned with the sensitivity of the analysis itself to the assumptions of the model adopted, in particular, to the prior distribution which is usually difficult to assess in practice.

The state of the art up to the year 2000 is reviewed in [1]; in particular, in [2] an overview of the robust Bayesian approach is presented; this latter is usually subdivided into the global robustness approach, in which the class of all priors compatible with the elicited prior information is considered, and the local robustness approach, where the interest is in the rate of change in inferences with respect to changes in the prior.

Most contributions to the subject are of a theoretical nature and not much work has been done about effective algorithms for robustness. As a result, this article is quite technical mathematically, reflecting the nature of the work done on the subject up to now. Developing effective and practical algorithms for Bayesian robustness that work in general situations has proved difficult to do.

In the case of global robustness with respect to the prior (in short robustness in the sequel), an important link with the possibility of providing effective algorithms is available when the class of priors is defined by generalized moment conditions. Indeed in this case the problem of analyzing Bayesian robustness reduces to a problem of linear semi-infinite programming (LSIP).

Generalized moment classes have been considered in connection with robustness first in [3] and then in several other papers (see [4] and its references). Such classes, whose features are recalled in the section titled “Generalized Moment Classes and Bayesian Robustness”, incorporate a number of interesting situations – the most common being the one in which bounds on **quantiles** of the prior distribution are available – and have been widely studied in other contexts, so that a rather comprehensive theory exists for optimization of linear functions defined over them, originated by [5–7], which we review in the section titled “Application

of the Generalized Moment Theory to Bayesian Robustness”. In robustness analysis, the function to be optimized is not linear but, as it is shown in the section titled “Generalized Moment Classes and Bayesian Robustness”, it is possible to obtain linearity by a suitable transformation, so that the above theory can be applied to robustness analysis as well.

Once the problem is reduced to an LSIP one, its numerical solution can be accomplished by algorithms developed in this latter context. A review of such algorithms is presented in the section titled “LSIP Algorithms and Bayesian Robustness”. In the section titled “An Example” the application to robustness problems of the so-called accelerated central cutting plane (ACCP) algorithm, which has the attractive feature of converging under mild conditions on the problem constraint functions, is illustrated by an example. A C++ implementation of the ACCP algorithm and its interface with robustness analysis has been developed by the author and is available on request.

Generalized Moment Classes and Bayesian Robustness

Let X be a random variable on a dominated statistical space $(\mathcal{X}, \mathcal{F}_X, \{P_\theta, \theta \in (\Theta, \mathcal{F})\})$, with density $f(x | \theta)$ with respect to the dominant measure λ , where the parameter space $(\Theta, \mathcal{F})$ is such that Θ is a subset of a finite-dimensional euclidean space; denote by $l_x(\theta)$ (or simply by $l(\theta)$) the likelihood function $f(x | \theta)$, which we assume is $\mathcal{F}_X \otimes \mathcal{F}$ measurable. Let $\mathcal{P}$ be the space of all probability measures on $(\Theta, \mathcal{F})$ and define

$$\Gamma = \{\pi \in \mathcal{P} : \int_{\Theta} H_i(\theta) \pi(d\theta) \leq \alpha_i, i = 1, \dots, m\} \quad (1)$$

where H_i are integrable with respect to any $\pi \in \mathcal{P}$ and α_i are fixed real constants, $i = 1, \dots, m$. Suppose that Γ is nonempty. We are interested in determining

$$\sup_{\pi \in \Gamma} \frac{\int_{\Theta} g(\theta) l(\theta) \pi(d\theta)}{\int_{\Theta} l(\theta) \pi(d\theta)} \quad (2)$$

where $g : \Theta \rightarrow \mathbb{R}$ is a given function such that $\int g(\theta) \pi(d\theta | x)$ exists for all $\pi \in \Gamma$, $\pi(d\theta | x)$ being

2 Bayesian Robustness: Theory and Computation

the posterior distribution of θ corresponding to the prior π .

The function to be optimized in expression (2) is not linear in π , however it is possible to build up a problem equivalent to the one given by expression (2) in which the objective function is linear. Consider the following map from $\mathcal{P}$ into $\mathcal{M}$, the set of finite measures on Θ : $\nu(A) = \int_A \pi(d\theta) / \int_\Theta l(\theta)\pi(d\theta)$ $A \in \mathcal{F}$, assuming that $0 < \int_\Theta l(\theta)\pi(d\theta) < +\infty$. Accordingly, expression (2) turns into

$$\begin{aligned} \sup_{\nu \in \mathcal{M}_1} \int_\Theta g(\theta)l(\theta)\nu(d\theta) \\ \mathcal{M}_1 = \mathcal{M}_1(\Theta) = \left\{ \nu \in \mathcal{M} : \int_\Theta f_i(\theta)\nu(d\theta) \right. \\ \left. \leq 0, i = 1, \dots, m \right. \\ \left. \int_\Theta l(\theta)\nu(d\theta) = 1 \right\} \\ f_i(\theta) := H_i(\theta) - \alpha_i, i = 1, \dots, m \end{aligned} \quad (3)$$

From now on the transformed problem formulation (3) is considered. For an extension to the non-parametric case see [4].

Application of the Generalized Moment Theory to Bayesian Robustness

Let Θ be some closed subset of $\mathbb{R}^n$ and let $g, l, f_i, i = 1, \dots, m$ in expression (3) be such that the following assumptions hold:

- A1. l is continuous and for $i = 1 \dots m$ f_i is lower semicontinuous;
- A2. g is upper semicontinuous;
- A3. $\forall i = 1 \dots m \exists C_i \in \mathbb{R}$ such that $f_i \geq C_i$;
- A4. l is bounded from above;
- A5. gl is bounded from above;
- A6. $\exists \tilde{\beta}_0 \in \mathbb{R}, \tilde{\beta}_1, \dots, \tilde{\beta}_{m+1} \geq 0 : \forall \theta \in \Theta$ it is $\tilde{\beta}_0 l(\theta) + \sum_{i=1}^m \tilde{\beta}_i f_i(\theta) - \tilde{\beta}_{m+1} g(\theta)l(\theta) > 1$;
- A7. $\forall \epsilon > 0 \exists K_\epsilon \subseteq \Theta, K_\epsilon$ compact, $\exists \beta_0^\epsilon \in \mathbb{R}, \beta_i^\epsilon \geq 0, i = 1, \dots, m+1$ with $|\beta_0^\epsilon| + \sum_{i=1}^{m+1} \beta_i^\epsilon \leq B, B > 0$ such that, for any $s \in K_\epsilon^c$, it is $\beta_0^\epsilon l(\theta) + \sum_{i=1}^m \beta_i^\epsilon f_i(\theta) - \beta_{m+1}^\epsilon g(\theta)l(\theta) > 1/\epsilon$.

According to [4], Theorem 5, the following result holds:

Theorem 1 Under conditions A1–A7, it is

$$\begin{aligned} \sup_{\nu \in \mathcal{M}_1} \int_\Theta g(\theta)l(\theta) d\nu(\theta) \\ = \inf \left\{ \beta_0 : \beta_0 \in \mathbb{R}, \beta_1, \dots, \beta_m \geq 0, \beta_0 l(\theta) \right. \\ \left. + \sum_{i=1}^m \beta_i f_i(\theta) \geq g(\theta)l(\theta), \theta \in \Theta \right\} \end{aligned}$$

when $\mathcal{M}_1$ is $\neq \emptyset$. Moreover the supremum is attained.

A6 can be rewritten as

$$\text{A6'} \quad \exists \tilde{\beta}_0, \tilde{\beta}_1, \dots, \tilde{\beta}_m > 0 : \forall \theta \in \Theta \text{ it is } \tilde{\beta}_0 l(\theta) + \sum_{i=1}^m \tilde{\beta}_i f_i(\theta) \geq g(\theta)l(\theta) + \eta, \eta > 0,$$

which is known in LSIP literature as strong Slater condition (SSC) (see, e.g., [8, p. 128]). It is also worth noticing that A6 is also equivalent to the following more intuitive condition:

$$\text{A6''} \quad \exists \bar{\epsilon} > 0, \exists \tilde{\beta}_1, \dots, \tilde{\beta}_{m+1} \geq 0 : \sum_{i=1}^m \tilde{\beta}_i f_i(\theta) - \tilde{\beta}_{m+1} g(\theta)l(\theta) > 1 \quad \forall \theta : l(\theta) \leq \bar{\epsilon},$$

which states that in the regions where the likelihood l is small (close to 0) some f_i or $-gl$ must be positive and away from 0.

As far as A7 is concerned, it essentially requires that when $\theta \rightarrow \infty$ some $f_i(\theta)$ or $-g(\theta)l(\theta)$ diverges to $+\infty$.

In the case of Θ compact, then Theorem 1 holds under A1, A2, and A6, the other assumptions being true as a consequence of the compactness of Θ .

LSIP Algorithms and Bayesian Robustness

Theorem 1 provides the link between prior robustness under generalized moment conditions and LSIP, which is concerned with optimization problems stated as, under standard notation,

$$(P) \quad I = \inf \{c'x : x \in \mathbb{R}^n, a'(\theta)x \geq b(\theta) \quad \forall \theta \in \Theta\} \quad (4)$$

where Θ is an arbitrary infinite index set.

This enables the numerical treatment of robustness analysis by algorithms developed in this latter context which has been widely considered in the optimization literature (see, e.g., [8] and [9], for comprehensive surveys). Among the algorithms for solving

LSIP problems, and hence for determining the bounds required by robustness analysis via Theorem 1, the classification in [8, p. 253], states that *discretization methods* (by *grids* and by *cutting planes*, *CPs*) have the highest reputation in terms of computational efficiency, according to most experts' opinions. Being also the most popular and easy to implement, they are the natural candidates for the numerical treatment of Bayesian robustness analysis, although algorithms not falling into this category have also been proposed in the literature.

Grid discretization methods are based on sequences of finite grids Θ_r (by grid we mean a finite or countable set of points belonging to Θ , not necessarily regularly displaced) such that $\Theta_r \subset \Theta_{r+1}$ (expansive grids) and dense in Θ . For each grid an ordinary linear programming (LP) program is solved, corresponding to the subproblem (P^r) of (P) obtained substituting Θ by Θ_r , until the solution of (P^r) , x^r , is such that $a'(\theta)x^r - b(\theta) \geq -\varepsilon \forall \theta \in \Theta$ for some prefixed accuracy $\varepsilon > 0$. Convergence is obtained if all the functions defining (P) are continuous, which rules out indicator functions, commonly involved in robustness analysis. The final x^r is not guaranteed to be feasible for (P) (in this case it would also be optimal), so that grid algorithms in general underestimate the optimal solution; when applied to robustness analysis, this means that the posterior range of the quantity of interest will be in general underestimated, leading to an erroneous conclusion about the presence of robustness.

Cutting plane (CP) discretization methods, first proposed in [10], starting from a finite set of constraints and solving the corresponding LP subproblem of (P) , proceed by adding to the current constraint set Θ_r a finite number of constraints of (P) which are violated at the current subproblem solution. Convergence can be proved under more general conditions than the ones required by grid methods, essentially the boundedness of the functions a in (P) , which can always be achieved by normalization.

Within the same scheme of things, but not included in the general setting of the basic CP algorithm is the central cutting plane (CCP) method in [11], originally proposed for general convex programming problems; the method is aimed at providing at each iteration a feasible point (unlike both grid discretization and ordinary CP methods) and at preventing the numerical instability of CP algorithms due to the fact that cuts tend to be parallel near

optimal points. CCP method proceeds by computing at each step the center x^r and the radius of the maximal sphere inscribed into the polytope $\{c'x \leq \alpha_r, a'(\theta)x \geq b(\theta), \theta \in \Theta_r\}$, cutting down α_r to $c'x^r$ if x^r is feasible for (P) (*objective cut*) and enlarging the constraint set to $\Theta_r \cup \theta_r$, where θ_r is any point in Θ giving a violated constraint, otherwise (*feasibility cut*). Convergence is achieved when the functions a are bounded and a SSC condition $a'(\theta)\hat{x} \geq b(\theta) + \eta, \forall \theta \in \Theta$ and some $\eta > 0$ and $\hat{x} \in \mathbb{R}^n$, holds.

In [12] a modification of the CCP algorithm, called ACCP algorithm, was introduced with the aim of accelerating its convergence. The ACCP algorithm shares with the CCP algorithm features that make it attractive for application to robustness analysis, i.e., its applicability under mild conditions (no regularity condition is required on the functions defining the constraints) and the fact that it produces a convergent sequence of points interior to the feasible region (so that the infimum can be approximated from above rather than from below as in the grid methods); its advantage over the CCP method is the superior rate of convergence (see [12, Theorem 3]).

In order to apply the ACCP algorithm to the LSIP problem stated by Theorem 1, all that is needed is to ensure that

1. A SSC holds: this is assumption A6', equivalent to A6.
2. $\sup_{\theta \in \Theta} |f_i(\theta)| < +\infty$ for $i = 1, \dots, m$; if some of the f_i 's are not bounded from above (recall that the f_i 's are bounded from below by assumption A3), then it is possible to transform the problem into an equivalent one in which all required boundedness conditions are satisfied. By the existence of positive $\tilde{\beta}_0, \tilde{\beta}_1, \dots, \tilde{\beta}_{m+1}$ such that $\tilde{s}(\theta) = \tilde{\beta}_0 l(\theta) + \sum_{i=1}^m \tilde{\beta}_i f_i(\theta) - \tilde{\beta}_{m+1} g(\theta) l(\theta) \geq 1 \forall \theta \in \Theta$, the minimization problem of Theorem 1 can be shown to be equivalent to the one $\inf\{\beta_0 : \beta_0 \in R, \beta_1, \dots, \beta_m \geq 0, \beta_0 \tilde{l}(\theta) + \sum_{i=1}^m \beta_i \tilde{f}_i(\theta) \geq g(\theta) \tilde{l}(\theta), \theta \in \Theta\}$, where $\tilde{l}(\theta) = l(\theta)/\tilde{s}(\theta)$ and $\tilde{f}_i(\theta) = f_i(\theta)/\tilde{s}(\theta)$. For this problem too a SSC holds (with coefficients $\tilde{\beta}_i, i = 0, \dots, m+1$); moreover, all the functions involved are now bounded.
3. There exists a suitable finite subset $\Theta_0 \subset \Theta$ such that the corresponding finite subproblem of (P) has nonempty bounded level sets; this is the case when there exists a finite measure $\bar{\nu}$ with finite

4 Bayesian Robustness: Theory and Computation

support such that $\int l d\bar{\nu} > 0$ and $\int f_i d\bar{\nu} \leq -\epsilon_i$, with $\epsilon_i > 0, i = 1, \dots, m$, as can be proved by Corollary 9.3.1 in [8], as restated by Corollary 3 in [13].

A CP algorithm is also introduced in [14] for calculating a Γ -minimax test, where Γ is given by a finite set of generalized moment conditions. The inner maximization problem which is required to solve is a generalized moment problem in the form of expression (3) with $l \equiv 1$, i.e., ν is a probability measure.

An Example

We consider the classical example of Bayesian analysis for failure data of 10 pumps in a nuclear plant reported in [15].

Let $X_i, i = 1, \dots, 10$ be the number of failures of the pump i during the observation period t_i of the pump. The observed values x_i and the respective t_i are reported in Table 1.

X_i is assumed to be distributed according to a Poisson distribution with mean $\lambda_i t_i$, where λ_i is the expected number of failures per unit of time. The pumps are assumed to work independently, so that $X = (X_1, \dots, X_{10})$ has density $f(x|\lambda) = \prod_{i=1}^{10} e^{-\lambda_i t_i} (\lambda_i t_i)^{x_i} / x_i!$.

The rates λ_i are assumed to be independent realizations from a gamma distribution with parameters α and θ , so that λ has density $p(\lambda|\alpha, \theta) = \prod_{i=1}^{10} \theta^\alpha \lambda_i^{\alpha-1} e^{-\theta \lambda_i} / \Gamma(\alpha)$.

We assume α known, $\alpha = 1.8$, and write $p(\lambda|\theta)$ instead of $p(\lambda|\alpha, \theta)$. Let $\pi(d\theta)$ the prior distribution of θ . Then, for the posterior density of λ it is

$$p(\lambda|x) = \frac{f(x|\lambda) \int p(\lambda|\theta) \pi(d\theta)}{\int f(x|\lambda) p(\lambda|\theta) d\lambda \pi(d\theta)} \quad (5)$$

and hence, for a given function of λ , say $g(\lambda)$, it is

$$E(g(\lambda)|x) = \frac{\int g_1(\theta) l(\theta) \pi(d\theta)}{\int l(\theta) \pi(d\theta)} \quad (6)$$

where

$$l(\theta) = \int f(x|\lambda) p(\lambda|\theta) d\lambda, \\ g_1(\theta) = \frac{\int g(\lambda) f(x|\lambda) p(\lambda|\theta) d\lambda}{l(\theta)} \quad (7)$$

After standard computations, we obtain

$$l(\theta) = C \frac{\theta^{10\alpha}}{\prod_{i=1}^{10} (\theta + t_i)^{\alpha+x_i}}, \\ C = \prod_{i=1}^{10} t_i^{x_i} \frac{\Gamma(\alpha + x_i)}{\Gamma(1 + x_i) \Gamma(\alpha)}, \\ g_1(\theta) = \sum_{i=1}^{10} \frac{\alpha + x_i}{\theta + t_i} \quad (8)$$

In order to apply the generalized moment theory, we assume $\theta \in [0, 5]$, so that Θ is compact and hence we need only to verify assumptions A1, A2, A6 (or A6').

We need now to consider some constraints. Suppose that prior information, rather than being reflected into an explicit form for π , is expressed through the marginal probability of no failure of one of the pumps, say pump 1, in a period of length t as a function of t :

$$p_0(t) = P(X_1(t) = 0) \\ = \int P(X_1(t) = 0 | \lambda_1) p(\lambda_1 | \theta) \pi(d\theta) \\ = \int e^{-\lambda_1 t} \frac{1}{\Gamma(\alpha)} \theta^\alpha \lambda_1^{\alpha-1} e^{-\theta \lambda_1} d\lambda_1 \pi(d\theta) \\ = \int \left(\frac{\theta}{\theta + t} \right)^\alpha \pi(d\theta) \quad (9)$$

where $X_1(t)$ represents the number of failures during a period of length t for the pump 1.

Table 1 Failure data in the nuclear plant test case

Pump	1	2	3	4	5	6	7	8	9	10
Failures	5	1	5	14	3	19	1	1	4	22
Period	94.32	15.72	62.88	125.76	5.24	31.44	1.05	1.05	2.10	10.48

Let us assume that lower bounds of $p_0(t)$ are elicited at m values of t , $p_0(\bar{t}_i) \geq \alpha_i$, $i = 1, \dots, m$. Then, $f_i(\theta) = -(\theta/(\theta + \bar{t}_i))^\alpha + \alpha_i$, $i = 1, \dots, m$.

As $f_i(\theta) > 0$ at $\theta = 0$, the only point at which $l(\theta) = 0$, then A6' holds. A1, A2 obviously hold too.

Suppose that $m = 6$, $\bar{t}_1 = 0.001$, $\bar{t}_2 = 0.01$, $\bar{t}_3 = 0.1$, $\bar{t}_4 = 0.5$, $\bar{t}_5 = 1$, $\bar{t}_6 = 3$ and $\alpha_1 = 0.4$, $\alpha_2 = 0.25$, $\alpha_3 = 0.10$, $\alpha_4 = 0.04$, $\alpha_5 = 0.02$, $\alpha_6 = 0.01$.

We remark that the above conditions define a generalized moment class which the usual prior for $\theta \Gamma(0.1, 1)$ belongs to, as can be seen by numerical integration.

The ACCP algorithm gave the following bounds: $4.74415 \leq E(g(\lambda)|x) \leq 11.41814$. The upper bound and the lower bound were obtained, respectively, in 3.5 s and in 0.1 s of CPU time on a computer with an Athlon MP 1600+ processor.

References

- [1] Berger, J.O., Rios Insua, D. & Ruggeri, F. (2000). Bayesian robustness, in *Robust Bayesian analysis, Lecture Notes in Statistics, Vol. 152*, D. Rios Insua, & F. Ruggeri, eds, Springer, New York, pp. 1–32.
- [2] Rios Insua, D. & Ruggeri, F. (eds) (2000). *Robust Bayesian Analysis*, Springer, New York.
- [3] Betró, B., Męczarski, M. & Ruggeri, F. (1994). Robust Bayesian analysis under generalized moments conditions, *Journal of Statistical Planning and Inference* **41**, 257–266.
- [4] Betró, B., Bodini, A. & Guglielmi, A. (2006). Generalized moment theory and robust Bayesian analysis: the nonparametric case, *Annals of the Institute of Statistical Mathematics* **58**, 721–738.
- [5] Kemperman, J.H.B. (1971). On a class of moment problems, *Proceedings Sixth Berkeley Symposium on Mathematical Statistics and Probability*, University of California, Vol. 2, pp. 101–106, June–July 1970.
- [6] Kemperman, J.H.B. (1983). On the role of duality in the theory of moments, *Semi-Infinite Programming and Applications, Lecture Notes in Economics and Mathematical Systems, Vol. 215*, Springer, Berlin, pp. 63–92.
- [7] Kemperman, J.H.B. (1987). Geometry of the moment problem, in *Moments in Mathematics, Proceedings of Symposia in Applied Mathematics*, H.J. Landau, ed, AMS, Providence, pp. 16–53.
- [8] Goberna, M.A. & López, M.A. (1998). *Linear Semi-Infinite Optimization*, John Wiley & Sons, New York.
- [9] Gustafson, S. & Kortanek, K.O. (1973). Numerical treatment of a class of semi-infinite programming problems, *Naval Research Logistics Quarterly* **20**, 477–504.
- [10] Reemsten, R. & Rückmann, J.J. (eds) (1998). *Semi-Infinite Programming*, Kluwer Academic Publishers, Dordrecht.
- [11] Elzinga, J. & Moore, T. (1975). A central cutting plane algorithm for the convex programming problem, *Mathematical Programming* **8**, 134–145.
- [12] Betró, B. (2004). An accelerated central cutting plane algorithm for semi-infinite linear programming, *Mathematical Programming. Series A* **101**, 479–495.
- [13] Betró, B. & Guglielmi, A. (2000). Methods for global prior robustness under generalized moment conditions, in *Robust Bayesian Analysis, Lecture Notes in Statistics, Vol. 152*, D. Rios Insua & F. Ruggeri, eds, Springer, New York, pp. 273–293.
- [14] Noubiap, F.R. & Seidel, W. (2001). An algorithm for calculating Γ -minimax decision rules under generalized moment conditions, *The Annals of Statistics* **29**, 1094–1116.
- [15] Gaver, D.P. & O’Muircheartaigh, I.G. (1987). Robust empirical Bayes analysis of event rates, *Technometrics* **29**, 1–15.

Related Articles

Bayesian Control Charts; Bayesian Design of Reliability Experiments; Bayesian Networks; Bayesian Networks in Reliability; Bayesian Reliability Analysis; Empirical Bayes Estimation of Reliability; Expert Opinion in Reliability; Hierarchical Markov Chain Monte Carlo (MCMC) for Bayesian System Reliability; Life Testing from a Bayesian Perspective; Nonparametric and Semi-parametric Bayesian Reliability Analysis; Prior Distribution Elicitation; Prior Information in Sampling Schemes; Repairable Systems: Bayesian Analysis; Software Reliability: Bayesian Analysis; Statistical Process Control, Bayesian; Survival Analysis, Nonparametric.

BRUNO BETRÒ

Survival Analysis, Nonparametric

Introduction

Statistical analysis of survival (failure time) data is of great importance in several biomedical and engineering fields. We consider survival data arising from a (continuous) distribution on $\mathbb{R}^+$ (the positive real line), with distribution function $F(t)$ and density function $f(t)$. Inferential objectives involve the survival function, $S(t) = 1 - F(t)$, the cumulative hazard function, $H(t) = -\log S(t)$, or the hazard function, $h(t) = f(t)/S(t)$. Inference approaches need to account for censoring (and/or truncation), the distinguishing feature of survival data. Moreover, several modeling formulations have been explored to incorporate **covariate** information.

There exists a substantial literature on classical (frequentist) semiparametric regression techniques for survival data analysis (see, e.g., [1, 2]). Here, the survival distribution is left unspecified with interest focusing on **estimation** for the regression coefficients. Hence, lacking a likelihood specification, point estimation proceeds through optimization of a particular estimating equation, with further inference relying on asymptotic or **resampling** methods.

Parametric probabilistic modeling for survival analysis applications has also received much attention. Such modeling approaches are in contrast to classical nonparametric methods, since, here, all the unknown functions and/or distributions of the model are described through parametric forms (specified up to a number of unknown parameters). Hence, either likelihood or Bayesian estimation techniques can be utilized for inference. Parametric models for survival data analysis are widely available in statistical packages and are routinely used by practitioners. However, they might yield inaccurate inference or misleading predictions when the postulated distributional assumptions are not supported by the data.

Bayesian nonparametric methods provide the means to enhance the flexibility of standard parametric models retaining a probabilistic framework for inference. Under the Bayesian nonparametric paradigm, the unknown functions or distributions of the model are treated as the random (infinite-dimensional) parameters. Hence, from a theoretical

point of view, Bayesian nonparametric approaches require priors (formally, random probability measures) over spaces of functions or distributions. The area of Bayesian nonparametrics has grown rapidly following the work in [3] on the Dirichlet process (DP), a random probability measure on spaces of distribution functions. As with most research areas in Bayesian statistics, Bayes nonparametrics has benefited enormously from advances in simulation-based model fitting, mainly Markov chain Monte Carlo (MCMC) methods (see **Markov Chain Monte Carlo, Introduction and Convergence and Mixing in Markov Chain Monte Carlo** for background on MCMC posterior simulation methods). With an increasing amount of emphasis on applications over the past 10 years or so, it is now easy to argue for the utility and potential of Bayesian nonparametric modeling. See [4–6] for general reviews of Bayesian nonparametrics.

Bayesian nonparametric methods are, arguably, very well suited for survival data analysis, since they: enable flexible modeling for the unknown survival function, cumulative hazard function, or hazard function; provide techniques to handle censoring under a unifying inferential framework; allow incorporation of prior information; and yield rich inference that does not rely on restrictive parametric specifications or asymptotic approximations for large sample sizes. There is, by now, a relatively large literature on Bayesian nonparametric approaches to survival analysis (and reliability analysis) problems, including different prior models for survival, cumulative hazard, and hazard functions. For related references, see the reviews in [7, 8] as well as articles **Nonparametric and Semiparametric Bayesian Reliability Analysis**; **Lévy Processes, Simulation of**; **Simulation of the Beta-Stacy Process**; **Dirichlet Process, Simulation of**; **Polya Trees and Their Use in Reliability and Survival Analysis**.

In this article, we review Bayesian nonparametric and semiparametric methods for survival analysis using mixture models for distributions. In the next section, we consider DP mixture models for survival distributions in a fully nonparametric setting. In the process, we provide a brief introduction to DPs and DP mixtures. The section titled “Semiparametric Accelerated Failure Time Regression Models” presents practically important extensions to regression settings. Finally, we conclude with a summary.

Dirichlet Process Mixture Models for Survival Distributions

Here, we consider fully nonparametric modeling for survival distributions for problems that do not include a regression component. The objective of such modeling is to increase the flexibility of commonly used parametric models, such as, the gamma, Weibull, or lognormal families of survival distributions.

In looking beyond standard parametric families, one is naturally led to mixture models. In particular, finite mixtures provide flexible modeling and are now feasible to implement using MCMC posterior simulation techniques. For instance, finite Weibull mixtures for survival analysis are studied in [9].

Though it may appear paradoxical, rather than handling the large number of parameters resulting from finite mixtures with a large number of mixture components, it may be easier to work with an infinite-dimensional specification by assuming a random mixing distribution that is not restricted to a specified parametric family. The DP is the most widely used in this context, following [10–12].

Under the DP mixture setting, the random distribution function for the survival distribution can be expressed as

$$F(t; G) = \int K(t; \theta) dG(\theta), \quad t \in \mathbb{R}^+,$$

$$G \sim \text{DP}(\alpha, G_0) \quad (1)$$

Here, the parametric kernel of the mixture, $K(\cdot; \theta)$, is a distribution function on $\mathbb{R}^+$ (with parameter vector θ), and $\text{DP}(\alpha, G_0)$ denotes the DP prior for the random mixing distribution G . The DP prior is specified through two hyperparameters, a positive scalar parameter α , and a distribution function $G_0 \equiv G_0(\psi)$, which, in general, depends on parameters ψ and defines the prior expectation of G . Hence, G_0 plays the role of the center of the DP prior whereas α can be viewed as a precision parameter; the larger α is the *closer* a DP realization is expected to be to G_0 .

A very useful characterization of the DP is given by its *stick-breaking* constructive definition (see [13]). On the basis of this definition, a realization G from $\text{DP}(\alpha, G_0)$ is (almost surely) of the form $G = \sum_{j=1}^{\infty} \omega_j \delta_{\theta_j}$, where δ_x denotes a point mass at x , $\omega_1 = z_1$, $\omega_j = z_j \prod_{s=1}^{j-1} (1 - z_s)$, $j = 2, 3, \dots$, with z_s independent from a $\text{beta}(1, \alpha)$ distribution, and,

independently, θ_j independent from G_0 . Evidently, DP realizations are discrete distributions. In particular, $F(t; G)$ in equation (1) admits the (almost sure) representation $\sum_{j=1}^{\infty} \omega_j K(t; \theta_j)$, that is, it can be expressed as a countable mixture with weights generated according to the stick-breaking scheme above.

The choice of the mixture kernel in equation (1) is discussed in [14], where a Weibull DP mixture is developed. In principle, any other two-parameter kernel, which yields unimodal as well as decreasing density shapes, should provide similar inferences. In particular, that should be the case for the gamma DP mixture proposed in [15].

Illustrating with the model from [14], let $K(t; \gamma, \lambda) = 1 - \exp(-\lambda^{-1} t^\gamma)$ be the Weibull distribution function, with shape parameter $\gamma > 0$ and scale parameter $\lambda > 0$, where $\theta = (\gamma, \lambda)$. For clarity of exposition, we consider right censoring; left and interval censored data can be handled in a similar fashion. Suppose that the study involves n subjects, and let t_i and c_i denote the survival time and a fixed censoring time, respectively, for the i th subject. Hence, the actual i th survival time is observed only if $t_i \leq c_i$, and it is right censored otherwise. The observed data consist of the pairs (y_i, d_i) , $i = 1, \dots, n$, where $y_i = \min\{t_i, c_i\}$ and $d_i = 1$, if $t_i \leq c_i$, with $d_i = 0$, if $t_i > c_i$.

The full Bayesian model can be expressed in a hierarchical form by introducing a latent mixing parameter vector (γ_i, λ_i) associated with each observation. Then, y_i , given (γ_i, λ_i) , are independent with first stage specification $\{\lambda^{-1} \gamma_i y_i^{\gamma_i - 1}\}^{d_i} \exp(-\lambda_i^{-1} y_i^{\gamma_i})$. Next, (γ_i, λ_i) , given G , are independent and identically distributed (i.i.d.) from G , with the $\text{DP}(\alpha, G_0)$ prior assigned to G . The model is completed with hyperpriors for α and for (a subset of) the parameters of G_0 . Thus, the full posterior involves the infinite-dimensional parameter G as well as (γ_i, λ_i) , α , and ψ .

Regarding prior specification, note that the prior expectation for $F(t; G)$ is given by $\int K(t; \theta) dG_0(\theta)$, a result that can be used to drive the choice for G_0 and for the prior on its parameters ψ . Moreover, the choice of the prior for α is facilitated by the role that α plays in the DP-induced clustering of (γ_i, λ_i) (see, e.g. [12]).

We refer to **Dirichlet Process, Simulation of** for a review of MCMC posterior simulation for DP mixtures. We simply note here that MCMC inference is based on the marginalized version of the

model, which results by integrating G over its DP prior. Hence, regardless of the MCMC method that might be used, inference for functionals of the random mixture $F(\cdot; G)$ is limited, especially, in this fully nonparametric survival analysis context. Without obtaining posterior samples for G , one can only estimate the posterior predictive survival function, which, evaluated at $t \in \mathbb{R}^+$, is the posterior expectation of $1 - F(t; G)$. An approach to full inference for any functional of the survival distribution that might be of interest is presented in [14]. The posterior distribution of G is sampled using a key result from [10], which provides the conditional posterior of G , given (γ_i, λ_i) , α , and ψ , as a DP with updated parameters. Samples from this DP can be obtained using its constructive definition with a truncation approximation. Inference for median survival times, survival functions, and hazard functions is illustrated in [14] with right-censored data.

Semiparametric Accelerated Failure Time Regression Models

In this section, we review the literature on Bayesian semiparametric regression approaches for survival data under the accelerated failure time (AFT) setting.

Denote by $\mathbf{x}_i = (x_{i1}, \dots, x_{ip})^T$ the vector of (known) covariates corresponding to survival time t_i , $i = 1, \dots, n$. Under the AFT setting, the covariates act multiplicatively on the survival time, the most common formulation being, $t_i = \exp(-\mathbf{x}_i^T \boldsymbol{\beta}) V_i$, where $\boldsymbol{\beta}$ is the vector of regression coefficients. The error terms V_i are assumed independent from a distribution function $F(\cdot)$ on $\mathbb{R}^+$. To avoid restrictions on the model for the error distribution, one can opt not to include an intercept term in $\boldsymbol{\beta}$.

Early semiparametric work for the AFT regression model was reported in [16] and [17]. This work was based on a DP prior (instead of a DP mixture prior) for $F(\cdot)$. These papers provide a careful study of the posterior distributions resulting from the DP prior. In particular, they highlight the practical difficulties of a fully Bayesian analysis, especially, in the presence of censoring.

More recent Bayesian work in [15, 18–22] utilizes DP mixture priors for $F(\cdot)$, taking advantage of MCMC techniques for DP mixtures. Here, the DP mixture model in equation (1) is used for the errors, that is, $V_i | G$ are assumed i.i.d. from $F(v; G) =$

$\int K(v; \boldsymbol{\theta}) dG(\boldsymbol{\theta})$, $v \in \mathbb{R}^+$, with a DP prior for G . In fact, as discussed below, some of the approaches involve a semiparametric error model with DP mixing on a portion of $\boldsymbol{\theta}$ only. Introducing latent mixing parameters $\boldsymbol{\theta}_i$ for the t_i , the generic DP mixed AFT regression model can be expressed hierarchically as follows:

$$\begin{aligned} t_i &= \exp(-\mathbf{x}_i^T \boldsymbol{\beta}) V_i, \quad V_i | \boldsymbol{\theta}_i \stackrel{\text{i.i.d.}}{\sim} K(v; \boldsymbol{\theta}_i), \\ i &= 1, \dots, n \\ \boldsymbol{\theta}_1, \dots, \boldsymbol{\theta}_n &| G \stackrel{\text{i.i.d.}}{\sim} G, \\ G | \alpha, \psi &\sim \text{DP}(\alpha, G_0(\psi)) \end{aligned} \quad (2)$$

with priors added for $\boldsymbol{\beta}$ and, possibly, for α and/or ψ .

Here, for simpler notation, we do not consider censoring, which, however, can be handled. Indeed, with the exception of [20], all the papers above develop MCMC methods for inference under right censoring. This is accomplished with either data augmentation techniques to impute the latent survival times for the subjects with right-censored observations, or, more efficiently, through direct inclusion of the right-censored times in the first stage specification of equation (2) as discussed in the section titled “Dirichlet Process Mixture Models for Survival Distributions”.

A normal DP mixture for the V_i is used in [18], implemented with a fixed variance for the normal kernel. A general DP mixture for the log V_i is also discussed. For instance, a normal DP mixture in this case would result in a DP mixture of lognormal distributions for the V_i , and might be a more natural choice, since it models the errors on $\mathbb{R}^+$. Moreover, this paper considers the practically important problem of variable selection using a stochastic search algorithm.

Weibull DP mixtures are utilized in [20, 21]. In both cases, the error model is semiparametric, involving DP mixing only on the scale parameter of the Weibull kernel with a common (random) shape parameter. The former paper focuses on applications in reliability analysis. In fact, the model is presented within the proportional hazards regression framework. However, given the choice of the Weibull distribution, it can also be represented as an AFT model. The latter paper studies applications in survival analysis. However, to enable use of existing software, the Weibull DP mixture in [21]

is, in fact, implemented approximately as a finite mixture. This paper also reports results on posterior consistency, albeit under the assumption that the *true* response distribution is a scale Weibull mixture.

Gamma DP mixture models for right-censored survival data are studied in [15] and [22]. Here, DP mixing is with respect to both the shape and scale parameters of the gamma kernel. To facilitate MCMC posterior simulation, G_0 is taken to be the product of two exponentials. In fact, the emphasis in [22] is on comparison of the modeling and resulting inference with an alternative nonparametric mixture model based on a normalized inverse-Gaussian prior (see [23]) for G .

Finally, extensions to median regression and median residual life regression are proposed in [19] and [24], respectively. In both cases, the methodology is based on the AFT setting (on a logarithmic scale) with structured prior models, which yield unimodal error distributions on $\mathbb{R}$ with zero median. Two error models are developed, a semiparametric scale DP mixture of skewed normals, and a more general nonparametric scale DP mixture of uniforms.

Summary

We have reviewed research work on Bayesian nonparametric and semiparametric survival analysis. We have considered approaches that are based on one of the most commonly used classes of nonparametric priors, specifically, DP mixture models. The discussion has focused on nonparametric mixture modeling for univariate survival data, in particular, modeling approaches for the survival distribution (“Dirichlet Process Mixture Models for Survival Distributions”) and for the error distribution under the AFT regression setting (“Semiparametric Accelerated Failure Time Regression Models”).

References

- [1] Kalbfleisch, J.D. & Prentice, R.L. (2002). *The Statistical Analysis of Failure Time Data*, Second Edition, John Wiley & Sons, Hoboken.
- [2] Klein, J.P. & Moeschberger, M.L. (2003). *Survival Analysis: Techniques for Censored and Truncated Data*, Second Edition, Springer, New York.
- [3] Ferguson, T.S. (1973). A Bayesian analysis of some nonparametric problems, *Annals of Statistics* **1**, 209–230.
- [4] Walker, S.G., Damien, P., Laud, P.W. & Smith, A.F.M. (1999). Bayesian nonparametric inference for random distributions and related functions (with discussion), *Journal of the Royal Statistical Society. Series B* **61**, 485–527.
- [5] Müller, P. & Quintana, F.A. (2004). Nonparametric Bayesian data analysis, *Statistical Science* **19**, 95–110.
- [6] Hanson T., Branscum A. & Johnson W. (2005). Bayesian nonparametric modeling and data analysis: an introduction, in *Bayesian Thinking: Modeling and Computation, Handbook of Statistics*, D.K. Dey & C.R. Rao, eds, Elsevier, Amsterdam, Vol. 25, pp. 245–278.
- [7] Sinha, D. & Dey, D.K. (1997). Semiparametric Bayesian analysis of survival data, *Journal of the American Statistical Association* **92**, 1195–1212.
- [8] Ibrahim, J.G., Chen, M.-H. & Sinha, D. (2001). *Bayesian Survival Analysis*, Springer, New York.
- [9] Marín, J.M., Rodríguez-Bernal, M.T. & Wiper, M.P. (2005). Using Weibull mixture distributions to model heterogeneous survival data, *Communications in Statistics: Simulation and Computation* **34**, 673–684.
- [10] Antoniak, C.E. (1974). Mixtures of Dirichlet processes with applications to nonparametric problems, *Annals of Statistics* **2**, 1152–1174.
- [11] Lo, A.Y. (1984). On a class of Bayesian nonparametric estimates: I. density estimates, *Annals of Statistics* **12**, 351–357.
- [12] Escobar, M.D. & West, M. (1995). Bayesian density estimation and inference using mixtures, *Journal of the American Statistical Association* **90**, 577–588.
- [13] Sethuraman, J. (1994). A constructive definition of Dirichlet priors, *Statistica Sinica* **4**, 639–650.
- [14] Kottas, A. (2006). Nonparametric Bayesian survival analysis using mixtures of Weibull distributions, *Journal of Statistical Planning and Inference* **136**, 578–596.
- [15] Hanson, T.E. (2006). Modeling censored lifetime data using a mixture of gammas baseline, *Bayesian Analysis* **1**, 575–594.
- [16] Christensen, R. & Johnson, W. (1988). Modelling accelerated failure time with a Dirichlet process, *Biometrika* **75**, 693–704.
- [17] Johnson, W. & Christensen, R. (1989). Nonparametric Bayesian analysis of the accelerated failure time model, *Statistics & Probability Letters* **8**, 179–184.
- [18] Kuo, L. & Mallick, B. (1997). Bayesian semiparametric inference for the accelerated failure-time model, *Canadian Journal of Statistics* **25**, 457–472.
- [19] Kottas, A. & Gelfand, A.E. (2001). Bayesian semiparametric median regression modeling, *Journal of the American Statistical Association* **96**, 1458–1468.
- [20] Merrick, J.R.W., Soyer, R. & Mazzuchi, T.A. (2003). A Bayesian semiparametric analysis of the reliability and maintenance of machine tools, *Technometrics* **45**, 58–69.
- [21] Ghosh, S.K. & Ghosal, S. (2006). Semiparametric accelerated failure time models for censored data, in *Bayesian*

-
- Statistics and its Applications*, S.K. Upadhyay, U.. Singh & D.K. Dey, eds, Anamaya Publishers, New Delhi, pp. 213–229.
- [22] Argiento, R., Guglielmi, A., Pievatolo, A. & Ruggeri, F. (2006). Bayesian semiparametric inference for the accelerated failure time model using hierarchical mixture modeling with N-IG priors, *Proceedings of the 2006 Joint Statistical Meetings, American Statistical Association Section on Bayesian Statistical Science*, Alexandria, pp. 1–8.
- [23] Lijoi, A., Mena, R.H. & Prünster, I. (2005). Hierarchical mixture modeling with normalized inverse-Gaussian priors, *Journal of the American Statistical Association* **100**, 1278–1291.
- [24] Gelfand, A.E. & Kottas, A. (2003). Bayesian semiparametric regression for median residual life, *Scandinavian Journal of Statistics* **30**, 651–665.

Related Articles

Convergence and Mixing in Markov Chain Monte Carlo; Dirichlet Process, Simulation of; Markov Chain Monte Carlo, Introduction; Polya Trees and Their Use in Reliability and Survival Analysis; Nonparametric and Semiparametric Bayesian Reliability Analysis; Lévy Processes, Simulation of; Simulation of the Beta-Stacy Process.

ATHANASIOS KOTTAS

Decision Search via Simulation

Introduction

Experimental design requires the specification of all aspects of an experiment, for example, the number of experimental units needed, which treatments to study in a **clinical trial**, or what levels of input/control variables should be used. When designing an experiment, decisions are often taken before **data collection**. Sample sizes are usually restricted because of costs, ethics, or other limitations on resources or time, hence making efficient use of all the available resources becomes crucial. The basic idea of experimental design consists of improving the inference on the quantity of interest by appropriately selecting the values of the design variables which are under the control of the investigator within the constraints of the available resources. Decision analysis provides a general framework for solving decision making problems under uncertainty. A major use of statistical inference is its application to decision making under uncertainty. Statistically designed experiments are an important tool for decision makers in all areas including engineers, analysts in planning offices and public agencies, project management consultants, manufacturing process planners, financial and economic analysts, experts supporting medical/technological diagnosis, etc. Since specific information is often available before the experiment is carried out, it can be helpfully used to design it. In this respect, the Bayesian approach can play an important role as it provides a natural framework to incorporate prior information into the design problem and account for various sources of variability [1, 2].

An **optimal design** problem is conceptually simple: we look for a design that maximizes the expected utility in a statistical experiment or equivalently minimizes the expected loss. The maximization is taken over some design parameter $d \in D$. The experiment is defined by a probability model for the observable data y taking a value in a set $\mathcal{Y}$ (sample space). Ultimately, the amount we lose as a result of our decision will depend on the observed quantities y . We denote with $p_d(y | \theta)$ the probability distribution of the observables y conditional on some unknown vector of parameter θ . We denote with Θ the parameter

space. In some applications the model may depend on the decision d and therefore we have introduced a subscript d in the probability model. The model is then completed by a prior distribution $p(\theta)$ on the parameter vector θ and

$$p_d(y, \theta) = p_d(y | \theta)p(\theta) \quad (1)$$

Sample sizes, levels of explanatory variables, or other characteristics of the design to be selected are implicitly captured by $p_d(y | \theta)$. The main goal of the experiment may include **estimation** of θ , or some functions of θ , prediction of future observations, model selection, or hypothesis testing. The purpose and the cost of the experiment are usually reflected in the utility function. A general utility function has the form $u(d, \theta, y)$. For example, the utility could be equal to the negative squared error loss $u = -\{\theta - E(\theta | y, d)\}^2$, or something more problem specific like $u =$ total number of successfully treated patients, etc. The utility function models preferences over consequences. It is fundamental to select a utility function that appropriately captures the goals and the various aspects of a given experiment. A design that is optimal for estimation is not necessarily optimal for prediction. For example, in quality control or in clinical trials prediction of future observations can be of special interest.

In most general types of decision problems, decisions must be taken before data collection, i.e. before observing the experimental data y , consequently we need to maximize the expected utility U , given by the expectation of $u(\cdot)$ with respect to (θ, y) . We can then formalize decision making under uncertainty as choosing $d^* \in D$ which maximizes the expected utility:

$$d^* = \arg \max_{d \in D} U(d), \quad \text{where}$$

$$U(d) = \int u(d, \theta, y) p_d(\theta, y) d\theta dy \quad (2)$$

Therefore, in this framework the optimal way to solve a decision problem is to specify a utility function that describes the essential features of the problem and then to select a design that maximizes the expected utility. While conceptually simple, the actual solution of the design problem may be extremely complex, e.g., model and likelihood are chosen in a way that does not allow for analytic evaluation

of the expected utility, the set of possible alternatives D is continuous, or a sequence of decisions is included (*sequential design*). In these scenarios approximations can be used. Most approximations to expected utility usually require a normal approximation of the posterior distribution $p(\theta | y, d)$. See Chaloner and Verdinelli [2] for a review of Bayesian experimental design, which covers analytical and approximate solutions to the optimal design problem.

A very powerful and widely applicable alternative in complex setups is offered by simulation-based strategies. See, for example, Müller [3] for a review. Recent advances in computing have enabled design and decision analysis based on highly structured, simulation-based stochastic models. In the following sections we present some of the most common simulation schemes.

Optimal Design by Monte Carlo Integration

The most simple scheme to evaluate $U(d)$ is to use straightforward **Monte Carlo** simulations. In many problems it is easy to simulate from $p_d(\theta, y) = p(\theta)p_d(y | \theta)$ using efficient random variate generation. Therefore we can use a sample $\{(\theta_i, y_i), i = 1, \dots, M\}$ generated from $\theta_i \sim p(\theta)$, $y_i \sim p_d(y | \theta)$ to approximate the marginal utility by the empirical average:

$$\hat{U}(d) = \frac{1}{M} \sum_{i=1}^M u(d, \theta_i, y_i) \quad (3)$$

and $\hat{U}(d)$ converges almost surely to $U(d)$ by the strong law of **large numbers**. Using the approximate evaluations, $\hat{U}(d)$, we could proceed using a suitable maximization method to find the optimal design d^* . The use of Monte Carlo integration is common in Bayesian optimal design problems. An appealing feature of Monte Carlo integration is that the probability model $p(\theta)$ is not restricted to any particular form; we only need to be able to simulate from it. For example, Sun *et al.* [4] consider a design problem in quantal response analysis, where an experimenter must choose a set of dose levels and number of independent observations to take at those levels, subject to some total sample size, in order to minimize the expected or predicted posterior variance of some

function ϕ (utility function) of a binary response curve. One of the proposed solutions consists of a combination of simulated outcomes and Monte Carlo integration to evaluate the predicted variance. Carlin *et al.* [5] discuss optimal design in a sequential experiment. In this scenario only up to K observations (data monitoring points) are allowed. The design defines a stopping rule, which gives at each time a decision of whether or not to stop the experiment. The authors offer a fully Bayesian approach to this problem, specifying not only the likelihood and **prior distributions** but appropriate loss functions as well. At each data monitoring point, the available decisions are enumerated. **Monte Carlo simulation** from $p_d(\theta, y)$ is used to evaluate expected utility for a given decision. In the context of sequential design, this approach is referred to as the *forward sampling algorithm*.

Müller and Parmigiani [6] propose numerical methods for stochastic optimization by curve fitting of Monte Carlo samples. The goal is to maximize the marginal utility function $U(d)$ with respect to d , when the evaluation of $U(d)$ requires Monte Carlo integration. Rather than estimating $U(d)$ for a design d , they exploit the smoothness of the expected utility surface and “borrow information” from Monte Carlo simulation at neighboring design points d . This is achieved by simulating draws from $p_d(\theta, y)$ and then evaluating the observed utilities. Fitting a smooth surface $\tilde{U}(d)$ through these simulated points serves as an estimate of $U(d)$. We describe here the basic implementation strategy:

Algorithm 1

1. Select $d_i \in D, i = 1, \dots, N$.
2. Draw (θ_i, y_i) from $p_{d_i}(\theta, y)$ for $i = 1, \dots, N$. Compute $u_i = u(d_i, \theta_i, y_i)$.
3. Fit a smooth k -dimensional curve $\tilde{U}(d)$ to the resulting points (d_i, u_i) , where k is the dimension of the decision space.
4. Find the optimal design d^* deterministically.

Simulation from p_d is usually straightforward and the evaluation of the maximum can be performed through standard optimization methods, as the optimal design is simply the mode of the smooth curve $\tilde{U}(d)$. On the other hand, the choice of the curve fitting method and of the number of design points N is fundamental for the performance of the scheme [6].

Augmented Probability Simulation

The methods described in the previous section evaluate expected utilities for given instantiations of the decision variables d (e.g., on a grid of values for d) and therefore they do not easily accommodate continuous decision variables, unless a discretization is carried out. Moreover, computations become more expensive as the dimensionality of the spaces considered increases. The algorithm described in Müller and Parmigiani [6] partly overcomes this limitation by fitting a continuous surface through the expected utility. An alternative and powerful strategy is to recast the optimal design problem equation (2) as a problem of simulation from an augmented probability model $h(d, \theta, y)$ defined on the product space of the alternatives, parameters, and observables [7, 8]. This method consists of defining a joint probability model on $D \times \Theta \times \mathcal{Y}$ so that the marginal distribution on the space of decisions/alternatives is proportional to the expected utility. Assuming that D is bounded, $u(d, \theta, y)$ is nonnegative and bounded, we can specify an artificial distribution

$$h(d, \theta, y) \propto p_d(\theta, y)u(d, \theta, y) \quad (4)$$

From equation (4), we find that the marginal distribution in d is proportional to the expected utility:

$$h(d) \propto \int h(d, \theta, y) d\theta dy = U(d) \quad (5)$$

Hence the optimal design d^* coincides with the mode of the marginal distribution $h(d)$. The underlying idea behind this approach is simple: to sample designs proportionally to their expected utility, we can construct a **Markov chain** with stationary distribution $h(d, \theta, y)$ and then we use the simulated values as an approximate sample from the desired distribution.

Algorithm 2

1. Start with a design d_0 . Simulate (θ_0, y_0) from $p(\theta)p_{d_0}(y | \theta)$. Evaluate $u_0 = u(d_0, \theta_0, y_0)$.
2. Generate a candidate $\tilde{d}$ from a proposal distribution $g(\tilde{d} | d_0)$.
3. Simulate $(\tilde{\theta}, \tilde{y})$ as in 1 under design $\tilde{d}$. Evaluate $\tilde{u} = u(\tilde{d}, \tilde{\theta}, \tilde{y})$.
4. Compute the acceptance probability

$$\alpha = \min \left\{ 1, \frac{h(\tilde{d}, \tilde{\theta}, \tilde{y})}{h(d_0, \theta_0, y_0)} \frac{g(d_0 | \tilde{d})}{g(\tilde{d} | d_0)} \frac{p_{d_0}(\theta_0, y_0)}{p_{\tilde{d}}(\tilde{\theta}, \tilde{y})} \right\}$$

$$= \min \left\{ 1, \frac{\tilde{u}}{u_0} \frac{g(d_0 | \tilde{d})}{g(\tilde{d} | d_0)} \right\}$$

5. With probability α accept the proposal and set $(d_1, u_1) = (\tilde{d}, \tilde{u})$, otherwise set $(d_1, u_1) = (d_0, u_0)$.
6. Repeat steps 2–5 until convergence.

This algorithm defines a Markov chain Monte Carlo (MCMC) scheme to simulate from $h(d, \theta, y)$, using a Metropolis–Hastings chain to update (d, θ, y) [9]. An independent proposal is used to update (θ, y) , while a proposal distribution $g(\tilde{d} | d)$ to draw d needs to be specified. The choice of a symmetric proposal distribution, i.e. $g(\tilde{d} | d) = g(d | \tilde{d})$, leads to a simple expression for the acceptance probability, $\alpha = \min \{1, \frac{\tilde{u}}{u}\}$. The choice of the proposal distribution affects the speed of convergence of the chain. The sample of simulated $d_i, i = 1, \dots, M$ can be used to estimate the marginal distribution $h(d)$. The mode of $h(d)$ corresponds to the optimal design d^* .

Optimal Design by Inhomogeneous Markov Chain Simulation

Algorithm 2 does not entirely solve the original stochastic optimization problem. Expected utility integration is only one of the computational challenges in optimal design. The other key issue is maximization over the design space. If D is low dimensional, we can find the mode from a density estimate, $\hat{h}(d)$. In the case of high-dimensional design spaces, density estimation becomes impracticable. Moreover, in many applications the expected utility surface is relatively flat over a wide range of designs and the sampler would end up drawing roughly uniformly from the dimension of d and information would accumulate slowly, sometimes requiring unfeasibly large sample sizes to estimate the mode of $h(d)$. In fact, Algorithm 2 works best when the utility function mimics a concentrated probability distribution. Müller *et al.* [10] propose an approach to optimal design based on inhomogeneous Markov chains simulations. The approach builds on simulated annealing algorithms [11] by blending simulated annealing and standard MCMC integration to develop an algorithm that simultaneously addresses maximization and integration. The central idea is to replace $h(d)$ with a

4 Decision Search via Simulation

power transformation $h(d)^J$, for a positive integer J . For sufficiently large J , $h_J(d)$ is tightly concentrated around the mode d^* (note that d^* does not need to be unique). In this setup, an artificial distribution h_J is defined as

$$h_J(d, \theta_1, y_1, \dots, \theta_J, y_J) \propto \prod_{j=1}^J p_d(\theta_j, y_j) \times u(d, \theta_j, y_j) \quad (6)$$

i.e., for each d we generate J experiments (θ_j, y_j) and we then consider the product of the observed utilities. Marginalizing equation (6) over $(\theta_1, y_1, \dots, \theta_J, y_J)$, we obtain

$$h_J(d) \propto \left\{ \int u(d, \theta, y) p_d(\theta, y) d\theta dy \right\}^J = U(d)^J \quad (7)$$

Fixing $J = 1$ leads to Algorithm 2. The algorithm starts by simulating from h ($J = 1$). In subsequent iterations of the MCMC algorithm, $n = 1, 2, \dots$, we increment $J = J_n$ as we proceed hence changing the target distribution to h_J and creating an inhomogeneous Markov chain. The chain is set up in such a way that the stationary distribution for fixed J is h_J . As J increases, $h_J(d)$ becomes more concentrated around the mode. To describe the algorithm we use the alternative state vector (d, v) , where v is the average log observed utility $v = 1/J \sum_{j=1}^J \log u(d, \theta_j, y_j)$.

Algorithm 3

1. Assume that the current state of the chain is $(J, d, \theta_1, y_1, \dots, \theta_J, y_J)$ or (d, v) .
2. Generate a candidate $\tilde{d}$ from a proposal distribution $g(\tilde{d} | d)$. For simplicity of explanation, assume that $g(\tilde{d} | d)$ is symmetric, with $g(\tilde{d} | d) = g(d | \tilde{d})$.
3. Simulate experiments $(\tilde{\theta}_j, \tilde{y}_j)$ from $p_{\tilde{d}}(\theta, y)$ for $j = 1, \dots, J$. Evaluate $\tilde{v} = 1/J \sum \log u(\tilde{d}, \tilde{\theta}_j, \tilde{y}_j)$. Thus the proposed new state of the Markov chain is $(\tilde{d}, \tilde{v})$.
4. Compute the acceptance probability

$$\alpha_J = \min \{1, \exp(J\tilde{v} - Jv)\}.$$

5. With probability α_J accept the proposal and set $(d, v) = (\tilde{d}, \tilde{v})$, otherwise leave (d, v) unchanged.
6. Let J_n denote the current value of J . Increase J to $J_{n+1} \geq J_n$, following a chosen cooling schedule $(J_n, n = 1, 2, \dots)$, such that $J_n \rightarrow \infty$.
7. Repeat steps 2–6 until convergence.

For practical implementation, J can be increased only to the point where $h_J(d)$ is sufficiently peaked. Technical details and proof of convergence of the MCMC scheme are given in Müller *et al.* [10].

In Müller *et al.* [10], the algorithm is applied to choose an optimal design for a network of ground-level ozone monitoring stations in the eastern United States. Figure 1(a) shows the location of the 463 monitoring stations. Each station reports daily maxima for 8 hours running average ozone concentrations. A formal solution to this design problem is complicated by the use of complex spatiotemporal models for the ozone measurements and the need to account for several competing goals in the utility function. The optimal network needs to satisfy some key requirements: (a) we want to identify locations with mean ozone level beyond the air quality standard of 85 parts per billion, (b) we wish to accurately estimate ozone levels at stations that are dropped, (c) we wish to make inference about the response surface of mean ozone level across the eastern United States, and (d) we want to minimize the cost. The optimal design for the ozone network problem is shown in Figure 1(b): 209 stations were selected by maximizing the expected utility.

Discussion

A complex decision problem may preclude the possibility of obtaining an exact solution to the optimal design problem. Therefore in many scenarios we need to consider approximation methods, including simulations. We have described various simulation schemes to implement simulation-based optimal Bayesian design. None of the described algorithms is sufficiently general to allow for routine application without adaptation to the specific problem. Choice of utility functions and probability models remain a key issue in any decision analysis. However, with increased computer power, simulation-based methods for finding optimal designs allow for more realistic model, prior, and utility specifications, although the

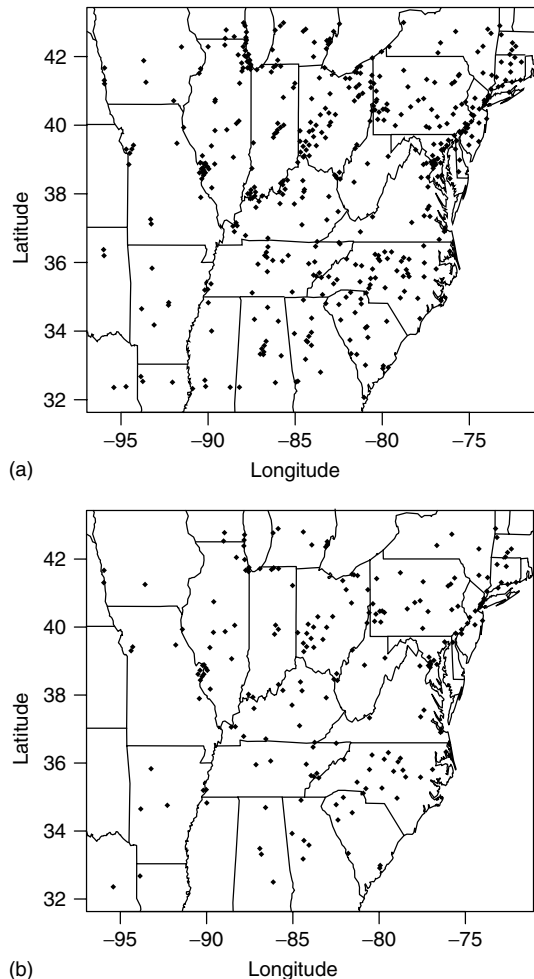


Figure 1 (a) Location of monitoring stations. (b) Optimal design: selected stations

main limitation of these strategies still remains the computationally intensive implementation in some scenarios.

References

- [1] Clyde, M. (2001). Experimental design: a Bayesian perspective, in *International Encyclopedia of the Social and Behavioral Sciences*, N.J. Smelser & P.B. Baltes, eds, Elsevier Science, New York.
- [2] Chaloner, K. & Verdinelli, I. (1995). Bayesian experimental design: a review, *Statistical Science* **10**, 273–304.
- [3] Müller, P. (1999). Simulation based optimal design, in *Bayesian Statistics*, J.O. Berger, J.M. Bernardo, A.P. Dawid & A.F.M. Smith, eds, University Press, Oxford, Vol. 6, pp. 459–474.
- [4] Sun, D., Tsutakawa, R.K. & Lu, W.-S. (1996). Bayesian design of experiment for quantal response: what is promised versus what is delivered, *Journal of Statistical Planning and Inference* **52**, 289–306.
- [5] Carlin, B.P., Kadane, J.B. & Gelfand, A.E. (1998). Approaches for optimal sequential decision analysis in clinical trials, *Biometrics* **54**, 964–975.
- [6] Müller, P. & Parmigiani, G. (1995). Optimal design via curve fitting of Monte Carlo experiments, *Journal of the American Statistical Association* **90**, 1322–1330.
- [7] Bielza, C., Müller, P. & Insua, D.R. (1999). Monte Carlo methods for decision analysis with applications to influence diagrams, *Management Science* **45**, 995–1007.
- [8] Clyde, M., Müller, P. & Parmigiani, G. (1995). *Exploring Expected Utility Surfaces by Markov Chains*, Discussion Paper 95-39, Duke University ISDS.
- [9] Gilks, W.R., Richardson, S. & Spiegelhalter, D.J. (eds) (1996). *Markov Chain Monte Carlo in Practice*, Chapman & Hall, New York.
- [10] Müller, P., Sansó, B. & De Iorio, M. (2004). Optimal Bayesian design by inhomogeneous Markov chain simulation, *Journal of the American Statistical Association* **99**, 788–798.
- [11] van Laarhoven, P.J.M. & Aarts, E.H.L. (1987). *Simulated Annealing: Theory and Applications*, D. Reidel Publishing, Dordrecht.

Related Articles

Assessment of Experimental Designs; Convergence and Mixing in Markov Chain Monte Carlo; Markov Chain Monte Carlo, Introduction; Optimal Design.

MARIA DE IORIO

Lifetime Distributions, Optimization Methods for

Maximum-Likelihood Estimation

Single-Parameter Estimation

Let y be a realization of the random variable Y with probability density $f_Y(y; \theta)$ where θ is an unknown parameter to be estimated. Define the likelihood function $L(\theta) = f_Y(y; \theta)$ as the probability of the observed data y expressed as a function of θ . Although the argument of $f_Y(y; \theta)$ is y , the argument of $L(\theta)$ is θ . When the random variables $Y_1, \dots, Y_n$ are mutually independent, the joint density is equal to the product of the marginal densities, and the likelihood function becomes $L(\theta) = \prod_{j=1}^n f_{Y_j}(y_j; \theta)$.

It is usually more convenient to work with the log-likelihood function $\ell(\theta) = \ln L(\theta)$. Defining the score function $S(\theta) = \partial \ell / \partial \theta$ as the derivative of the log-likelihood, the maximum-likelihood estimate (MLE) $\hat{\theta}$ of θ is the solution to the score equation $S(\theta) = 0$ and is the most plausible value of θ in the sense that it makes the observed sample most probable (see **Estimation; Maximum Likelihood**). We can verify that a solution $\hat{\theta}$ of the equation $S(\theta) = 0$ is actually a maximum by checking that $I(\hat{\theta}) > 0$, where $I(\theta) = -\partial^2 \ell / \partial \theta^2$ is called the *Fisher information function*. The scalar quantity $I(\hat{\theta})$ is referred to as the *observed Fisher information*. A large value of $I(\hat{\theta})$ is associated with a well-defined maximum, indicating less uncertainty about θ .

Example 1: The time to failure T of components has an exponential distribution with mean θ . Suppose that n components are tested for 100 hours and that m components failed at times $t_1, \dots, t_m$, with $n - m$ components surviving the 100-h test. The likelihood function can be written as

$$L(\theta) = \underbrace{\prod_{i=1}^m \frac{1}{\theta} e^{-t_i/\theta}}_{\text{components failed}} \underbrace{\prod_{j=m+1}^n P(T_j > 100)}_{\text{components survived}} \quad (1)$$

Clearly $P(T \leq t) = 1 - e^{-t/\theta}$ implies $P(T > 100) = e^{-100/\theta}$ is the probability of a component surviving

the 100-h test. Then

$$L(\theta) = \left(\prod_{i=1}^m \frac{1}{\theta} e^{-t_i/\theta} \right) (e^{-100/\theta})^{n-m} \quad (2)$$

$$\ell(\theta) = -m \ln \theta - \frac{1}{\theta} \sum_{i=1}^m t_i - \frac{1}{\theta} 100(n-m) \quad (3)$$

$$S(\theta) = -\frac{m}{\theta} + \frac{1}{\theta^2} \sum_{i=1}^m t_i + \frac{1}{\theta^2} 100(n-m) \quad (4)$$

Setting $S(\hat{\theta}) = 0$ suggests the estimator $\hat{\theta} = [\sum_{i=1}^m t_i + 100(n-m)]/m$. Also, $I(\hat{\theta}) = m/\hat{\theta}^2 > 0$, and $\hat{\theta}$ is indeed the MLE. Although failure times were recorded for just m components, this example usefully demonstrates that all n components contribute to the estimation of the mean failure parameter θ . The $n - m$ surviving components are often referred to as *right censored* (see **Reliability Sampling**).

Multiparameter Estimation

Suppose that a statistical model specifies that the data $\mathbf{y}$ has a probability distribution $f(\mathbf{y}; \boldsymbol{\theta})$ depending on a vector of m unknown parameters $\boldsymbol{\theta} = (\theta_1, \dots, \theta_m)$. In this case, the likelihood function is a function of the m parameters $\theta_1, \dots, \theta_m$ and having observed the value of $\mathbf{y}$ is defined as $L(\boldsymbol{\theta}) = f(\mathbf{y}; \boldsymbol{\theta})$ with $\ell(\boldsymbol{\theta}) = \ln L(\boldsymbol{\theta})$.

The MLE of $\boldsymbol{\theta}$ is a value $\hat{\boldsymbol{\theta}}$ for which $L(\boldsymbol{\theta})$, or equivalently $\ell(\boldsymbol{\theta})$, attains its maximum value. For $r = 1, \dots, m$ define $S_r(\boldsymbol{\theta}) = \partial \ell / \partial \theta_r$. Then we can (usually) find the MLE $\hat{\boldsymbol{\theta}}$ by solving the set of m simultaneous equations $S_r(\boldsymbol{\theta}) = 0$ for $r = 1, \dots, m$. The information matrix $\mathbf{I}(\boldsymbol{\theta})$ is defined to be the $m \times m$ matrix whose (r, s) element is given by I_{rs} , where $I_{rs} = -\partial^2 \ell / \partial \theta_r \partial \theta_s$. The conditions for a value $\hat{\boldsymbol{\theta}}$ satisfying $S_r(\hat{\boldsymbol{\theta}}) = 0$ for $r = 1, \dots, m$ to be an MLE are that all the eigenvalues of the matrix $\mathbf{I}(\hat{\boldsymbol{\theta}})$ are positive.

Example 2: A Weibull random variable with “shape” parameter $a > 0$ and “scale” parameter $b > 0$ has density $f_T(t) = (a/b)(t/b)^{a-1} \exp\{-(t/b)^a\}$ for $t \geq 0$. The (cumulative) distribution function is $F_T(t) = 1 - \exp\{-(t/b)^a\}$ on $t \geq 0$. Suppose that

2 Lifetime Distributions, Optimization Methods for

the time to failure T of components has a Weibull distribution and after testing n components for 100 h, m components fail at times $t_1, \dots, t_m$, with $n - m$ components surviving the 100-h test. The likelihood function can be written

$$L(\theta) = \underbrace{\prod_{i=1}^m \frac{a}{b} \left(\frac{t_i}{b}\right)^{a-1} \exp\left\{-\left(\frac{t_i}{b}\right)^a\right\}}_{\text{components failed}} \underbrace{\prod_{j=m+1}^n \exp\left\{-\left(\frac{100}{b}\right)^a\right\}}_{\text{components survived}} \quad (5)$$

Then the log-likelihood function is

$$\ell(a, b) = m \ln(a) - ma \ln(b) + (a-1) \sum_{i=1}^m \ln(t_i) - \sum_{i=1}^n (t_i/b)^a \quad (6)$$

where for convenience we have written $t_{m+1} = \dots = t_n = 100$. This yields score functions

$$S_a(a, b) = \frac{m}{a} - m \ln(b) + \sum_{i=1}^m \ln(t_i) - \sum_{i=1}^n (t_i/b)^a \ln(t_i/b) \quad (7)$$

and

$$S_b(a, b) = -\frac{ma}{b} + \frac{a}{b} \sum_{i=1}^n (t_i/b)^a \quad (8)$$

It is not obvious how to solve $S_a(a, b) = S_b(a, b) = 0$ for a and b .

Newton–Raphson Optimization

When the m equations $S_r(\theta) = 0$, $r = 1, \dots, m$ cannot be solved directly, numerical optimization is required. Let $\mathbf{S}(\theta)$ be the $m \times 1$ vector whose r th element is $S_r(\theta)$. Let $\hat{\theta}$ be the solution to the set of equations $\mathbf{S}(\theta) = \mathbf{0}$ and let θ_0 be an initial guess at

$\hat{\theta}$. Then a first-order Taylor's series approximation to the function S about the point θ_0 is given by

$$S_r(\hat{\theta}) \approx S_r(\theta_0) + \sum_{j=1}^m (\hat{\theta}_j - \theta_{0j}) \frac{\partial S_r}{\partial \theta_j}(\theta_0) \quad (9)$$

for $r = 1, 2, \dots, m$ which may be written in matrix notation as

$$\mathbf{S}(\hat{\theta}) \approx \mathbf{S}(\theta_0) - \mathbf{I}(\theta_0)(\hat{\theta} - \theta_0) \quad (10)$$

Requiring $\mathbf{S}(\hat{\theta}) = \mathbf{0}$, this last equation can be reorganized to give

$$\hat{\theta} \approx \theta_0 + \mathbf{I}(\theta_0)^{-1} \mathbf{S}(\theta_0) \quad (11)$$

Thus given θ_0 this is a method for finding an improved guess at $\hat{\theta}$. We then replace θ_0 by this improved guess and repeat the process. We keep repeating the process until we obtain a value θ^* for which $|S_r(\theta^*)|$ is less than ϵ for $r = 1, 2, \dots, m$ where ϵ is some small number like 0.0001. θ^* will be an approximate solution to the set of equations $\mathbf{S}(\theta) = \mathbf{0}$. We then evaluate the matrix $\mathbf{I}(\theta^*)$ and if all m of its eigenvalues are positive we set $\hat{\theta} = \theta^*$.

Example 2 (continued): For the Weibull distribution

$$I_{aa} = \frac{m}{a^2} + \sum_{i=1}^n (t_i/b)^a [\ln(t_i/b)]^2 \quad (12)$$

$$I_{bb} = -\frac{ma}{b^2} + \frac{a(a+1)}{b^2} \sum_{i=1}^n (t_i/b)^a \quad (13)$$

$$I_{ab} = I_{ba} = \frac{m}{b} - \frac{1}{b} \sum_{i=1}^n (t_i/b)^a \times [\ln(t_i/b) + 1] \quad (14)$$

Given initial values a_0 and b_0 , equation (11) can be used to obtain updated estimates via

$$\begin{pmatrix} a_{\text{new}} \\ b_{\text{new}} \end{pmatrix} = \begin{pmatrix} a_0 \\ b_0 \end{pmatrix} + \begin{pmatrix} I_{aa}(a_0, b_0) & I_{ab}(a_0, b_0) \\ I_{ba}(a_0, b_0) & I_{bb}(a_0, b_0) \end{pmatrix}^{-1} \times \begin{pmatrix} S_a(a_0, b_0) \\ S_b(a_0, b_0) \end{pmatrix} \quad (15)$$

Table 1 Newton–Raphson estimates of the Weibull parameters

a_0	b_0	a^*	b^*	Steps	Eigenvalues
1.8	74.5	1.924941	78.12213	4	All positive
2.36	72	1.924941	78.12213	5	All positive
2	70	1.924941	78.12213	5	All positive
1	2	1.924941	78.12213	15	All positive
80	1	1.924941	78.12213	387	All positive
2	100	4.292059	−34.591322	2	Crashed

Having solved for a and b , we check that a maximum has been achieved by calculating $\mathbf{I}(a^*, b^*)$ and verifying that both of its eigenvalues are positive.

Suppose that $n = 10$ and $m = 8$ items fail at times 17, 29, 32, 55, 61, 74, 77, and 93 (in hours) with the remaining two items surviving the 100-h test. Table 1 shows estimates (a^*, b^*) obtained from equation (15) using various starting values (a_0, b_0) along with the number of iteration steps taken until convergence using $\epsilon = 0.000001$. Four of the starting pairs considered converged to estimates $a^* = 1.924941$ and $b^* = 78.12213$ for a and b . The starting pairs (1, 2) and (80, 1) might misleadingly suggest that the algorithm is robust to the choice of starting values. However, the starting pair (2, 100) produced a negative estimate of b , and as $\ln(b)$ is required in the computation of the score function, caused the algorithm to crash. Other starting pairs not reported failed to yield estimates due to a noninvertible $\mathbf{I}(\theta_0)$ matrix being produced during the running of the algorithm. The remainder of this section describes commonly applied modifications of the Newton–Raphson method.

Initial Values

The Newton–Raphson method depends on the initial guess being close to the true value. If this requirement is not satisfied, the procedure might converge to a minimum instead of a maximum, or just simply diverge and fail to produce any estimates at all. Methods of finding good initial estimates depend very much on the problem at hand and may require some ingenuity.

The distribution function of a Weibull random variable T can be linearized as $\log(-\log[1 - F(t)]) = -a \log(b) + a \log(t)$. A regression (see **Mean Square Error; Least-Squares Estimation**) of the empirical distribution function $\hat{F}(t)$ (see **Quantiles**)

on $\log(t)$ can be used to recover initial estimates for a and b . The starting pair (1.8, 74.5) in Table 1 was obtained in this way.

A more general method for finding initial values is the method of moments. This method estimates the k th moment (see **Moments**) about the origin $E(T^k)$ by the sample moment $\tilde{m}_k = n^{-1} \sum t_i^k$. For a Weibull distribution, the k th moment about the origin equals $b^k \Gamma(1 + k/a)$ where $\Gamma(c) = \int_0^\infty u^{c-1} e^{-u} du$ is the usual gamma function satisfying $\Gamma(c) = (c-1)\Gamma(c-1)$. It is easy to show that $E(T^2)/[E(T)]^2 = 2a\Gamma(2/a)/[\Gamma(1/a)]^2$. An estimate, a_0 , of a can be obtained by solving $\tilde{m}_2^2/(\tilde{m}_1)^2 = 2a\Gamma(2/a)/[\Gamma(1/a)]^2$ for a . The corresponding estimate of b is then $b_0 = \tilde{m}_1/\Gamma(1 + 1/a_0)$. The starting pair (2.36, 72) in Table 1 was obtained in this way.

Fisher’s Method of Scoring

One modification to the Newton–Raphson method is to replace the matrix $\mathbf{I}(\theta)$ in equation (11) with $\bar{\mathbf{I}}(\theta) = E[\mathbf{I}(\theta)]$. The matrix $\bar{\mathbf{I}}(\theta)$ is positive definite, thus overcoming many of the problems regarding matrix inversion. Like Newton–Raphson, there is no guarantee that Fisher’s method of scoring will avoid producing negative parameter estimates or converging to local minima. Unfortunately, calculating $\bar{\mathbf{I}}(\theta)$ can often be mathematically difficult.

The Profile Likelihood

Setting $S_b(a, b) = 0$ from equation (8) we can solve for b in terms of a as

$$b = \left[\frac{1}{m} \sum_{i=1}^n t_i^a \right]^{1/a} \quad (16)$$

This reduces our two-parameter problem to a search over the “new” one-parameter *profile* log-likelihood

$$\begin{aligned} \ell_a(a) &= \ell(a, b(a)) \\ &= m \ln(a) - m \ln \left(\frac{1}{m} \sum_{i=1}^n t_i^a \right) \\ &\quad + (a-1) \sum_{i=1}^m \ln(t_i) - m \end{aligned} \quad (17)$$

4 Lifetime Distributions, Optimization Methods for

Given an initial guess a_0 for a , an improved estimate a_1 can be obtained using a single-parameter Newton–Raphson updating step $a_1 = a_0 + S(a_0)/I(a_0)$, where $S(a)$ and $I(a)$ are now obtained from $\ell_a(a)$. The estimates $\hat{a} = 1.924941$ and $\hat{b} = 78.12213$ were obtained by applying this method to the Weibull data using starting values $a_0 = 0.001$ and $a_0 = 5.8$ in 16 and 13 iterations, respectively. However, the starting value $a_0 = 5.9$ produced the sequence of estimates $a_1 = -5.544163$, $a_2 = 8.013465$, $a_3 = -16.02908$, $a_4 = 230.0001$ and subsequently crashed.

Reparameterization

Negative parameter estimates can be avoided by reparameterizing the profile log-likelihood in equation (17) using $\alpha = \ln(a)$. Since $a = e^\alpha$ we are guaranteed to obtain $a > 0$. The reparameterized profile log-likelihood becomes

$$\begin{aligned} \ell_\alpha(\alpha) = m\alpha - m \ln \left(\frac{1}{m} \sum_{i=1}^n t_i^{e^\alpha} \right) \\ + (e^\alpha - 1) \sum_{i=1}^m \ln(t_i) - m \end{aligned} \quad (18)$$

implying score function

$$\begin{aligned} S(\alpha) = m - \frac{me^\alpha \sum_{i=1}^n t_i^{e^\alpha} \ln(t_i)}{\sum_{i=1}^n t_i^{e^\alpha}} \\ + e^\alpha \sum_{i=1}^m \ln(t_i) \end{aligned} \quad (19)$$

and information function

$$I(\alpha) = \frac{me^\alpha}{\sum_{i=1}^n t_i^{e^\alpha}} \times \left[\sum_{i=1}^n t_i^{e^\alpha} \ln(t_i) - \frac{e^\alpha \left[\sum_{i=1}^n t_i^{e^\alpha} \ln(t_i) \right]^2}{\sum_{i=1}^n t_i^{e^\alpha}} + e^\alpha \sum_{i=1}^n t_i^{e^\alpha} [\ln(t_i)]^2 \right] - e^\alpha \sum_{i=1}^m \ln(t_i) \quad (20)$$

The estimates $\hat{a} = 1.924941$ and $\hat{b} = 78.12213$ were obtained by applying this method to the Weibull data using starting values $a_0 = 0.07$ and $a_0 = 76$ in 103 and 105 iterations, respectively. However, the starting values $a_0 = 0.06$ and $a_0 = 77$ failed due to division by computationally tiny ($1.0e^{-300}$) values.

The Step-Halving Scheme

The Newton–Raphson method uses the (first and second) derivatives of $\ell(\theta)$ to maximize the function $\ell(\theta)$, but the function itself is not used in the algorithm. The log-likelihood can be incorporated into the Newton–Raphson method by modifying the updating step to

$$\theta_{i+1} = \theta_i + \lambda_i \mathbf{I}(\theta_i)^{-1} \mathbf{S}(\theta_i) \quad (21)$$

where the search direction has been multiplied by some $\lambda_i \in (0, 1]$ chosen so that the inequality

$$\ell(\theta_i + \lambda_i \mathbf{I}(\theta_i)^{-1} \mathbf{S}(\theta_i)) > \ell(\theta_i) \quad (22)$$

holds. This requirement protects the algorithm from converging toward minima or saddle points. At each iteration, the algorithm sets $\lambda_i = 1$, and if equation (22) does not hold, λ_i is replaced with $\lambda_i/2$. The process is repeated until the inequality in equation (22) is satisfied. At this point, the parameter estimates are updated using equation (21) with the value of λ_i for which equation (22) holds. If the function $\ell(\theta)$ is concave and unimodal, convergence is guaranteed. Finally, when $\bar{\mathbf{I}}(\theta)$ is used in place of $\mathbf{I}(\theta)$, convergence to a (local) maxima is guaranteed, even if $\ell(\theta)$ is not concave.

Approximate Standard Errors

Let $\hat{\theta} = (\hat{\theta}_1, \dots, \hat{\theta}_m)$ be the MLE of $\theta = (\theta_1, \dots, \theta_m)$. Let V_r^* and V_r be, respectively, the r th

diagonal elements of the matrices $[\tilde{\mathbf{I}}(\hat{\theta})]^{-1}$ and $[\mathbf{I}(\hat{\theta})]^{-1}$. The quantities $\sqrt{V_r^*}$ and $\sqrt{V_r}$ are often used as approximate standard errors of the MLE $\hat{\theta}_r$. Thus an approximate, 95% confidence interval (see **Confidence Intervals**) for θ_r is given by $[\hat{\theta}_r - 1.96\sqrt{V_r^*}, \hat{\theta}_r + 1.96\sqrt{V_r^*}]$. Another approximate, 95% confidence interval for θ_r is given by $[\hat{\theta}_r - 1.96\sqrt{V_r}, \hat{\theta}_r + 1.96\sqrt{V_r}]$.

Regression Problems

The time to failure T_i of electric fuse i , for $i = 1, \dots, n$ is thought to have a γ density $f_{T_i}(t_i; \omega, \lambda_i) = \lambda_i(\lambda_i y)^{\omega-1} e^{-\lambda_i y} / \Gamma(\omega)$ whose mean is given by $E(T_i) = \omega/\lambda_i = \mu_i = \beta_0 + \beta_1 x_i$ where x_i is the voltage applied across the fuse, and the parameter ω is constant across all i . The density of T_i can also be written as $f_{T_i}(t_i; \omega, \theta_i) = \exp\{\omega[t_i \theta_i + \ln(-\theta_i)] + c(t_i, \omega)\}$ for $t_i > 0$, where $\theta_i = -1/\mu_i$ and $c(t_i, \omega) = (\omega - 1) \ln y_i - \ln \Gamma(\omega) + \omega \ln \omega$. For the purposes of this example we will not consider inferences about ω .

A probability distribution is said to belong to the exponential family (of distributions) if its density function $f_Y(y; \theta)$ can be written in the form $f_Y(y; \theta) = \exp\{(\omega/\phi)[y\theta - b(\theta)] + c(y, \phi)\}$, where θ, ϕ are scalar parameters and ω is a known constant. Clearly this is satisfied by the gamma distribution with $\phi = 1$ and $b(\theta) = -\ln(-\theta)$. The exponential family of distributions give rise to a wide class of regression problems, collectively known as *generalized linear models* (see **Generalized Linear Models**). These have three components: (a) a parent distribution for the response variable Y , which in our example is gamma, having expected value $E(Y_i) = \mu_i$; (b) a linear predictor $\eta_i = \mathbf{x}_i^T \boldsymbol{\beta}$, which in our example is $\eta_i = \beta_0 + \beta_1 x_i$; and (c) a link function $g(\mu_i) = \eta_i$, which in our example is $g(\mu_i) = -1/\mu_i$. Using this formulation, a general method can be derived for fitting the models using MLE and Fisher's scoring algorithm, which further reduce to fitting using iterative weighted least squares.

Stochastic Search

Although the Newton–Raphson method is sufficient for most problems, it does not perform well in the

presence of likelihood functions containing multiple optima. Stochastic search routines (for example, see **Hierarchical Markov Chain Monte Carlo (MCMC) for Bayesian System Reliability; Genetic Algorithms in Reliability**) should be considered in such circumstances. One such routine, simulated annealing, works on the basis that at each iteration step the search makes a random jump to a different part of the likelihood function. Downward jumps are always possible, thus allowing the algorithm to escape from local maxima, but such jumps become less likely as the updated value of the estimated maximum increases.

Concluding Remarks

This article gives only a brief account of optimization methods in statistics. For more details see any of the texts [1–4].

The popular statistical software package R contains a general-purpose optimization function `optim()` that implements Newton–Raphson, quasi-Newton, and conjugate gradient algorithms. It also includes an option for simulated annealing. Alternatively, FORTRAN and C code can be obtained from [5] and [6].

References

- [1] Everitt, B.S. (1987). *Introduction to Optimization Methods and their Applications in Statistics*, Chapman & Hall, London.
- [2] Lange, K. (1999). *Numerical Analysis for Statisticians*, Springer, London.
- [3] Monahan, J.F. (2001). *Numerical Methods in Statistics*, Cambridge University Press.
- [4] Nocedal, J. & Wright, S.J. (1999). *Numerical Optimization*, Springer.
- [5] Press, W.H., Flannery, B.P., Teukolsky, S.A. & Vetterling, W.T. (1992). *Numerical Recipes in FORTRAN 77: The Art of Scientific Computing*, Cambridge University Press, New York.
- [6] Press, W.H., Flannery, B.P., Teukolsky, S.A. & Vetterling, W.T. (1992). *Numerical Recipes in C: The Art of Scientific Computing*, Cambridge University Press, New York.

KEVIN HAYES AND DON BARRY

Bootstrap and Jackknife, Overview

Introduction

About a quarter of a century ago, the age of fast, cheap computing dawned. Now, with ever faster computers on the desktop, we live in an age of which previous generations of scientists can only have dreamed. The classical ways of the past now seem rather restrictive – the reliance on simplifying distributional (usually normal) assumptions fails a basic test of usability as the conditions on which they depend do not seem to hold in many important real-world problems. Supported by the advent of unprecedented computational power, new, computer-intensive approaches were developed based on the simple idea that the sample data itself must hold more information than is explicitly used by classical methods, particularly when classical assumptions fail – so-called sample reuse or “resampling” was born. Several types of resampling are now commonplace, with whimsical names like *jackknife* and *bootstrap* capturing the imaginations of statisticians and non-statisticians alike. These methods replace the overt use of tools like central limit theorems to approximate sampling distributions of certain statistics by “generating” these distributions on the basis of the behavior of a statistic under “resampled data” derived from the original data using a resampling rule. In many cases of practical interest, the distributions so generated converge to the true, unknown sampling distributions as sample size grows.

Some sample reuse methods were proposed well before the advent of fast, modern computing. Permutation procedures were developed as early as the 1930s by Fisher and others [1], but their widespread use was delayed by the lack of available computing power. The jackknife was first proposed in 1949 by Quenouille [2] as a method for **bias** reduction in a **time series estimation** context, and it was later popularized by Tukey in a famous 1958 abstract [3] promoting its use beyond its original, restricted setting to variance estimation (see **Variance**). The bootstrap, a method now virtually synonymous with the term *resampling*, was introduced by Efron in 1979 [4], and leveraged the explosion of computational power occurring at the time to allow thousands

of bootstrap resamples to be used in generating the sampling distributions of complex statistics. Efron’s realization in marrying an “old” idea, that of using the empirical cumulative distribution function (CDF) (see **Cumulative Distribution Function (CDF)**) in place of the true unknown CDF in thinking about the sampling distribution of a statistic of interest, with the means to carry out the necessary calculations – resampling from the data with replacement – was a watershed moment in modern statistics.

The predominant resampling methods in common usage are the jackknife, the bootstrap, and permutation tests. Apart from the bootstrap, which is the most utilitarian of the techniques, other techniques have attracted use for specific purposes, the jackknife primarily for bias and variance estimation, and permutation tests for assessing the equality of two distributions. In this article, we will describe some common and important uses for the jackknife and the bootstrap. Another article, *Sampling from Virtual Populations* (see **Sampling from Virtual Populations**) places the jackknife and the bootstrap more generally in the context of sampling algorithms. In the context of statistical reliability, see **Resampling for Lifetime Distributions**, which focuses in particular on the theory behind resampling methods specifically in a lifetime data setting.

The Jackknife

Quenouille [2] introduced the jackknife as a method to correct for bias in time series estimation (see **Time Series Analysis**). Tukey [3] popularized the idea and coined its name in perhaps statistics’ most famous abstract (there was no paper). The basic idea of the jackknife is to mimic a kind of repeated sampling by creating so-called jackknife samples, and then to use the values of the relevant statistic calculated for those samples in forming bias (see **Bias of an Estimator**) and variance estimates. Specifically, given a sample $\mathbf{X} = (X_1, \dots, X_n)$ and an estimator $\hat{\theta} = \hat{\theta}(X) = \hat{\theta}(X_1, \dots, X_n)$ that is symmetric in the data, define jackknife (leave-1-out) samples:

$$\begin{aligned}\mathbf{X}_1 &= (X_2, \dots, X_n), \dots, \\ \mathbf{X}_i &= (X_1, \dots, X_{i-1}, X_{i+1}, \dots, X_n), \dots, \\ \mathbf{X}_n &= (X_1, \dots, X_{n-1})\end{aligned}\tag{1}$$

Now, define the jackknife replicates $\hat{\theta}_{(i)} = \hat{\theta}(X_1, \dots, X_{i-1}, X_{i+1}, \dots, X_n)$, the value of $\hat{\theta}$ computed using the i th jackknife sample. Define $\hat{\theta}_{(\bullet)} = n^{-1} \sum_{i=1}^n \hat{\theta}_{(i)}$. Then, the jackknife estimate of the bias of $\hat{\theta}$ in estimating θ is $\hat{b}_{\text{Jack}}(\hat{\theta}) = (n-1)[\hat{\theta}_{(\bullet)} - \hat{\theta}]$, and so a bias-corrected estimate of θ is $\hat{\theta}_{\text{corrected}} = \hat{\theta} - \hat{b}_{\text{Jack}}(\hat{\theta}) = n\hat{\theta} - (n-1)\hat{\theta}_{(\bullet)}$. For the sample mean, $\hat{\theta}_{(\bullet)}$ is equal to the sample mean, and the jackknife estimate of bias is zero. For the (biased) estimate $\hat{\theta} = n^{-1} \sum_{i=1}^n (X_i - \bar{X})^2$ of a variance, it turns out that $\hat{\theta}_{(\bullet)} = (n-1)^{-2}(n-2) \sum_{i=1}^n (X_i - \bar{X})^2$, and so $\hat{\theta}_{\text{corrected}} = (n-1)^{-1} \sum_{i=1}^n (X_i - \bar{X})^2$, the usual unbiased variance estimator. Basically, the jackknife removes the first-order ($O(n^{-1})$) term in the bias expansion of the original estimator. In the case of the variance discussed above, it removes *all* of the bias, and, indeed, it always does so for quadratic functionals of the underlying distribution function.

In suggesting its use beyond bias estimation, Tukey [3] developed an alternative way to view the jackknife, in terms of so-called pseudovalues, $\tilde{\theta}_i = n\hat{\theta} - (n-1)\hat{\theta}_{(i)}$, based on the jackknife replicates. In terms of the pseudovalues, the jackknife bias-corrected estimate of θ is simply the average of the pseudovalues $\hat{\theta}_{\text{corrected}} = n^{-1} \sum_{i=1}^n \tilde{\theta}_i = n\hat{\theta} - (n-1)\hat{\theta}_{(\bullet)}$. The pseudovalues were constructed so that in the case of the mean, the pseudovalues were just the original data points and the jackknife estimator returned the sample mean. Informally, the i th pseudovalue is a measure of the “influence” the i th data point has on the estimator $\hat{\theta}$, and so the pseudovalues act as “data points” on the same scale as the parameter estimate.

Tukey further reasoned that if pseudovalues act like data points, their use may be generalized to include variance estimation: $\hat{V}_{\text{Jack}}(\hat{\theta}) = n^{-1}(n-1) \sum_{i=1}^n (\hat{\theta}_{(i)} - \hat{\theta}_{(\bullet)})^2$, or, in terms of pseudovalues, $\hat{V}_{\text{Jack}}(\hat{\theta}) = \{n(n-1)\}^{-1} \sum_{i=1}^n (\tilde{\theta}_i - \hat{\theta}_{\text{corrected}})^2$, the latter formula reminiscent of the case of the variance estimator of the sample mean, $\hat{V}(\bar{X}) = \{n(n-1)\}^{-1} \sum_{i=1}^n (X_i - \bar{X})^2 = s^2/n$. Uses of the jackknife beyond bias correction and variance estimation have been disappointing. For example, its use in constructing confidence intervals (see **Confidence Intervals**) has been limited and generally resulted in intervals with poor coverage properties.

Permutation Tests

Permutation tests are the oldest resampling method, introduced in the 1930s by Fisher [1] but essentially shelved until recent times when computing power became sufficient to allow their regular use. Suppose that we have data $X_1, X_2, \dots, X_n$ and $Y_1, Y_2, \dots, Y_m$, m not necessarily equal to n . We wish to test whether X and Y arise from the same distribution (so if F is the distribution function of X , and G is the distribution of Y , we want to test $F = G$). Under the hypothesis that F and G are the same, then we can think of $X_1, X_2, \dots, X_n$ and $Y_1, Y_2, \dots, Y_m$ as a single sample of size $n+m$ from the common distribution. The idea is to sample from all $n+m$ data points at random without replacement, allocating the first n points to an “ X ” sample, and the remaining m points to a “ Y ” sample. This process simulates repeated sampling *under the assumption that the parent distributions are the same*. Now for each permutation, compute the value of some test statistic which measures disparity between the distributions of X and Y . Obviously, there are many possible choices for such a statistic, for example, $\bar{X} - \bar{Y}$, (close to 0 if $F = G$), $\bar{X}/\bar{Y}$, (close to 1 if $F = G$), or $\text{median}(X) - \text{median}(Y)$. Extreme values of the test statistic support the hypothesis that the two distributions differ, and “extreme” is defined in terms of the null distribution of the test statistic. The null distribution is generated using the simulated “repeated sampling” – that is, the permutations. For each permutation, the value of the chosen test statistic is calculated, and the Actual Significance Level (ASL), defined as the proportion of permutation values not exceeding the observed test statistic, is computed. A two-sided α -level test rejects the null hypothesis if $\min\{\text{ASL}, (1 - \text{ASL})\} < \alpha/2$.

The Bootstrap

The bootstrap, introduced by Efron in 1979 [4], has become almost synonymous with the term *resampling*, mainly because of its wide applicability and the flexibility it offers in allowing practitioners to avoid restrictive parametric assumptions in their analyses. Efron’s genius lay in his realization that the simple idea of using the empirical distribution function in place of the unknown population distribution function in estimating a parameter of interest could be implemented through a computationally intensive scheme

of sampling from the data itself. Efron also coined the term *bootstrap*, perhaps referring to the legendary Baron von Munchausen of German folklore, who was said to have saved himself from drowning by hoisting himself up by his own bootstraps! Many statistics of interest are functionals of the underlying CDF (see **Cumulative Distribution Function (CDF)**): $\theta = \theta(F)$. For example, the mean, $\mu = \int x dF$. A natural “plug-in” estimator of θ replaces F by the empirical CDF $\hat{F}(x) = n^{-1}\{\#(X_i \leq x)\}$ for data $X_1, \dots, X_n$, the proportion of observed data not exceeding x . The “bootstrap estimator” of θ is $\hat{\theta} = \theta(\hat{F})$, the same functional of $\hat{F}$. For independent and identically distributed (i.i.d.) data and under mild conditions, $|\hat{F}(x) - F(x)| = O_p(n^{-1/2})$, and $\hat{F}(x)$ estimates the true underlying distribution F well (in “normal deviations” regions) even for small values of n . Informally, because $\hat{F} \rightarrow F$, provided θ is a reasonably smooth functional, then $\hat{\theta} \rightarrow \theta$. The bootstrap paradigm involves sampling from $\hat{F}(x)$ in order to simulate repeated sampling from F . Resamples $X_1^*, \dots, X_n^*$ are drawn from $X_1, \dots, X_n$ using the rule $P(X_i^* = X_j) = 1/n$, for all i, j . Bootstrap resampling involves sampling *with* replacement from n equally likely data points. By contrast, permutation tests use sampling without replacement. Each resample, conditional on the data, has distribution function $\hat{F}$, so “repeated sampling” versions of $\hat{\theta} = \theta(\hat{F}) = \theta(X_1, \dots, X_n)$ are $\hat{\theta}^* = \theta(X_1^*, \dots, X_n^*)$, and the bootstrap distribution of $\hat{\theta}$ (the empirical distribution of the $\hat{\theta}^*$ values, conditional on the original data) estimates the sampling distribution of $\hat{\theta}$.

The nature of bootstrap resampling makes clear the importance of two assumptions: in resampling,

the points sampled can appear *in any order* and *in roughly equal proportions*. *In any order* suggests that independence (at least exchangeability) is an important assumption; *in roughly equal proportions* suggests that the data arise from a common distribution. As a result, the basic assumption is that the data is i.i.d. Both assumptions can be relaxed, but doing so requires significant modification to the bootstrap algorithm to allow for resampling independent “units” with like structure.

In principle, “repeated sampling” in the traditional sense reflects an *ad infinitum* thought experiment. But the finite nature of the data means that bootstrapping cannot be conducted this way. The total number of possible distinct resamples is $\binom{2n-1}{n}$ (assuming distinct data), but exhaustive bootstrap sampling is rarely used. In practice, a random sample of resamples is drawn, and the number of resamples required for reasonable performance varies depending on the application. Hundreds of resamples are typically sufficient for point estimation (bias, variance), while thousands of resamples are necessary for interval estimation where tail information is needed.

Figure 1 illustrates the bootstrap paradigm, the idea that the constructed relationship relating the sample and resample “worlds” can be effectively transplanted to allow inferences about model parameters, thereby eliminating the need for classical distributional assumptions through the use of computational effort.

The bootstrap has found wide usage for bias and variance estimation and in the construction of confidence intervals and hypothesis tests (see **Hypothesis Testing**). It has also been applied much

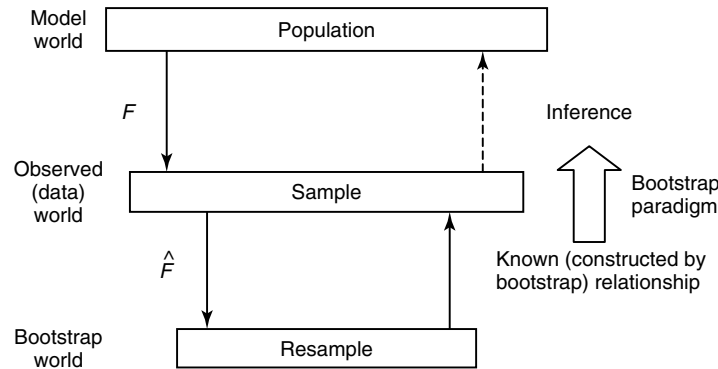


Figure 1 A schematic diagram for the bootstrap

more broadly, for example to setting confidence bands for curvilinear relationships using nonparametric curve estimators, and to problems such as spatial kriging. In a reliability context, bootstrap methods have been employed for setting tolerance intervals for process capability indices (*see Process Capability Indices, Nonnormal*; and Chernick [5, p. 134]) for a discussion of this and other uses in this area, with bootstrap confidence interval construction a key element. Here, we will describe the major uses of the bootstrap, leaving specific discussion to related, more specific articles (*see Resampling for Lifetime Distributions*).

Bootstrap Bias and Variance Estimation

The bias of an estimator $\hat{\theta}$, $b(\hat{\theta}) = E_F(\hat{\theta}) - \theta = E_F\{\theta(\hat{F})\} - \theta(F)$ is a functional of the underlying distribution and of the empirical CDF. Under the bootstrap paradigm, replacing population quantities by sample quantities and sample quantities by resample quantities,

$$\begin{aligned}\hat{b}_{\text{Boot}}(\hat{\theta}) &= E_{\hat{F}}(\hat{\theta}^*) - \hat{\theta} \\ &= E_{\hat{F}}\{\theta(\hat{F}^*)\} - \theta(\hat{F}) \\ &= E_{\hat{F}}\{\theta(X_1^*, \dots, X_n^*)\} - \theta(X_1, \dots, X_n)\end{aligned}\quad (2)$$

where $\hat{F}^*$ is the empirical CDF of a resample $(X_1^*, \dots, X_n^*)$. The bootstrap bias formula involves the quantity $E_{\hat{F}}(\hat{\theta}^*)$, which could be calculated exactly using all possible resamples. In practical terms, we resort to a Monte Carlo algorithm (*see Monte Carlo Methods*) to approximate it: resample $\{X_1^*, \dots, X_n^*\}$ B times from $\{X_1, \dots, X_n\}$ and compute $\hat{\theta}^* = \hat{\theta}(X_1^*, \dots, X_n^*)$ to obtain B values $\hat{\theta}_1^*, \dots, \hat{\theta}_B^*$. The approximation for $E_{\hat{F}}(\hat{\theta}^*)$ is the average $\hat{E}_{\hat{F}}(\hat{\theta}^*) = B^{-1} \sum_{b=1}^B \hat{\theta}_b^*$. The bootstrap bias estimate is $B^{-1} \sum_{b=1}^B \hat{\theta}_b^* - \hat{\theta}$ and the bootstrap bias-corrected estimate is $2\hat{\theta} - B^{-1} \sum_{b=1}^B \hat{\theta}_b^*$, which typically has bias an order of magnitude lower than that of $\hat{\theta}$.

For variance estimation, note that the variance of $\hat{\theta}$ is $V(\hat{\theta}) = E_F(\hat{\theta}^2) - \{E_F(\hat{\theta})\}^2$. The bootstrap estimator of $V(\hat{\theta})$ is thus $\hat{V}_{\text{Boot}}(\hat{\theta}) = E_{\hat{F}}\{(\hat{\theta}^*)^2\} - \{E_{\hat{F}}(\hat{\theta}^*)\}^2$. A Monte Carlo approach can be employed using “bootstrap values” $\hat{\theta}_1^*, \hat{\theta}_2^*, \dots, \hat{\theta}_B^*$

to approximate expectations under $\hat{F}$. Then, $\hat{V}_{\text{Boot}}(\hat{\theta}) = B^{-1} \sum_{b=1}^B \{\hat{\theta}_b^*\}^2 - \left(B^{-1} \sum_{b=1}^B \hat{\theta}_b^*\right)^2 = B^{-1} \sum_{b=1}^B (\hat{\theta}_b^* - \bar{\theta}^*)^2$, where $\bar{\theta}^* = B^{-1} \sum_{b=1}^B \hat{\theta}_b^*$. Note that $\hat{V}_{\text{Boot}}(\hat{\theta})$ is biased for the true variance as it uses divisor B rather than $B-1$. An unbiased, bootstrap variance estimator is $\hat{V}_{\text{Boot}}(\hat{\theta}) = (B-1)^{-1} \sum_{b=1}^B (\hat{\theta}_b^* - \bar{\theta}^*)^2$, the sample variance of the bootstrap quantities $\hat{\theta}_1^*, \hat{\theta}_2^*, \dots, \hat{\theta}_B^*$.

Bootstrap Confidence Intervals

One of the major uses of the bootstrap is confidence interval construction (*see Confidence Intervals*). The standard $100(1-\alpha)\%$ normal theory confidence interval for a parameter θ , $\hat{\theta} \pm z_{1-\alpha/2} \text{s.e.}(\hat{\theta})$, makes the assumption that the standardized quantity $(\hat{\theta} - \theta) / \text{s.e.}(\hat{\theta})$ has a standard normal distribution (*see Normal Distribution*). Our description of bootstrap confidence intervals is based on the concept of pivots, quantities whose (asymptotic) distributions do not depend on unknown parameters. The various bootstrap methods replace the standard normal assumption by bootstrap estimates of the sampling distributions of various “pivots”. We will describe several kinds of bootstrap confidence intervals, each based on different approximate “pivots”: percentile method, percentile- t method (or bootstrap- t), and Efron’s accelerated bias-corrected BC_a method.

The percentile method is based on the premise that the bootstrap distribution of $\hat{\theta}^*$ (that is, the distribution of $\hat{\theta}^*$ under repeated sampling from $\hat{F}$) is close to the distribution of $\hat{\theta}$ under F . **Quantiles** of the bootstrap distribution of $\hat{\theta}^*$ are then used as interval endpoints: if $\hat{G}$ represents the CDF of the bootstrap distribution of $\hat{\theta}^*$, then the interval $[\hat{G}^{-1}(\alpha/2), \hat{G}^{-1}(1-\alpha/2)]$ is a $100(1-\alpha)\%$ percentile method interval for θ . In principle, the percentile interval performs no better than the standard interval, at least in terms of coverage accuracy. It does, however, allow for asymmetry, which means its shape may be more appropriate than that of the standard interval. If there exists a monotone transformation g of $\hat{\theta}$ for which $\hat{\phi} = g(\hat{\theta}) \sim N(\phi, \tau^2)$, where $\phi = g(\theta)$ and τ is a known constant, then the percentile method works well. In this case, we can think of the percentile method as being based on the pivot $(\hat{\phi} - \phi) / \tau$. Moreover, the percentile method is transformation respecting, so the practitioner need

not know what the right transformation is, just that one exists. This property is a major improvement on the standard interval. An excellent example where this phenomenon applies, at least approximately, is the correlation coefficient (see **Correlation**), where Fisher's $\tanh^{-1}$ function effectively normalizes the correlation to a constant variance scale.

To construct a percentile method interval for a parameter θ given data $X_1, \dots, X_n$ and an estimator $\hat{\theta}$ of θ , resample B times $\{X_1^*, \dots, X_n^*\}$ from $\{X_1, \dots, X_n\}$ each time computing $\hat{\theta}^* = \hat{\theta}(X_1^*, \dots, X_n^*)$ to obtain B values $\hat{\theta}_1^*, \dots, \hat{\theta}_B^*$. The $100(1 - \alpha)\%$ equal-tailed percentile interval for θ is $[\hat{\theta}_{[\alpha B/2+1]}^*, \hat{\theta}_{[(1-\alpha)B/2+1]}^*]$, where $\hat{\theta}_{[1]}^* \leq \dots \leq \hat{\theta}_{[B]}^*$ denote the ordered $\hat{\theta}^*$ values.

The percentile- t method is based on the premise that the bootstrap distribution of $(\hat{\theta}^* - \hat{\theta})/s^*$ is close to the distribution of $(\hat{\theta} - \theta)/s$, where $s = s(X_1, \dots, X_n)$ is an estimate of the **standard error** of $\hat{\theta}$ and $s^* = s(X_1^*, \dots, X_n^*)$ is its bootstrap version. Quantiles of the bootstrap distribution of $(\hat{\theta}^* - \hat{\theta})/s^*$ are then used to construct interval endpoints. Formally, if $\hat{K}$ represents the CDF of the bootstrap distribution of $(\hat{\theta}^* - \hat{\theta})/s^*$, then the interval $[\hat{\theta} - s\hat{K}^{-1}(1 - \alpha/2), \hat{\theta} - s\hat{K}^{-1}(\alpha/2)]$ is a $100(1 - \alpha)\%$ percentile- t method interval for θ .

The percentile- t interval performs well provided an adequate estimate of the standard error (see **Standard Error**) of $\hat{\theta}$ is available. Nevertheless, the interval endpoints produced by the percentile- t method are *second-order correct*, meaning that they agree with theoretically “correct” endpoints to terms of order $O(n^{-1})$. In contrast, the percentile method interval is only first-order correct, as is the classical “standard interval”.

If the scale estimate s is not stable enough, especially in small samples, the percentile- t method can go awry. For example, in the case of the correlation coefficient, percentile- t intervals can extend well outside the $[-1, 1]$ range of the parameter. Unfortunately, like the standard interval, the percentile- t interval is not range respecting. It is also not transformation respecting, so it does not share the advantage of the percentile interval in automatically benefiting from the existence of an appropriate normalizing transformation.

To construct a percentile- t method interval for a parameter θ given data $X_1, \dots, X_n$ and an estimator $\hat{\theta}$ of θ and an estimator s of the

standard error of $\hat{\theta}$, resample B times $\{X_1^*, \dots, X_n^*\}$ from $\{X_1, \dots, X_n\}$, each time computing $\hat{\tau}^* = \{\hat{\theta}(X_1^*, \dots, X_n^*) - \hat{\theta}(X_1, \dots, X_n)\}/(s(X_1^*, \dots, X_n^*))$ to obtain B values $\hat{\tau}_1^*, \hat{\tau}_2^*, \dots, \hat{\tau}_B^*$. The $100(1 - \alpha)\%$ equal-tailed percentile- t interval for θ is $[\hat{\theta} - s\hat{\tau}_{[(1-\alpha)B/2+1]}^*, \hat{\theta} - s\hat{\tau}_{[\alpha B/2+1]}^*]$, where $\hat{\tau}_{[1]}^* \leq \hat{\tau}_{[2]}^* \leq \dots \leq \hat{\tau}_{[B]}^*$ denote the ordered $\hat{\tau}^*$ values.

BC_a intervals modify percentile method intervals to improve their coverage accuracy. The percentile method tacitly assumes that there exists a monotone transformation g of $\hat{\theta}$ for which $\hat{\phi} = g(\hat{\theta}) \sim N(\phi, \tau^2)$, where $\phi = g(\theta)$ and τ is a known constant. The BC_a interval assumes that a normalizing transformation g exists, but that the resultant model is potentially biased and has nonconstant variance. Instead, it allows that $(\hat{\phi} - \phi)/(\tau(1 + a\phi)) \sim N(-z_0, 1)$ for constants z_0 (biasing constant) and a (acceleration constant, called so because it measures the rate of change of the standard error with respect to the normalized parameter). The BC_a method adjusts the quantiles of the distribution of $\hat{\theta}^*$ to account for the bias and accelerated variance of this normal approximation. The resultant BC_a interval is

$$\left[\hat{G}^{-1} \left(\Phi \left\{ \hat{z}_0 + (1 - \hat{a}(\hat{z}_0 + z_{\alpha/2}))^{-1} (\hat{z}_0 + z_{\alpha/2}) \right\} \right), \right. \\ \left. \hat{G}^{-1} \left(\Phi \left\{ \hat{z}_0 + (1 - \hat{a}(\hat{z}_0 + z_{(1-\alpha)/2}))^{-1} \right. \right. \right. \\ \left. \left. \left. \times (\hat{z}_0 + z_{(1-\alpha)/2}) \right\} \right) \right] \quad (3)$$

Efron [6] developed the BC_a interval as an adjustment to the percentile interval to allow for a more flexible model for the distribution of $(\hat{\phi} - \phi)/\tau$. Thought of in this way, the BC_a method is an attempt to approximate the distribution of this “pivot” better than the percentile method does by allowing the limiting distribution to incorporate appropriate bias and variance terms. A key part of implementing the BC_a interval is the estimation of the biasing and acceleration constants. The biasing constant is estimated as $\hat{z}_0 = \Phi^{-1}\{\hat{G}(\hat{\theta})\}$, while the acceleration constant is estimated as $\hat{a} = \sum_{i=1}^n (\hat{\theta}_{(i)} - \hat{\theta}_{(i)})^3 / [6\{\sum_{i=1}^n (\hat{\theta}_{(i)} - \hat{\theta}_{(i)})^2\}^{3/2}]$. DiCiccio and Efron [7] developed an analytical approximation to BC_a intervals, called the *ABC method*, designed to avoid the need for significant resampling. The algorithm replaces bootstrap quantities in the BC_a algorithm

by Taylor series approximations. The ABC method shares the advantages of the BC_a method, but can be computed in a small fraction of the time that BC_a intervals require.

Use of the bootstrap has also been proposed in a variety of contexts beyond those already described here. Bootstrap **hypothesis testing** has been developed both as an extension of bootstrap confidence interval methods and also in specific contexts where resampling under relevant hypotheses has been developed. A particularly important use of the bootstrap has been in the contexts of regression and time series models. For regression models, two main approaches have been developed: residual-based resampling, where i.i.d. model residuals (*see Residuals*) are resampled and the fitted model is added back to the resampled residuals to create resampled data; and case-based resampling, where individual cases (**covariate**/response pairs) are resampled under the notion that they represent i.i.d. variates drawn from an appropriate multivariate distribution. Time series settings pose particularly difficult questions for the bootstrap, as serial dependence (*see Dependence*) in data contravenes the exchangeability assumption required by the basic bootstrap. Again, two approaches have been developed to deal with this problem, the first based on resampling residuals from an appropriate Autoregressive Integrated Moving Average (ARIMA) (*see Large Autoregressive Integrated Moving Average (ARIMA) Modeling*), or other, model, and the second based on resampling “blocks” of contiguous data points under the notion that while data may suffer short-range dependence, data points sufficiently far apart (that is, in different blocks) should be close to independent.

Where to Look for More Information

In an article of this length, it is only possible to skim the surface of a topic as broad as resampling, and only a very limited set of references is possible. A number of excellent monographs provide much additional detail. Efron and Tibshirani [8] is an excellent starting point, providing a very good, accessible account of the bootstrap in practice. Davison and Hinkley [9] gives a very thorough overview of resampling in a very broad set of contexts. More theoretical

accounts are given by Shao and Tu [10] and Hall [11], while Chernick [5] gives an exhaustive bootstrap bibliography with about 1600 references. Young [12] offers an excellent article outlining the achievements of, and the challenges for, the bootstrap.

References

- [1] Fisher, R.A. (1935). *The Design of Experiments*, Hafner Publications, New York.
- [2] Quenouille, M. (1949). Approximate tests of correlation in time series, *Journal of the Royal Statistical Society. Series B* **11**, 18–84.
- [3] Tukey, J.W.T. (1958). Bias and confidence in not-quite large samples, *Annals of Mathematical Statistics* **29**, 614.
- [4] Efron, B. (1979). Bootstrap methods: another look at the jackknife, *Annals of Statistics* **7**, 1–26.
- [5] Chernick, M. (1999). *Bootstrap Methods: A Practitioner's Guide*, John Wiley & Sons, New York.
- [6] Efron, B. (1987). Better bootstrap confidence intervals (with discussion), *Journal of the American Statistical Association* **82**, 171–200.
- [7] DiCiccio, T.J. & Efron, B. (1996). Bootstrap confidence intervals (with discussion), *Statistical Science* **11**, 189–228.
- [8] Efron, B. & Tibshirani, R. (1993). *An Introduction to the Bootstrap*, Chapman & Hall, New York.
- [9] Davison, A.C. & Hinkley, D.V. (1997). *Bootstrap Methods and Their Application*, Cambridge University Press.
- [10] Shao, J. & Tu, D. (1995). *The Jackknife and Bootstrap*, Springer, New York.
- [11] Hall, P. (1992). *The Bootstrap and Edgeworth Expansion*, Springer, New York.
- [12] Young, G.A. (1994). Bootstrap: more than a stab in the dark (with discussion), *Statistical Science* **9**, 382–415.

Related Articles

Bias of an Estimator; Confidence Intervals; Correlation; Cumulative Distribution Function (CDF); Dependence; Hypothesis Testing; Large Autoregressive Integrated Moving Average (ARIMA) Modeling; Monte Carlo Methods; Normal Distribution; Process Capability Indices, Nonnormal; Resampling for Lifetime Distributions; Residuals; Sampling from Virtual Populations; Standard Error; Time Series Analysis; Variance.

MICHAEL A. MARTIN AND STEVEN ROBERTS

Asymptotic Reliability Analysis of Very Large Structural Systems

Introduction

Uncertainties need to be taken into account for optimal and reliable design of complex safety-critical structural systems such as aircraft, helicopters, and marine vehicles. Such uncertainties include, but are not limited to, uncertainties in specifying material properties, geometric parameters, boundary conditions, and applied loadings. Other articles in this encyclopedia that deal with related problems are **Fault Tree Analysis for Large Systems**, **Quadrature and Numerical Integration**, and **Normal Distribution**, respectively. Suppose that the random variables describing the uncertainties in the structural properties and loading are considered to form a vector $\mathbf{x} \in \mathbb{R}^n$ where n is the number of random variables. The statistical properties of the system are fully described by the joint **probability density function** $p(\mathbf{x}) : \mathbb{R}^n \mapsto \mathbb{R}$. For a given set of variables $\mathbf{x}$, the structure will either fail under the applied (random) loading or be safe. The condition of the structure for every $\mathbf{x}$ can be described by a safety margin $g(\mathbf{x}) : \mathbb{R}^n \mapsto \mathbb{R}$ such that the structure fails if $g(\mathbf{x}) \leq 0$ and is safe if $g(\mathbf{x}) > 0$. Thus, the probability of failure is given by

$$P_f = \int_{g(\mathbf{x}) \leq 0} p(\mathbf{x}) d\mathbf{x} \quad (1)$$

The function $g(\mathbf{x})$ is also known as the *limit-state function* and the $(n - 1)$ -dimensional surface $g(\mathbf{x}) = 0$ is known as the *failure surface*. The central theme of a structural reliability analysis is to evaluate the multidimensional integral (1). The exact evaluation of this integral, either analytically or numerically, is not possible for most practical problems because n is usually large and $g(\mathbf{x})$ is a highly nonlinear function of $\mathbf{x}$ which may not be available explicitly. Over the past three decades, there has been extensive research (see, e.g., the books [1–4]) to develop approximate numerical methods for the efficient calculation of the reliability integral. The approximate reliability methods can be broadly grouped into (a) first-order reliability method

(FORM) and (b) second-order reliability method (SORM). In FORM and SORM, it is assumed that all the basic random variables are transformed and scaled so that they are uncorrelated Gaussian random variables, each with zero mean and unit standard deviation.

The difficulty in computing the failure probability using classical FORM and SORM increases rapidly with the number of variables or “dimension”. There are two main reasons behind this. The first is the increase in computational time with the increase in the number of random variables. In principle, this problem can be handled with superior computational tools and powerful computing machines. The second, which is perhaps more fundamental, is that there are some conceptual difficulties associated typically with high dimensions. In the context of FORM, using the Tchebysheff bound, Veneziano [5] has shown that the probability of failure depends on the dimension n although the reliability index does not explicitly depend on n . In the context of SORM, using the χ^2 distribution, Fiessler *et al.* [6] have shown that there can be a significant difference in the probability of failure in higher dimensions for a fixed value of the reliability index. This means that even if one manages to carry out the necessary computations, the application of existing FORM and SORM may still lead to incorrect results in high dimensions. In this chapter, two asymptotic approximations are derived for the case when the number of random variables $n \rightarrow \infty$.

Asymptotic Distribution of Quadratic Forms

We suppose that the random variables are scaled so that $\mathbf{x}$ is a standard n -dimensional Gaussian random vector. Using classical SORM (see, e.g., [2, 7, 8]), the failure surface can be approximated as

$$\tilde{g} \approx -y_n + \beta + \mathbf{y}^T \mathbf{A} \mathbf{y} \quad (2)$$

Here $\mathbf{y}$ is an $(n - 1)$ -dimensional standard Gaussian random vector, y_n is a standard Gaussian random variable, and $\mathbf{A}$ is a symmetric, positive definite matrix. The scalar β is known as the *reliability index* [9], which is the minimum distance of the failure surface from the origin in $\mathbb{R}^n$. The eigenvalues of $\mathbf{A}$ can be related to the principal curvatures of the failure surface at the design point. With the

2 Asymptotic Reliability Analysis of Very Large Structural Systems

approximation in equation (2), the failure probability is given by

$$\begin{aligned} P_f &\approx \text{Prob}[\tilde{g} \leq 0] \approx \text{Prob}[y_n \geq \beta + \mathbf{y}^T \mathbf{A} \mathbf{y}] \\ &= \text{Prob}[y_n \geq \beta + U] \end{aligned} \quad (3)$$

where

$$U : \mathbb{R}^{n-1} \mapsto \mathbb{R} = \mathbf{y}^T \mathbf{A} \mathbf{y} \quad (4)$$

is a central quadratic form in standard Gaussian random variables.

Using quadratic approximation of the failure surface together with asymptotic analysis, Breitung [10] proved that

$$\begin{aligned} P_f &\longrightarrow \Phi(-\beta) \|\mathbf{I}_{n-1} + 2\beta \mathbf{A}\|^{-1/2} \\ \text{when } \beta &\longrightarrow \infty \end{aligned} \quad (5)$$

Here $\Phi(\bullet)$ is the standard Gaussian cumulative distribution function. Later, Hohenbichler and Rackwitz [11] proposed the following formula

$$P_f \approx \Phi(-\beta) \left\| \mathbf{I}_{n-1} + 2 \frac{\varphi(\beta)}{\Phi(-\beta)} \mathbf{A} \right\|^{-1/2} \quad (6)$$

where $\varphi(\bullet)$ is the standard Gaussian probability density function. This expression is more accurate than equation (5) for lower values of β (although both are asymptotically equivalent). In FORM, the failure surface is approximated by a hyperplane at the design point and consequently the probability of failure is approximated as

$$P_f \approx \Phi(-\beta) \quad (7)$$

This is the simplest approximation to the integral (1). Breitung's formula (5) and the formula by Hohenbichler and Rackwitz (6) can be viewed as corrections to the FORM formula (7) to take into account the curvature of the failure surface at the design point.

If n is very large, the computation of P_f using any available methods will be computationally expensive. Nevertheless, it is of interest to know if we can expect the same level of accuracy from the classical FORM/SORM formula in high dimensions as we do in low dimensions. Moreover, what we exactly mean by "high dimension" is also an open question. We tried to address these issues using the asymptotic distribution of the quadratic form (4) as $n \rightarrow \infty$.

Overview of SORM Approximations

Discussions on asymptotic distribution of quadratic forms may be found in [12, Section 4.6b]. The moment generating function of U is given by

$$\begin{aligned} M_U(s) &= E[e^{sU}] = E[e^{s\mathbf{y}^T \mathbf{A} \mathbf{y}}] \\ &= \|\mathbf{I}_{n-1} - 2s\mathbf{A}\|^{-1/2} \end{aligned} \quad (8)$$

Suppose n is large such that

$$\begin{aligned} \sum_{k=1}^{n-1} \frac{(2a_k/\sqrt{n})^2}{2} &< \infty \\ \text{or } \frac{2}{n} \text{Trace}(\mathbf{A}^2) &< \infty \end{aligned} \quad (9)$$

$$\begin{aligned} \text{and } \sum_{k=1}^{n-1} \frac{(2a_k/\sqrt{n})^r}{r} &\longrightarrow 0 \\ \text{or } \frac{2^r}{n^{r/2}} \text{Trace}(\mathbf{A}^r) &\longrightarrow 0, \quad \forall \quad r \geq 3 \end{aligned} \quad (10)$$

Under these assumptions, using the series expansion of the logarithm of the moment generating function, it can be shown that U asymptotically approaches a Gaussian random variable with mean $\text{Trace}(\mathbf{A})$ and variance $2 \text{Trace}(\mathbf{A}^2)$; that is, when $n \rightarrow \infty$

$$\begin{aligned} U &\simeq \mathbb{N}_1(m, \sigma^2), \quad \text{with } m = \text{Trace}(\mathbf{A}) \\ \text{and } \sigma &= \sqrt{2 \text{Trace}(\mathbf{A}^2)} \end{aligned} \quad (11)$$

Suppose ϵ is the allowable error for the asymptotic approximation. Using the error in neglecting more than quadratic order terms in the Taylor series of $\ln(M_q(s))$, it can be shown that we should have

$$\begin{aligned} \left| \frac{1}{r} \frac{(-2\beta)^r}{n^{r/2}} \text{Trace}(\mathbf{A}^r) \right| &< \epsilon \\ \text{or } n^{r/2} &> \frac{(2\beta)^r}{r\epsilon} \text{Trace}(\mathbf{A}^r) \end{aligned} \quad (12)$$

$$\text{or } n_{\min} = \frac{4\beta^2}{\sqrt[3]{9\epsilon^2}} \left(\sqrt[3]{\text{Trace}(\mathbf{A}^3)} \right)^2 \quad (13)$$

Since $\mathbf{A}$ is a positive definite matrix, the critical value of n is obtained for $r = 3$. From equation (13), the following points may be observed: (a) the minimum number of random variables required would be more if ϵ (error) is considered to be small (as expected)

and $n_{\min} \propto \frac{1}{\epsilon^{2/3}}$ and (b) if β is large, more random variables are needed to achieve a desired accuracy and $n_{\min} \propto \beta^2$.

Failure Probability Using Strict Asymptotic Formulation

In the strict asymptotic formulation we start with equation (3). The probability of failure can be rewritten as

$$P_f \approx \text{Prob}[y_n \geq \beta + U] = \text{Prob}[y_n - U \geq \beta] \quad (14)$$

Since the asymptotic pdf of U is Gaussian, the variable $z = y_n - U$ is also a Gaussian random variable with mean $(-\text{Trace}(\mathbf{A}))$ and variance $(1 + 2 \text{Trace}(\mathbf{A}^2))$. The probability of failure can be obtained from equation (14) as

$$P_{f\text{Strict}} \longrightarrow \Phi(-\beta_1), \quad \beta_1 = \frac{\beta + \text{Trace}(\mathbf{A})}{\sqrt{1 + 2 \text{Trace}(\mathbf{A}^2)}} \quad (15)$$

when $n \longrightarrow \infty$

This is the exact asymptotic expression of P_f after making the parabolic failure surface assumption and it cannot be improved or changed asymptotically. If the failure surface is close to linear and the number of random variables is not very large, then it is expected that $\text{Trace}(\mathbf{A}) = \text{Trace}(\mathbf{A}^2) \rightarrow 0$, and it is easy to see that equation (15) reduces to the classical FORM formula (7). Therefore, the expressions derived here can be viewed as the “correction” that needs to be applied to the classical FORM formula when a large number of random variables are considered.

The value of n given by equation (13) can be viewed as the borderline between the low and high dimension. Beyond this value of n , significant “trace effect” can be observed and consequently the modified reliability index β_1 should be used instead of β .

Failure Probability Using Weak Asymptotic Formulation

The expression of P_f given by equation (15) cannot be improved asymptotically. However, there is scope for improvements if one does not strictly apply the asymptotic condition $n \rightarrow \infty$. The advantage of such nonasymptotic approximations is that

the approximations may work well even when the asymptotic condition is not met. The disadvantage is that a nonasymptotic approximation will have unquantified errors and one cannot, in general, prove that such errors will vanish when the asymptotic condition is fulfilled. Nevertheless, it is worth perusing a nonasymptotic approximation since real-life structural systems have a finite number of random variables.

Rewriting equation (3), the failure probability can be expressed as

$$P_f \approx \text{Prob}[y_n \geq \beta + U] = \int_{\mathbb{R}} \left\{ \int_{\beta+u}^{\infty} \varphi(y_n) dy_n \right\} \times p_U(u) du = E[\Phi(-\beta - U)] \quad (16)$$

where $p_U(u)$ is the probability density function of U and $E[\bullet]$ is the expectation operator. From the definition of U in equation (4), note that $u \in \mathbb{R}^+$ since $\mathbf{A}$ is a positive definite matrix. We rewrite equation (16) as

$$P_f \approx \int_{\mathbb{R}^+} \Phi(-\beta - u) p_U(u) du = \int_{\mathbb{R}^+} e^{\ln[\Phi(-\beta - u)] + \ln[p_U(u)]} du \quad (17)$$

The aim here is to expand the integrand in a first-order Taylor series about the most probable point or optimal point, say $u = u^*$. The optimal point is the point where the integrand in equation (17) reaches its maxima in $u \in \mathbb{R}^+$. The asymptotic approximation of $p_U(u)$ will be used only to find the maxima of the integrand and will *not* be used subsequently to calculate the expectation. The expectation operation will be carried out exactly by utilizing the expression of the moment generating function in equation (8). For this reason this approach is called *weak asymptotic formulation*.

For the maxima of the integrand in equation (17), we must have

$$\frac{\partial}{\partial u} \{ \ln[\Phi(-\beta - u)] + \ln[p_U(u)] \} = 0 \quad (18)$$

Recalling that

$$p_U(u) = (2\pi)^{-1/2} \sigma^{-1} e^{-(u-m)^2/(2\sigma^2)} \quad (19)$$

4 Asymptotic Reliability Analysis of Very Large Structural Systems

where m and σ are given in equation (11), equation (18) results in

$$\frac{\varphi(\beta + u)}{\Phi(-(\beta + u))} = \frac{m - u}{\sigma^2} \quad (20)$$

Because this relationship holds at the optimal point u^* , we define a constant η as

$$\eta = \frac{\varphi(\beta + u^*)}{\Phi(-(\beta + u^*))} = \frac{m - u^*}{\sigma^2} \quad (21)$$

Taking a first-order Taylor series expansion of $\ln[\Phi(-\beta - u)]$ about $u = u^*$ and simplifying, we have

$$\Phi(-\beta - u) \approx \Phi(-\beta_2)e^{\eta u^*} e^{-\eta u} \quad (22)$$

where

$$\beta_2 = \beta + u^* \quad (23)$$

Taking the expectation of equation (22) and utilizing the expression of the moment generating function in equation (8), the probability of failure can be expressed as

$$P_f \approx \Phi(-\beta_2)e^{\eta u^*} \|\mathbf{I}_{n-1} + 2\eta\mathbf{A}\|^{-1/2} \quad (24)$$

The optimal point u^* should be obtained by solving the nonlinear equation (21). An exact closed-form solution of this equation does not exist. By considering the asymptotic expansion of the ratio $\varphi(\beta + u^*)/\Phi(-(\beta + u^*))$, equation (21) can be solved approximately and β_2 in equation (23) can be obtained as

$$\begin{aligned} \beta_2 &\approx \eta \approx \beta + u^* \approx \beta + \frac{m - u^*}{\sigma^2} \\ &\approx \frac{\beta + m}{1 + \sigma^2} = \frac{\beta + \text{Trace}(\mathbf{A})}{1 + 2 \text{Trace}(\mathbf{A}^2)} \end{aligned} \quad (25)$$

Now substituting η and β_2 in equation (24), the failure probability using weak asymptotic formulation can be obtained as

$$P_{f\text{Weak}} \longrightarrow \frac{\Phi(-\beta_2)e^{-(2\beta_2^2\text{Trace}(\mathbf{A}^2) - \beta_2\text{Trace}(\mathbf{A}))}}{\sqrt{\|\mathbf{I}_{n-1} + 2\beta_2\mathbf{A}\|}}, \quad (26)$$

when $n \longrightarrow \infty$

where β_2 is defined in equation (25). If the number of random variables is not very large, then it is expected that $\text{Trace}(\mathbf{A}) = \text{Trace}(\mathbf{A}^2) \rightarrow 0$. In that case, it is easy to see that $\beta_2 \rightarrow \beta$ and equation (26)

reduces to the classical SORM formula (5). Therefore, equation (26) can be viewed as the “correction” that needs to be applied to the classical SORM formula when a large number of random variables are considered. Unlike the strict formulation, a simple geometric explanation of this expression cannot be given. Also note that the modified reliability indices β_1 and β_2 for the two formulations are not identical.

A Numerical Example

We consider a problem for which the failure surface is *exactly* parabolic in the normalized space, as given by equation (2). In numerical calculations, we have fixed the number of random variables n and the trace of the coefficient matrix $\mathbf{A}$. It is assumed that the eigenvalues of $\mathbf{A}$ are uniform positive random numbers. Figure 1 shows the probability of failure (normalized by dividing with $\Phi(-\beta)$) for values of β ranging from 3 to 6 and values of n ranging from 10 to 1000. Results obtained from the expressions derived here are compared with two existing expressions given in equations (5) and (6). The following points may be noted for these results:

- For a fixed value of β and $\mathbf{A}$, the weak asymptotic formulation is more accurate for smaller values of n compared to the strict asymptotic formulation.
- For a fixed value of $\mathbf{A}$ and n , the results from both approaches will be more accurate if β is small.

Conclusions

The demands of modern engineering design have led structural engineers to model a structure using random variables in order to handle uncertainties. Two approximations to calculate the probability of failure of an engineering structure when the number of random variables used for mathematical modeling, $n \rightarrow \infty$ are provided. It is assumed that the basic random variables are Gaussian and the failure surface is approximated by a parabolic hypersurface in the neighborhood of the design point. The new approximations are based on the asymptotic distribution of a central quadratic form in Gaussian random variables. Two formulations, namely, *strict asymptotic formulation* and *weak asymptotic formulation* are presented. Both approximations result in simple

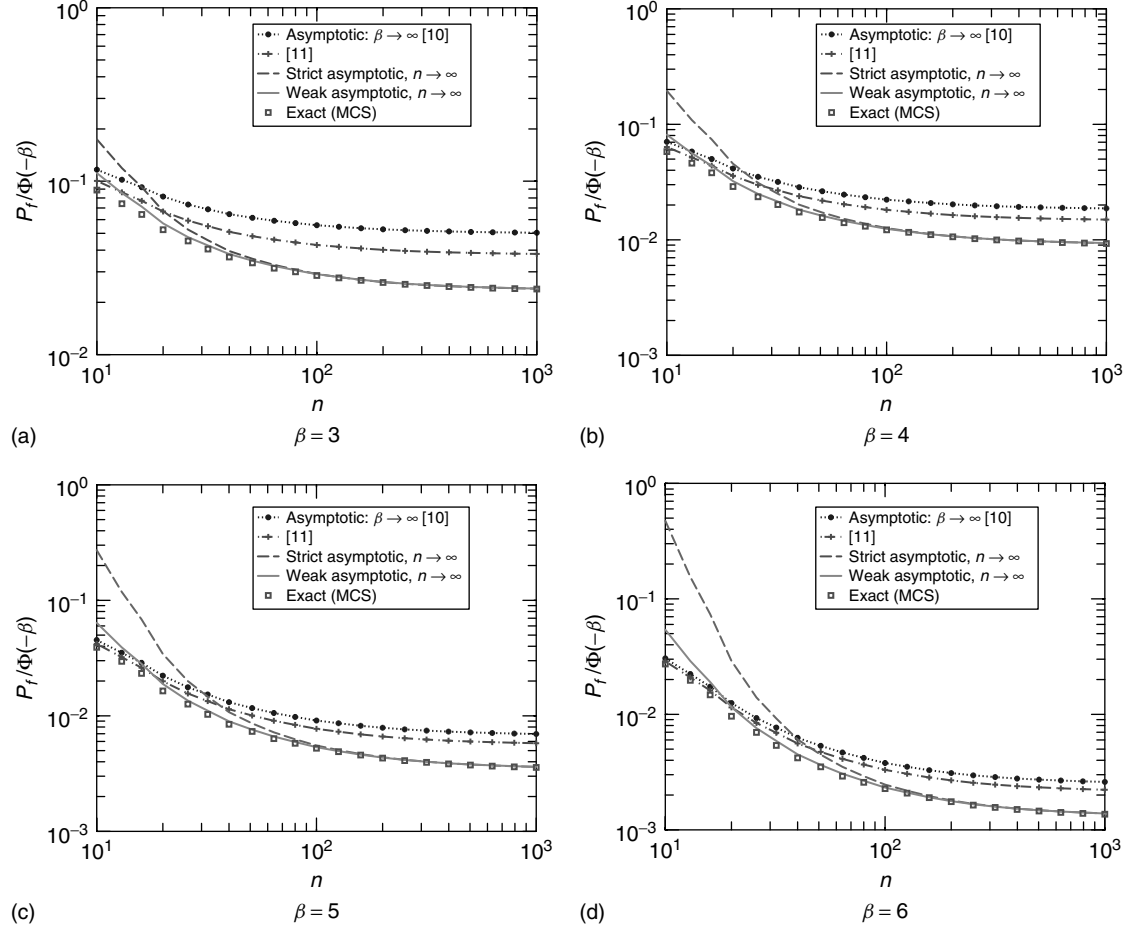


Figure 1 Normalized failure probability for different values of β when $\text{Trace}(\mathbf{A}) = 1$

closed-form expressions in terms of the trace of the curvature matrix. For a small number of variables, the trace effects may not be significant, and in such cases the proposed formulas reduce to the classical FORM and SORM formulas, respectively. A closed-form expression for the minimum number of random variables required to apply these asymptotic formulas is derived. Beyond this value of n , the reliability problem can be considered high dimensional since the trace effects become significant. The proposed approximations are compared with some existing approximations and **Monte Carlo** simulations using numerical examples. The results obtained from both asymptotic formulations match well with the Monte Carlo simulation results in high dimensions. Numerical studies show that the weak formulation is in

general applicable for a small number of variables compared to the strict formulation. In many real-life problems, the number of random variables is expected to be large. In such situations, the asymptotic results derived here will be useful.

References

- [1] Thoft-Christensen, P. & Baker, M.J. (1982). *Structural Reliability Theory and its Applications*, Springer-Verlag, Berlin.
- [2] Madsen, H.O., Krenk, S. & Lind, N.C. (1986). *Methods of Structural Safety*, Prentice-Hall, Englewood Cliffs.
- [3] Ditlevsen, O. & Madsen, H.O. (1996). *Structural Reliability Methods*, John Wiley & Sons, Chichester, January.

- [4] Melchers, R.E. (1999). *Structural Reliability Analysis and Prediction*, 2nd Edition, John Wiley & Sons, Chichester, February.
- [5] Veneziano, D. (1979). New index of reliability, *Journal of the Engineering Mechanics Division ASCE* **105**(EM2), 277–296.
- [6] Fiessler, B., Neumann, H.J. & Rackwitz, R. (1979). Quadratic limit states in structural reliability, *Journal of the Engineering Mechanics Division, Proceedings of the ASCE* **105**(EM4), 661–676.
- [7] Adhikari, S. (2004). Reliability analysis using parabolic failure surface approximation, *Journal of Engineering Mechanics ASCE* **130**(12), 1407–1427.
- [8] Adhikari, S. (2005). Asymptotic distribution method for structural reliability analysis in high dimensions, *Proceedings of the Royal Society of London Series A* **461**(2062), 3141–3158.
- [9] Hasofer, A.M. & Lind, N.C. (1974). Exact and invariant second moment code format, *Journal of the Engineering Mechanics Division, Proceedings of the ASCE* **100**(EM1), 111–121.
- [10] Breitung, K. (1984). Asymptotic approximations for multinormal integrals, *Journal of Engineering Mechanics ASCE* **110**(3), 357–367.
- [11] Hohenbichler, M. & Rackwitz, R. (1988). Improvement of second-order reliability estimates by importance sampling, *Journal of Engineering Mechanics ASCE* **114**(12), 2195–2199.
- [12] Mathai, A.M. & Provost, S.B. (1992). *Quadratic Forms in Random Variables: Theory and Applications*, Marcel Dekker, New York.

SONDIPON ADHIKARI

Perfect Sampling

Introduction

A common requirement of many problems in a variety of disciplines is for a sample to be drawn from some probability distribution π on a space $\mathcal{X}$. Often this distribution cannot be specified completely: for example, it may only be known up to a normalization constant. One widely used method for doing this is Markov chain **Monte Carlo** (MCMC). Such an approach involves designing an ergodic Markov chain X , which has π as its stationary distribution, and then running a simulation of X until it is near equilibrium (see **Markov Chain Monte Carlo, Introduction; Convergence and Mixing in Markov Chain Monte Carlo**).

However, an obvious problem with MCMC is the following: for how long should we run the chain before sampling (in other words, what is an appropriate burn-in time)? If we sample from X when it is not ‘close enough’ to equilibrium then the algorithm output may be from a distribution that is far from π .

An attractive alternative to estimating a sufficient burn-in period is to adapt the MCMC algorithm to form a *perfect simulation* algorithm. Such a procedure has two desirable features:

1. The algorithm determines for itself when it should stop.
2. If the algorithm is successful, then it returns a sample drawn *exactly* from π .

In [1], an algorithm known as *coupling from the past* (CFTP) was introduced and used to sample from the exact equilibrium distribution of the critical Ising model on a finite lattice. Although there do exist other perfect simulation algorithms (such as the FMMR algorithm [2]), in this article we shall principally present the CFTP algorithm, and a variant of this method named *dominated CFTP* (also called *CIAFTP*).

A friendly warning: perfect simulation is a technical area that has mainly evolved in the stochastic process literature, and which involves a lot of stochastic process theory. Unfortunately, there’s no escaping from this fact, but hopefully the simple examples below will help newcomers to this field gain some insight into the beauty of the approach.

Coupling from the Past

As the name suggests, CFTP is based on the idea of coupling Markov chains together (see [3] for a comprehensive introduction). For our purpose, a coupling of two chains X and X' (with common transition kernel) is a prescription for specifying their joint transition kernel such that if the two chains meet, then they stay together thereafter.

A general method of producing such a coupling is to define the updates of X and X' using a common update function. More specifically, it is possible to describe the transitions of X using a deterministic function $f : \mathcal{X} \times [0, 1] \rightarrow \mathcal{X}$ and a sequence of *Uniform*[0, 1] random variables $\{\xi_n\}$, such that

$$X_{n+1} = f(X_n, \xi_n) \quad (1)$$

(This construction is known as a *stochastic recursive sequence*.) Using this representation we can couple X and X' such that they stay together once they agree: define their updates by

$$X_{n+1} = f(X_n, \xi_n), \quad X'_{n+1} = f(X'_n, \xi_n) \quad (2)$$

The beauty of this representation is that it enables us to couple *any number* of chains in this way, using the same f and the same randomness $\{\xi_n\}$ for all the chains.

The basic concept behind CFTP is the following: Consider a hypothetical copy of our chain of interest, say $\tilde{X}$, which has been running since time $-\infty$. At time zero this chain will be in equilibrium, i.e., $\tilde{X}_0 \sim \pi$. This would be of obvious use if we could determine $\tilde{X}_0$, but clearly we cannot run a chain from time $-\infty$ in practice! However, it may be possible to determine $\tilde{X}_0$ by looking back only a *finite number of steps* into the past of $\tilde{X}$.

To describe the algorithm, suppose that we have available to us an independent and identically distributed (i.i.d.) sequence $\{\xi_n\}_{-\infty}^0$ of *Uniform*[0, 1] random variables and a deterministic update function f . For $v \geq -u$, define the random composite *input-output* map $F_{(-u,v]}(x) : \mathcal{X} \rightarrow \mathcal{X}$ by

$$F_{(-u,v]}(x) = f(f(\dots f(f(x, \xi_{-u}), \xi_{-u+1}) \dots, \xi_{v-2}), \xi_{v-1}) \quad (3)$$

Thus if X is begun at time $-u$ with the value $X_{-u} = x$, then we may set $X_v = F_{(-u,v]}(x)$. We write

2 Perfect Sampling

$X_v^{x,-u}$ for the value of the chain X at time v , when X is started at time $-u$ from state x .

Definition 1 For a given update function f , the *backwards coalescence time* T^* of the chain X is the random variable defined by

$$T^* = \min\{n \geq 0 : F_{(-t,0]}(x) = F_{(-s,0]}(y), \\ \text{for all } x, y \in \mathcal{X} \text{ and for all } s, t \geq n\} \quad (4)$$

Now consider starting chains $X^{x,-n}$ from *all* states $x \in \mathcal{X}$ at time $-n$. Using f we can couple all of these chains, as well as the hypothetical $\tilde{X}$, using the same source of randomness $\{\xi_n\}_{n=0}^\infty$. If $n \geq T^*$ then $X_0^{x,-n} =: X_0$ is the same for all x , and so $X_0 = \tilde{X}_0 \sim \pi$. That is, if we start our chains $X^{x,-n}$ far enough back into the past that they have all coalesced by time zero, then their common value at zero is an exact draw from π , as required! Of course, we do not know the value of T^* , and so the CFTP algorithm simply repeats the above procedure for larger and larger n until $n \geq T^*$ is achieved:

CFTP Algorithm

```

set  $n \leftarrow 1$ 
while  $F_{(-n,0]}(\mathcal{X})$  not constant
   $n \leftarrow 2n$ 
return  $F_{(-n,0]}(\mathcal{X})$ 

```

The most important aspect of the algorithm is that the random sequence $\{\xi_n\}$ is *reused* in each execution of the `while` loop. This implies a possible issue with computer memory in practice, but there is a trick that avoids this, known as *read-once CFTP* [4]. The use of binary search for T^* contained within the `while` loop is also noteworthy: we are free to increase n in any way we like, but the binary search means that the total number of simulated Markov chain steps is linear in T^* , whereas if we used $n \leftarrow n + 1$ (for example) it would grow quadratically. Furthermore, this strategy means that the number of simulation steps comes within a factor of 4 of the true value of T^* [1, 4].

Theorem 1 *If the backward coalescence time T^* is almost surely finite, then CFTP samples from equilibrium.*

The proof of Theorem 1 is very simple [1, 5]. Note that if $\mathcal{X}$ is finite, then T^* is almost surely finite (but very large) for the trivial independence coupling, where chains evolve independently until they meet, after which they agree forever. In fact, $\mathbb{P}(T^* < \infty)$ is always either zero or one, whatever the choice of f [6]. A good CFTP algorithm uses a function f that has a high probability of making target chains coalesce quickly.

A Simple Example

The following example was considered in [7], with variants appearing in [1, 8]. Consider a symmetric reflecting random walk X on $\mathcal{X} = \{1, 2, 3, 4\}$, satisfying $\pi(i) = 1/4$ for all $i \in \mathcal{X}$ (see Figure 1).

For our update function f we may choose the following:

$$f(x, u) = \begin{cases} \min(x + 1, 4) & \text{if } u \leq 1/2 \\ \max(x - 1, 1) & \text{if } u > 1/2 \end{cases} \quad (5)$$

Figure 2 demonstrates the CFTP algorithm for this example. The gray arrows indicate the transitions determined by f and the sequence $\{\xi_n\}$. The algorithm runs chains $X^{x,-n}$ started from all states x , with $n = 1, 2, 4, \dots$. When $n = 8$ is reached, all of

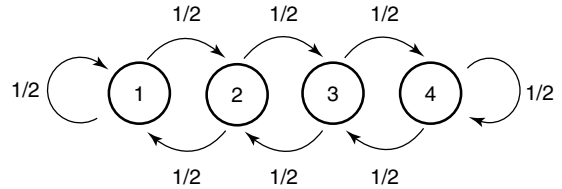


Figure 1 Simple symmetric reflecting random walk

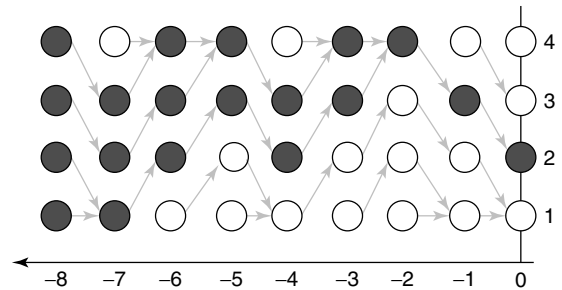


Figure 2 CFTP for a simple symmetric reflecting random walk

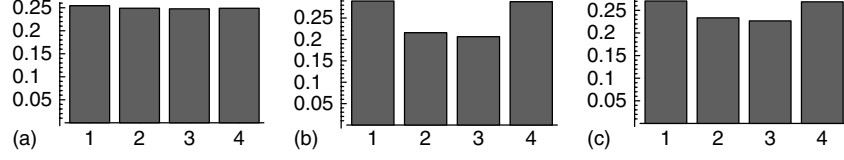


Figure 3 Output of 10 000 runs of the CFTP algorithm for the random walk described above: (a) a proper implementation, reusing randomness; (b) not reusing randomness; (c) reusing randomness, but with simulations with “long” run times interrupted and discarded

the target chains have the same value at time zero: $X_0^{x,-8} = 2$ for this realization (for which $T^* = 7$). The shaded circles highlight the possible values of the chains $X^{x,-8}$ at each step: coalescence occurs the first time that $X^{1,-n}$ hits four or $X^{4,-n}$ hits one.

Implementation of CFTP for this example is trivial, and Figure 3(a) shows a histogram of the results of 10 000 runs of the algorithm. A χ^2 test to compare this output to π yields a p value of 0.79: it is clear that the algorithm is drawing from the correct distribution. In contrast, Figure 3(b) shows the output of a wrongly implemented algorithm, in which the random sequence $\{\xi_n\}$ is not reused when the chains are restarted further into the past. Here a definite departure from uniformity is visible, and indeed a χ^2 test gives a tiny p value of 2×10^{-16} . The reader may like to see if his/her intuition can explain this bias toward a distribution with more mass at the extremal states! For more examples of such bias see [1, 4, 8]. Finally, Figure 3(c) shows another possible source of bias: that which is introduced by *user impatience*. This happens when simulations with long run times are interrupted and discarded. Since the output of the CFTP algorithm is in general not independent of the algorithm run time, this introduces a bias. In this implementation, any simulation that had not coalesced when starting chains from time -8 were discarded. The p value for this sample is again very small: $p = 2.6 \times 10^{-11}$. It should be noted that the FMMR algorithm (under suitable implementation) does not suffer from this bias [9].

Although this is clearly a toy example, it does highlight the ease with which CFTP may be applied to state spaces with a partial order. More specifically, suppose that $\mathcal{X}$ admits a partial order $\leq$, which is respected by the update function f . That is,

$$x \leq y \quad \Rightarrow \quad f(x, u) \leq f(y, u) \quad \text{for all } x, y \in \mathcal{X} \quad (6)$$

Furthermore, suppose that $\mathcal{X}$ contains maximal and minimal elements, $x^{\max}$ and $x^{\min}$, satisfying $x^{\min} \leq x \leq x^{\max}$ for all $x \in \mathcal{X}$. With this setup it becomes simple to check for coalescence in the CFTP algorithm, since the monotonicity of f guarantees that

$$\begin{aligned} F_{(-n,0]}(x^{\min}) &= F_{(-n,0]}(x^{\max}) \\ &\implies F_{(-n,0]}(\mathcal{X}) \text{ is constant} \end{aligned} \quad (7)$$

Thus, instead of running target chains, $X^{x,-n}$, starting from all states $x \in \mathcal{X}$ at time $-n$, we now only need to simulate chains started from $x^{\min}$ and $x^{\max}$. This is exactly the case in our example above, where f is monotonic (using the normal ordering on the integers): coalescence occurs exactly when the two chains $X^{1,-n}$ and $X^{4,-n}$ meet.

Evading Monotonicity

Monotonicity, while useful for CFTP, is by no means essential. Consider modifying the random walk above by setting $\mathbb{P}(X_{n+1} = 2 \mid X_n = 1) = 1$, and leaving all other transitions the same. This chain, say $\hat{X}$, is no longer reversible, and there is no monotonic update function under the usual ordering on the integers (although a different ordering, or the use of subsampling, may reintroduce monotonicity). However, a CFTP algorithm for this chain can be constructed using the *crossover trick* [7]: this first appeared in [10] to deal with *antimonotone chains*, and is also used in [11].

A variant of this approach is to use the original monotonic random walk X as an *envelope process* for $\hat{X}$. Suppose we define an update function $\hat{f}$ as follows:

$$\begin{aligned} \hat{f}(1, u) &= 2, \quad \text{and} \\ \hat{f}(x, u) &= \begin{cases} \min(x+1, 4) & \text{if } u \leq 1/2 \\ \max(x-1, 1) & \text{if } u > 1/2 \end{cases} \quad \text{for } x = 2, 3, 4 \end{aligned} \quad (8)$$

4 Perfect Sampling

Note that $\hat{f}$ is a valid update function for $\hat{X}$, and that $f(x, u) \leq \hat{f}(x, u)$ for all $x \in \mathcal{X}$ and $u \in [0, 1]$ (where f is defined in equation 5). This suggests the following perfect simulation algorithm:

- run a target chain $X^{1,-n}$ using the update function f and random sequence $\{\xi_i\}$, until the first time $S \leq 0$ that state 4 is hit (if $S \not\leq 0$, increase n and repeat, reusing $\{\xi_i\}$);
- due to the ordering of f and $\hat{f}$, $X_S^{1,-n} = \hat{X}_S^{x,-n} = 4$ for all x , and so all target chains $\hat{X}^{x,-n}$ have coalesced by time S ;
- run the chain $\hat{X}^{4,S}$ up to time zero, using $\hat{f}$ and the same $\{\xi_i\}$. Return $\hat{X}_0^{4,S}$.

The idea of using an envelope process was studied in more depth, under the name *bounding chains*, in [12]: the technique can be used in situations which are neither monotonic nor antimonotonic, and can also give bounds on the expected run time of CFTP. This trick, along with the crossover trick mentioned above, increases the possibility of efficient implementation of CFTP for a variety of chains. Furthermore, there are other variants of CFTP, which make the technique applicable to even more general chains, even those on infinite state spaces. For example, the method of *small set CFTP* [13] uses the idea of small set regeneration to determine when target chains have coalesced.

Dominated CFTP

The classic CFTP algorithm of [1] has a major drawback: a successful CFTP algorithm for X exists if and only if X is *uniformly ergodic* [6]. However, there does exist a major extension of CFTP, known as *dominated CFTP* [14, 15], or *coupling into and from the past* (CIAFTP), which can be applied to chains not satisfying this restriction. Indeed, dominated coupling from the past (domCFTP) can (in principle) be implemented for any *geometrically ergodic* chain [16], and a class of *subgeometrically ergodic* chains for which a domCFTP algorithm exists is introduced in [17], greatly extending the possible applicability of perfect simulation.

Whereas CFTP works by checking for *vertical* coalescence (target chains started from all states have coalesced by time zero), domCFTP checks for *horizontal* coalescence (all sufficiently early starts

from a specific state lead to the same result at time zero). A second chain Y is used to identify how far into the past one has to go to determine that this coalescence has occurred. To describe the algorithm in the simplest possible way, we consider here only a monotonic chain X on $\mathcal{X} = [0, \infty)$: see [18] for a much more general formulation.

Suppose that copies of X can be coupled such that, for each $x, t, u \geq 0$, and $s \geq -t$, we can construct $X^{x,-t}$ (begun at state x at time $-t$) satisfying

$$X_s^{x,-t} \leq X_s^{x,-u} \implies X_{s+1}^{x,-t} \leq X_{s+1}^{x,-u} \quad (9)$$

Suppose too that we can construct a dominating process Y on $\mathcal{X}$ which is stationary, defined for all time, and may be coupled to the target chains X such that

$$X_s^{x,-t} \leq Y_s \implies X_{s+1}^{x,-t} \leq Y_{s+1} \quad (10)$$

The domCFTP algorithm then proceeds as follows:

1. Draw Y_0 from its stationary distribution.
2. Simulate Y backwards to time $-n$.
3. Set $y = Y_{-n}$. Simulate $X^{y,-n}$ and $X^{0,-n}$ forwards to time zero (coupled to each other and to Y so that equations (9) and (10) are satisfied).
4. If $X_0^{y,-n} = X_0^{0,-n}$, then return this value as a perfect draw from π . If not, extend the realization of Y back to time $-2n$, set $n \leftarrow 2n$, and go to step 3.

A proof that this algorithm returns a perfect draw from π (so long as it terminates almost surely) may be found in [5, 15]. A consequence of the coupling in equation (9) is that the target processes are *funneled*: the earlier the two target chains (X^0 and X^y) are started, the closer they will be at time zero. As a simple example of the algorithm in practice, consider a birth–death process X with transitions $x \rightarrow x+1$ at rate $\alpha_x \leq \alpha < \infty$, and $x \rightarrow x-1$ at rate μx . This chain is clearly monotonic, and may be dominated by the birth–death chain Y , which has births at rate α and deaths at rate μ . It is easy to see that Y is reversible and has a *Poisson*(α/μ) equilibrium distribution. A realization of the domCFTP algorithm for this chain when $\alpha_x = 10 - \log(x+1)/(x+1)$ and $\mu = 1$ is shown in Figure 4. This algorithm can also be made to work if multiple simultaneous deaths are allowed to occur, or if the birth rate μ varies with x , as is the case with many problems in stochastic geometry [10, 15].

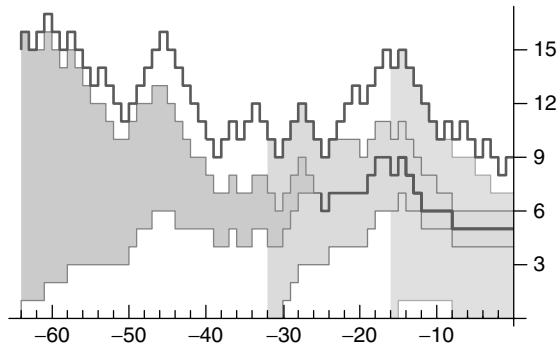


Figure 4 Implementation of domCFTP for a birth-death process. The topmost line shows the evolution of the dominating process into the past, and the shaded regions demonstrate the funneling of target processes. In this realization, all target chains started beneath the dominating process at time -64 have coalesced by time zero

Summary

Perfect simulation algorithms are both a useful practical tool and an exciting area of research. Unlike with MCMC, there is no issue of estimating a burn-in time, and the algorithm output is an exact draw from the required distribution. We have seen above how to apply CFTP and domCFTP to a couple of toy examples: generalization to more interesting chains will often require little additional theory. However, for complicated Markov chains such algorithms are typically harder to implement than MCMC, mainly due to the potential difficulty in finding an efficient method of coupling target chains together. That said, since the arrival of CFTP, more and more variants on this and other techniques mean that there is now a wide range of perfect simulation algorithms available. A comprehensive resource for further investigation of perfect simulation is the on-line bibliography of David Wilson: <http://research.microsoft.com/~dbwilson/exact/>.

References

- [1] Propp, J.G. & Wilson, D.B. (1996). Exact sampling with coupled Markov chains and applications to statistical mechanics, *Random Structures and Algorithms* **9**, 223–252.
- [2] Fill, J.A., Machida, M., Murdoch, D.J. & Rosenthal, J.S. (2000). Extension of Fill's perfect rejection sampling algorithm to general chains, *Random Structures and Algorithms* **17**(3–4), 290–316.
- [3] Lindvall, T. (2002). *Lectures on the Coupling Method*, Dover Publications.
- [4] Wilson, D.B. (2000). How to couple from the past using a read-once source of randomness, *Random Structures and Algorithms* **16**(1), 85–113.
- [5] Kendall, W.S. (2005). Notes on perfect simulation, in *Markov Chain Monte Carlo: Innovations and Applications*, Vol. 7 of *Lecture Notes Series*, W.S. Kendall, F. Liang & J.-S. Wang, eds, Institute for Mathematical Sciences, National University of Singapore, World Scientific.
- [6] Foss, S.G. & Tweedie, R.L. (1998). Perfect simulation and backward coupling, *Stochastic Models* **14**, 187–203.
- [7] Thönnies E. (2000). A primer on perfect simulation, *Statistical Physics and Spatial Statistics (Wuppertal, 1999)*, Vol. 554 of *Lecture Notes in Physics*, Springer, Berlin, pp. 349–378.
- [8] Häggström, O. (2002). *Finite Markov Chains and Algorithmic Applications*, Vol. 52 of *London Mathematical Society Student Texts*, Cambridge University Press.
- [9] Thönnies, E. (1999). Perfect simulation of some point processes for the impatient user, *Advances in Applied Probability* **31**(1), 69–87.
- [10] Kendall, W.S. (1998). Perfect simulation for the area-interaction point process, in *Probability Towards 2000*, L. Accardi & C. Heyde, eds, Springer-Verlag, pp. 218–234.
- [11] Häggström, O. & Nelander, K. (1998). Exact sampling from anti-monotone systems, *Statistica Neerlandica* **52**(3), 360–380.
- [12] Huber, M. (2004). Perfect sampling using bounding chains, *Annals of Applied Probability* **14**(2), 734–753.
- [13] Murdoch, D.J. & Green, P.J. (1998). Exact sampling from a continuous state space, *Scandinavian Journal of Statistics* **25**(3), 483–502.
- [14] Kendall, W.S. (1997). Perfect simulation for spatial point processes, in *Proceedings of the 51st Session of the ISI*, Istanbul, August 1997, Vol. 3, pp. 163–166.
- [15] Kendall, W.S. & Møller, J. (2000). Perfect simulation using dominating processes on ordered spaces, with application to locally stable point processes, *Advances in Applied Probability* **32**, 844–865.
- [16] Kendall, W.S. (2004). Geometric ergodicity and perfect simulation, *Electronic Communications in Probability* **9**, 140–151 (electronic).
- [17] Connor, S.B. & Kendall, W.S. Perfect simulation for a class of positive recurrent Markov chains, *Annals of Applied Probability* **17**(3), 781–808.
- [18] Cai, Y. & Kendall, W.S. (2002). Perfect simulation for correlated Poisson random variables conditioned to be positive, *Statistics and Computing* **12**(3), 229–243.

STEPHEN B. CONNOR

Resampling for Lifetime Distributions

Noncensored Lifetimes

Data with lifetimes are common in engineering where it is important to evaluate the reliability of technical systems, [1], and in medical sciences where data with lifetimes are collected during observations of patients, [2]. Various statistical methods are used in the analysis of such data. We consider the resampling methods. They can be used if certain assumptions are valid. It is especially important to know these assumptions because otherwise we will obtain erroneous decisions. We start with the following example. Let $\mathbf{dat}_n = \{t_1, \dots, t_n\}$ be different failure times of n similar elements that are tested. Suppose that we can consider them as values of *independent identically distributed* (i.i.d.) *random variables* (r.v.'s) $T_1, \dots, T_n$. We use capital letters for r.v.'s and small letters for their values. The *survival function* (s.f.) $S_0(t) = P[T_i > t]$ is the probability that the failure time of the i th element occurs after time t . Assume that the mean time and the variance of T_i are finite. The mean time $\mu_0 = E[T_i]$ and the α -quantile q_α are useful characteristics. They are defined by the relations $\mu_0 = \int_0^\infty S_0(t) dt$ and $S_0(q_\alpha) = 1 - \alpha$. One can use the estimators $\hat{\mu}_n = 1/n \sum_{i=1}^n T_i$ and $\hat{q}_{\alpha n} = T_{(k_{\alpha n})}$, where $k_{\alpha n} = \min_{1 \leq k \leq n} (k : \alpha \leq (k/n + 1))$, $T_{(k)} = T_{i_k}$ is the k th value in the list of ordered failure times $T_{i_1} < T_{i_2} < \dots < T_{i_n}$. Accuracy of the estimators can be characterized by the distributions of their deviations from their true values. We are interested in estimating the distributions of $\hat{\mu}_n - \mu_0$ and $\hat{q}_{\alpha n} - q_\alpha$. We need to have a sufficiently good random number generator, which can generate independent, (pseudo)random numbers J_{in}^* that are uniformly distributed on $\{1, 2, \dots, n\}$, i.e. $P^*[J_{in}^* = h] = 1/n$, $h = 1, \dots, n$, (see **Sampling from Virtual Populations**). We use “*” to stress the values related to the usage of the random number generator. Let $\mathbf{J}_n^* = \{J_{1n}^*, J_{2n}^*, \dots, J_{nn}^*\}$ and $M_{hn}^* = \sum_{i=1}^n I(J_{in}^* = h)$. $I(\cdot)$ is the indicator function. M_{hn}^* is the number of events $J_{in}^* = h$ in the list $\mathbf{J}_n^*$. We simulate independent lists $\mathbf{j}_n^{*b} = \{j_{1n}^{*b}, \dots, j_{nn}^{*b}\}$, $b = 1, \dots, B$, where j_{in}^{*b} and $m_{hn}^{*b} = \sum_{i=1}^n I(j_{in}^{*b} = h)$ are the simulated values of the r.v.'s J_{in}^* and M_{hn}^* , respectively. We obtain the b th resampled copy $\mathbf{dat}_n^{*b} = \{t_{j_{1n}^{*b}}, \dots, t_{j_{nn}^{*b}}\}$

of the original $\mathbf{dat}_n$. Let us for each $\mathbf{dat}_n^{*b}$ calculate the “copy” – estimate $\hat{\mu}_n^{*b} = 1/n \sum_{i=1}^n t_{j_{in}^{*b}} = 1/n \sum_{h=1}^n m_{hn}^{*b} t_h$ and $\hat{q}_{\alpha n}^{*b} = t_{(k_{\alpha n}^{*b})}$ which is the $k_{\alpha n}^{*b}$ failure time in the list of ordered resampled failure times $\mathbf{dat}_n^{*b}$. If certain assumptions hold as $n \rightarrow \infty$, then the numbers in the following two lists

$$\mathbf{dev}_\mu^* = \{\mu_n^* - \mu_n, \dots, \mu_n^{*B} - \mu_n\},$$

$$\mathbf{dev}_q^* = \{\hat{q}_{\alpha n}^{*1} - \hat{q}_{\alpha n}, \dots, \hat{q}_{\alpha n}^{*B} - \hat{q}_{\alpha n}\} \quad (1)$$

can be used as if they are values of i.i.d. deviations of the estimators $\hat{\mu}_n$ and $\hat{q}_\alpha$ from their true values. Efron was the first to recognized this universal utility of resampling methods [3]. We consistently estimate the distributions of deviations $\hat{\mu}_n - \mu_0$ and $\hat{q}_\alpha(n) - q_\alpha$ because then

$$\frac{1}{B} \sum_{b=1}^B I(\sqrt{n}(\mu_n^{*b} - \hat{\mu}_n) \leq x) - P[\sqrt{n}(\hat{\mu}_n - \mu_0) \leq x] \xrightarrow{P} 0 \quad (2)$$

$$\frac{1}{B} \sum_{b=1}^B I(\sqrt{n}(\hat{q}_{\alpha n}^{*b} - \hat{q}_{\alpha n}) \leq x) - P[\sqrt{n}(\hat{q}_{\alpha n} - q_\alpha) \leq x] \xrightarrow{P} 0 \quad (3)$$

for any real x as $n \rightarrow \infty$ and $B \rightarrow \infty$. The notation $\xrightarrow{P} 0$ means convergence to zero in probability.

The following variant of the central limit resampling theorem (CLRT) can be applied to justify relations (2) and (3).

Let $\{X_{in}, i = 1, \dots, n, n = 1, 2, 3, \dots\}$ be a triangular array of r.v.'s. For each $n = 2, 3, \dots$ $\{X_{1n}, \dots, X_{nn}\}$ are independent real-valued r.v.'s with finite second moments $E[X_{in}^2] < \infty$. Suppose that for any small $\varepsilon > 0$ and a number $C < \infty$,

- (i) $\sum_{i=1}^n E[(X_{in}/\sqrt{n})^2] \leq C < \infty$,
- (ii) $\sum_{i=1}^n E[(X_{in}/\sqrt{n})^2 I(|X_{in}/\sqrt{n}| > \varepsilon)] \rightarrow 0$, and
- (iii) $\sum_{i=1}^n (E[X_{in}/\sqrt{n}])^2 \rightarrow 0$ as $n \rightarrow \infty$

Then for each real number x

$$\frac{1}{B} \sum_{b=1}^B I\left(\frac{1}{\sqrt{n}} \sum_{h=1}^n (m_{hn}^{*b} - 1)x_{hn} \leq x\right) - P\left[\frac{1}{\sqrt{n}} \sum_{i=1}^n (X_{in} - E[X_{in}]) \leq x\right] \xrightarrow{P} 0 \quad (4)$$

2 Resampling for Lifetime Distributions

along a sequence of values x_{hn} of r.v.'s X_{hn} as $n \rightarrow \infty$ and $B \rightarrow \infty$.

A close variant of the CLRT is given in [4] where the more general notion of weakly approaching distributions is used. The CLRT with necessary and sufficient assumptions for consistency of resampling is proved in [5]. The sufficient for relation (2) assumptions (i)–(iii) hold with $x_{in} = t_i - \mu_0$ if r.v.'s T_i have finite variance. Relation (3) holds if the s.f. $S_0(t)$ has continuous nonzero derivative around $t = q_\alpha$. The influence functions for **quantiles** can be found by the δ -method [6]. The CLRT can be applied to the list of influence functions to justify relation (3). The convergence in probability to zero in relation (3) is rather slow but can be improved [7]. The Edgeworth expansions can be used if $\hat{\mu}_n$ is the statistic of interest. The above resampling method is inconsistent in assessing the accuracy of $\hat{\mu}_n$ if **dat**_n are n values of i.i.d. r.v.'s with infinite variance [8].

Nonparametric Estimators for Censored Lifetimes

Often lifetimes are not completely observed, i.e., their exact values T_i are not known. We say that these lifetimes are censored. There are several types of censoring. We consider right-censored lifetimes. Suppose that $T_1, \dots, T_n$ and the censoring times $U_1, \dots, U_n$ are i.i.d. r.v.'s., T_i and U_i are independent and may have different distributions. Observations are stopped at times $Y_i = U_i \wedge T_i = \min\{U_i, T_i\}$. The right-censored data **dat**_n = $\{\{Y_1, \delta_1\}, \dots, \{Y_n, \delta_n\}\}$, where $\delta_i = 1$ if and only if $T_i \leq U_i$, otherwise $\delta_i = 0$. To simplify some of the relations, we suppose that the s.f. $S_0(t)$ is continuous. Then there are no ties in the noncensored lifetimes and the *cumulative hazard function* (c.h.f.) $H_0(t) = -\log(S_0(t))$. $H_0(t)$ is often used in applications because its differential $dH_0(t)$ is the conditional probability $P[t \leq T_i < t + dt \mid t \leq T_i \wedge U_i]$.

Let us introduce the random processes $D_i(z) = I(T_i \leq z \wedge u_i)$, $R_i(z) = I(z \leq T_i \wedge u_i)$, $D_{\cdot n}(z) = \sum_{i=1}^n D_i(z)$, and $R_{\cdot n}(z) = \sum_{i=1}^n R_i(z)$. $D_{\cdot n}(z)$ is the number of lifetimes that have occurred before or at time z . $R_{\cdot n}(z)$ is the number of lifetimes which are not less z and have not been censored before z . One says that at the time z , there are $R_{\cdot n}(z)$ elements at risk. These processes have stepwise realizations. $D_{\cdot n}(z)$ is a counting process [9] and all its jumps

$\Delta D_{\cdot n}(z) = \Delta D_i(z) = 1$ if $z = T_i \leq u_i$. The conditional maximum-likelihood (ML) estimators of $S_0(t)$ and $H_0(t)$ given the censoring times $U_i = u_i$, $i = 1, \dots, n$ are

$$\hat{S}_n(t) = \prod_{z \leq t} \left(1 - \frac{\Delta D_i(z)}{R_{\cdot n}(z)}\right),$$

$$\hat{H}_n(t) = \sum_{i=1}^n \int_0^t \frac{dD_i(z)}{R_{\cdot n}(z)} = \sum_{z \leq t} \frac{\Delta D_{\cdot n}(z)}{R_{\cdot n}(z)} \quad (5)$$

$\hat{S}_n(t)$ is called the *product-limit* (PL) estimator or the *Kaplan–Meier estimator*, and $\hat{H}_n(t)$ is called the *A–estimator* or the *Aalen–Nelson estimator*. There is a one-to-one relation between these estimators. If $Q(t) = S_0(t)P[U_i \geq t] > 0$ then the estimators given in equations (5) are consistent as $n \rightarrow \infty$ [10]. A clear and short introduction to the theory of counting processes and related martingales is given in [11]. The Aalen identity holds

$$\hat{H}_n(t) = H_0(t) + \int_0^t \frac{dM_{\cdot n}(z)}{R_{\cdot n}(z)} - \int_0^t I(R_{\cdot n}(z) = 0) dH_0(z) \quad (6)$$

where $M_{\cdot n}(z) = \sum_{i=1}^n M_i(z)$, $M_i(z) = D_i(z) - \int_0^t R_i(z) dH_0(z)$ are martingales. From identity (6) it follows that $\hat{H}_n(t)$ has a negligible negative **bias** because $n \int_0^t I(R_{\cdot n}(z) = 0) dH_0(z) \xrightarrow{P} 0$ as $n \rightarrow \infty$.

Then from equation (6) and $\sum_{i=1}^n (M_{in}^* - 1) = 0$ it follows that

$$\sqrt{n}(\hat{H}_n(t) - H_0(t)) - \frac{1}{\sqrt{n}} \sum_{i=1}^n \left(\int_0^t \frac{dD_i(z)}{Q(z)} - H_0(t) \right) \xrightarrow{P} 0 \quad \text{as } n \rightarrow \infty \quad (7)$$

$$\sqrt{n} \sum_{i=1}^n (M_{in}^* - 1) \int_0^t \frac{dD_i(z)}{R_{\cdot n}(z)} - \frac{1}{\sqrt{n}} \sum_{i=1}^n (M_{in}^* - 1) \times \int_0^t \frac{dD_i(z)}{Q(z)} \xrightarrow{P} 0 \quad \text{as } n \rightarrow \infty \quad (8)$$

All assumptions in the CLRT hold for the last sum in relation (8). Hence, from relations (7) and (8) we obtain that values of the left sums in relation (8)

will have asymptotically the same distributions as $\sqrt{n}(\hat{H}_n(t) - H_0(t))$, $n \rightarrow \infty$.

Let us consider the generalized ML estimator $\hat{H}_n^{*b}(t) = \int_0^t (dD_n^{*b}(z)/R_n^{*b}(z))$, where $D_n^{*b}(z) = \sum_{i=1}^n M_{in}^{*b} D_i(z)$ and $R_n^{*b}(z) = \sum_{i=1}^n M_{in}^{*b}(z) R_i(z)$, $b = 1, \dots, B$. Then for any real x

$$\frac{1}{B} \sum_{i=1}^B \mathbb{I}(\sqrt{n}(\hat{H}_n^{*b}(t) - \hat{H}_n(t)) \leq x) - \mathbb{P}[\sqrt{n}(\hat{H}_n(t) - H_0(t)) \leq x] \xrightarrow{P} 0 \text{ as } n \rightarrow \infty, B \rightarrow \infty \quad (9)$$

These results can be applied in assessing the accuracy of the PL estimator $\hat{S}_n(t)$ given in equation (5). In a similar way we evaluate the accuracy of the nonparametric ML estimators of the s.f.'s and the c.h.f.'s of renewal times when independent renewal processes are observed during a series of times intervals. We assume that the ends of the time intervals and the renewal times are independent. Here, we have to resample from the subsets of data containing renewals which occurred during the same time interval. Evaluation of the accuracy of estimated reliability characteristics is actual in constructing some mechanical systems with redistributions of applied loads between survived system elements. Cables with many fibers are examples of such systems [12]. Applications of resampling methods to such mechanical systems are given in [13].

Parametric Families of Lifetime Distributions

Parametric families of distributions in lifetime analysis can be useful in obtaining approximated estimators of reliability or survival characteristics and in assessing their accuracy with data containing explanatory variables. Let the elements $\{y_i, \delta_i\}$ of $\mathbf{dat}_n$ be values of independent r.v.'s $\{Y_i, \delta_i\}$ the probability densities $p(y_i, \delta_i, \theta_0)$ of which are known functions of the unknown true parameter $\theta_0 \in \Theta \subseteq \mathbb{R}^k$, then the log-likelihood function [14] is

$$\begin{aligned} \text{lik}(\theta | \mathbf{dat}_n) &= \sum_{i=1}^n l_i(\theta | y_i, \delta_i), \\ l_i(\theta | y_i, \delta_i) &= \log(p(y_i, \delta_i, \theta)) \end{aligned} \quad (10)$$

Suppose that a consistent and asymptotically unbiased ML estimator $\hat{\theta}_n$ of the true parameter θ_0 exists and is defined by the equation

$$\text{lik}(\hat{\theta}_n | \mathbf{dat}_n) = \max(\text{lik}(\theta | \mathbf{dat}_n) : \theta \in \Theta) \quad (11)$$

It is possible to consider resampled copies of the log-likelihood defined by the relations

$$\text{lik}^{*b}(\theta | \mathbf{dat}_n) = \sum_{h=1}^n m_{hn}^{*b} l_h(\theta | \mathbf{dat}_n), \quad b = 1, \dots, B \quad (12)$$

$\{m_{hn}^{*b}\}$ were introduced in the section titled “Noncensored Lifetimes”. The resampled ML estimators $\hat{\theta}_n^{*b}$ are defined by the equations

$$\begin{aligned} \text{lik}^{*b}(\hat{\theta}_n^{*b} | \mathbf{dat}_n) &= \max(\text{lik}^{*b}(\theta | \mathbf{dat}_n) : \theta \in \Theta), \\ b &= 1, \dots, B \end{aligned} \quad (13)$$

The k -dimensional distributions of the components in the list $\{\hat{\theta}_n^{*1} - \hat{\theta}_n, \dots, \hat{\theta}_n^{*B} - \hat{\theta}_n\}$ asymptotically approach the distributions of deviations $\hat{\theta}_n - \theta_0$ under nearly usual regularity assumptions [15] as $n \rightarrow \infty$.

The resampling methods can be efficiently applied to evaluate the accuracy of estimators of lifetime parameters if the data with lifetimes are sufficiently large. The resampling methods can also be used to construct parameter **confidence intervals** (see **Bootstrap and Jackknife, Overview**).

Related influence functions, the CLRT and the functional CLRT can be used in proofs. Two variants of the functional CLRT theorems are given in [13] and [16].

References

- [1] Kalbfleisch, J. & Prentice, R. (1980). *The Statistical Analysis of Failure Time Data*, John Wiley & Sons, New York.
- [2] Klein, J. & Moeschberger, M. (2003). *Survival Analysis*, 2nd Edition, Springer, New York.
- [3] Efron, B. (1979). Bootstrap methods: another look at the jackknife, *Annals of Statistics* 7, 1–26.
- [4] Belyaev, Y. & Sjöstedt de-luna, S. (2000). Weakly approaching sequences of random distributions, *Journal of Applied Probability* 37(3), 807–822.
- [5] Belyaev, Y. (2003). Necessary and sufficient conditions for consistency of resampling, Research Report 2003-1, Center of Biostochastics, Swedish University of

4 Resampling for Lifetime Distributions

- Agricultural Sciences, Electronic Publications, <http://biostochastics.slu.se/publikationer>.
- [6] Davison, J. & Hinkley, D. (1997). *Bootstrap Methods and their Application*, Cambridge University Press, Cambridge.
 - [7] Shao, J. & Tu, D. (1995). *The Jackknife and Bootstrap*, Springer, New York.
 - [8] Athreya, K. (1987). Bootstrap of the mean in the infinite variance case, *Annals of Statistics* **17**, 724–731.
 - [9] Andersen, P., Borgan, O., Gill, R. & Keiding, N. (1993). *Stochastic Models Based on Counting Processes*, Springer, New York.
 - [10] Belyaev, Y. (1991). Asymptotical properties of nonparametric point estimators based on complexly structured reliability data with right-censoring, *Statistics* **25**(4), 589–612.
 - [11] Andersen, P. & Borgan, O. (1985). Counting process models for life history data, *Scandinavian Journal of Statistics* **12**, 97–158.
 - [12] Crowder, M., Kimber, A., Smith, R. & Sweeting, T. (1988). *Extreme Value Theory in Engineering*, Academic Press, San Diego.
 - [13] Rydén, P. (2000). Statistical analysis and simulation methods related to load-sharing models, Doctoral Dissertation, Department of Mathematical Statistics, Umeå University: Umeå, Sweden, ISBN 91-7191-965-1.
 - [14] Lehman, E. & Casella, G. (1998). *Theory of Point Estimation*, 2nd Edition, Springer, New York.
 - [15] Nilsson, L. (1998). Parametric estimation using computer intensive methods, Doctoral Dissertation, Department of Mathematical Statistics, Umeå University: Umeå, Sweden, ISBN 91-7191-483-8.
 - [16] Belyaev, Y. & Seleznev, O. (2000). Approaching in distribution with applications to resampling of stochastic processes, *Scandinavian Journal of Statistics* **27**, 371–384.

Related Articles

Bootstrap and Jackknife, Overview; Sampling from Virtual Populations.

YURI K. BELYAEV

Lévy Processes, Simulation of

Lévy Processes

Lévy processes arise in many areas of applied probability. (See, e.g., relevant articles **Monte Carlo Methods, Univariate and Multivariate; Markov Chain Monte Carlo, Introduction; Survival Analysis, Nonparametric; Dirichlet Process, Simulation of; Convergence and Mixing in Markov Chain Monte Carlo; Simulation of the Beta-Stacy Process.**) For theory of Lévy processes the reader is referred to books of Bertoin [1] and Sato [2], which are the classical references on this topic. A stochastic process $\{X(t)\}_{t \geq 0}$ is said to be a *Lévy process* if it has stationary and independent increments and satisfies $X(0) = 0$. One may always assume that sample paths of a Lévy process are right continuous and have left limits.

Standard **Brownian motion** $\{B(t)\}$ is a Lévy process, so is a linear Brownian motion $\{\mu t + \sigma B(t)\}$. A **Poisson process** $\{N(t)\}$ with rate λ is a Lévy process, so is a compound Poisson process. A linear combination of independent Lévy processes is again a Lévy process. In particular, the sum of a linear Brownian motion and a compound Poisson process

$$X(t) = \mu t + \sigma B(t) + \sum_{j=1}^{N(t)} V_j \quad (1)$$

is a Lévy process, where V_j are independent and identically distributed (i.i.d.) and $\{B(t)\}$, $\{N(t)\}$, and $\{V_j\}$ are independent of each other.

The key to understanding a Lévy process $\{X(t)\}$ is the structure of its jumps $J(t) = X(t) - X(t-)$. It occurs that for any Borel subset B of $\mathbb{R} \setminus \{0\}$, the counting process of B , $N^B(t) = \text{Card}\{s \leq t : J(s) \in B\}$ is a Poisson process with rate $\nu(B)$, where ν is a measure satisfying $\nu(\{0\}) = 0$ and

$$\int \min\{x^2, 1\} \nu(dx) < \infty \quad (2)$$

ν is called the *Lévy measure* of $\{X(t)\}$. Thus $\nu(dx)$ is the intensity of jumps of size x . The *Lévy-Itô*

representation reveals the path structure of a Lévy process:

$$X(t) = at + \sigma B(t) + \lim_{\epsilon \downarrow 0} \left\{ \sum_{s \leq t} J(s) \mathbf{1}_{\{|J(s)| > \epsilon\}} - t \int_{\epsilon < |x| \leq 1} x \nu(dx) \right\} \quad (3)$$

where $\{B(t)\}$ is a standard Brownian motion independent of the process of jumps $\{J(t)\}_{t \geq 0}$. The convergence on the right-hand side holds almost surely (a.s.) uniformly in t on each finite interval. When $\nu(\mathbb{R}) < \infty$, equation (3) reduces to equation (1) with $\mu = a - \int_{|x| \leq 1} x \nu(dx)$. If

$$\int \min\{|x|, 1\} \nu(dx) < \infty \quad (4)$$

then $\sum_{s \leq t} |J(s)| < \infty$ a.s. and equation (3) reduces to

$$X(t) = \mu t + \sigma B(t) + \sum_{s \leq t} J(s) \quad (5)$$

where μ is as above. Sample paths of $\{X(t)\}$ have *bounded variation* on each finite interval if and only if $\sigma = 0$ and equation (4) holds. If additionally $\mu \geq 0$ and $\nu((-\infty, 0)) = 0$, then $\{X(t)\}$ has nondecreasing sample paths and is called a *subordinator*.

The characteristic function of $X(t)$ is of the form $E e^{iuX(t)} = \exp(t\psi(u))$, where $\psi(u)$ is the cumulant of $X(1)$ given by the *Lévy-Khintchine formula*

$$\psi(u) = iau - \frac{\sigma^2 u^2}{2} + \int_{-\infty}^{\infty} (e^{iux} - 1 - iux \mathbf{1}_{\{|x| \leq 1\}}) \nu(dx) \quad (6)$$

where (a, σ^2, ν) are the same as in equation (3).

The *generating triplet* (a, σ^2, ν) or the cumulant $\psi(u)$ of $X(1)$ are usually used as identifiable parameterizations of Lévy processes.

Simulation Methods: An Overview

Simulation of a Brownian motion and/or of a compound Poisson process can be found in many textbooks and will not be discussed here. In view of the representation in equation (3), Brownian

2 Lévy Processes, Simulation of

and Poissonian-type components of a Lévy process can be generated independently of each other. The main problem is to simulate a Poissonian-type component of a Lévy process having infinite Lévy measure.

If the Lévy measure is infinite, then by equation (3) sample paths of $\{X(t)\}$ have infinitely many jumps in each finite interval. Exact simulation of such process is obviously impossible. A process that is close to the original one is generated instead. The following methods will be considered:

1. random walk approximation
2. series representations of Lévy processes
3. Poisson and Gaussian approximations.

1. A discrete skeleton $\{X(jh) : j = 0, 1, \dots\}$ is a random walk. Therefore, if a simulation method (at least an approximate one) of $X(h)$ is available, then one can simulate $\{X(t)\}$ approximately. A disadvantage of such a method is that one cannot precisely identify the location and magnitude of the large jumps. Large jumps are important, especially in the heavy-tailed case, because they determine various functionals of a Lévy process. Another complication may be that a simulation of $X(h)$ may be numerically tedious (e.g., when its density involves special functions).

2. Series representations provide uniform approximations of Lévy processes along the sample paths. They are often easy to simulate. Usually the largest jumps of a Lévy process are included in the first few terms of the series. A disadvantage of this method is that some series may converge slowly. Therefore, one may need a huge number of terms in a series to reach a desired accuracy of the approximation. On the other hand, with ever increasing computational speed, a slow convergence may not be an issue of practical importance for some applications.

3. If one discards small jumps from the right-hand side of equation (5) or if one replaces them by their mean value on the right-hand side of equation (3), then the resulting process is a compound Poisson process with a drift ($\sigma = 0$). This is a Poisson approximation of a Lévy process. It converges uniformly in t on each finite interval, as the size of the discarded jumps tends to zero. The large jumps are precisely simulated. Series representations of Lévy processes provide a consistent way to construct

such approximations. However, when the intensity of small jumps is high, discarding them may produce a substantial error. In such case, instead of discarding a small-jump part of a Lévy process, one can often approximate it by a Brownian motion with a small variance. This approximation complements method 2 because it is applicable when the series converges slowly.

Methods based on 1, 2, and 3 will be discussed in the following sections. Further information on simulation of Lévy processes and related topics can be found in the books of Cont and Tankov [3] and Asmussen and Glynn [4]. The author is grateful to Professor Søren Asmussen for making the manuscript of his forthcoming book available [4].

Random Walk Approximation

Suppose $\{X(t)\}$ is a Lévy process determined by equation (6) with $\sigma = 0$ and an infinite Lévy measure ν . Fix the time domain $[0, T)$, $n \geq 1$, and put $h = T/n$. Generate the increments $\Delta_j^h X = X(jh) - X((j-1)h)$ as i.i.d. random variables with the distribution $P_h(\cdot) = \mathbb{P}(X(h) \in \cdot)$, $j = 1, \dots, n-1$, and let

$$X^h(t) = \begin{cases} 0 & \text{if } 0 \leq t < h \\ \Delta_1^h X + \dots + \Delta_j^h X & \text{if } jh \leq t < (j+1)h \end{cases} \quad (7)$$

Process $\{X^h(t)\}_{0 \leq t < T}$ is a random walk approximation to $\{X(t)\}_{0 \leq t < T}$.

Example 1: Gamma process If $X(t) \sim \Gamma(\alpha t, \lambda)$, then $\{X(t)\}$ is called a *gamma process* with shape parameter $\alpha > 0$ and scale $\lambda > 0$. The density of $X(h)$ is given by

$$p_h(x) = \frac{\lambda^{\alpha h}}{\Gamma(\alpha h)} x^{\alpha h - 1} e^{-\lambda x} \quad (8)$$

There exist many algorithms for generating $\Delta_j^h X$'s with this density e.g., Johnk's and Best's algorithms; see the book of Devroye [5] for a survey on this topic.

Example 2: Stable process A stable process is determined by its cumulant function $\psi(u)$ in equa-

tion (6) of the form

$$\psi(u) = \begin{cases} -\theta^\alpha |u|^\alpha \left(1 - i\beta \operatorname{sgn} u \tan \frac{\pi\alpha}{2}\right) + i\mu u & \text{if } 0 < \alpha < 2, \alpha \neq 1 \\ -\theta |u| \left(1 + i\beta \frac{2}{\pi} \operatorname{sgn} u \log |u|\right) + i\mu u & \text{if } \alpha = 1 \end{cases} \quad (9)$$

where $0 < \alpha < 2$, $\theta > 0$, and $-1 \leq \beta \leq 1$. Densities $p_h(x)$ are known in only a few cases of α . However, the increments $\Delta_j^h X$ can be simulated using an algorithm due to Chambers *et al.* [6] (see also [7] and [8]). In the symmetric case, i.e. $\psi(u) = -\theta^\alpha |u|^\alpha$, this algorithm has a specially simple form

$$\Delta_j^h X = \theta h^{1/\alpha} \frac{\sin \alpha U_j}{(\cos U_j)^{1/\alpha}} \left(\frac{\cos((1-\alpha)U_j)}{V_j} \right)^{(1-\alpha)/\alpha} \quad (10)$$

where U_j are i.i.d. Uniform $(-\pi/2, \pi/2)$ random variables independent of i.i.d. standard exponential random variables V_j .

Example 3: Damien, Laud, and Smith algorithm

This algorithm provides an approximation to an infinitely divisible random variable. Let ν be a Lévy measure and define $\theta(dx) = x^2/(1+x^2) \nu(dx)$, $c = \int_{-\infty}^{\infty} \theta(dx)$. Fix $m \geq 1$ and let (U_i, V_i) , $i = 1, \dots, m$, be i.i.d. pairs such that U has distribution $\theta(dx)/c$ and given $U = u$, $V \sim \text{Poisson}(\lambda_m(u))$, where $\lambda_m(u) = c(1+u^2)/(mu^2)$. Then, as $m \rightarrow \infty$,

$$\sum_{i=1}^m \left(U_i V_i - \frac{c}{n U_i} \right) \quad (11)$$

converges in distribution to an infinitely divisible random variable with the generating triplet $(\int (x \mathbf{1}_{\{|x| \leq 1} - (x/1+x^2)) \nu(dx), 0, \nu)$. The random centers in equation (11) are not needed when equation (4) holds; see [9] for the proofs and further details. Using this algorithm for an approximate simulation of a Lévy process $\{X(t)\}$, one needs to fix n , the number of partitions of the time domain $[0, T]$ and m , the number of terms in equation (11) to approximately generate the increment $\Delta_j^h X$. Two types of errors are committed, one from using a discrete skeleton and another one from each approximation at a grid point.

Series Representations of Lévy Processes

Series representations provide uniform approximations of Lévy processes along sample paths. It occurs that any Lévy process $\{X(t)\}$ with the generating triplet $(a, 0, \nu)$ can be represented almost surely as uniform in t convergent series of the form

$$X(t) = \sum_{j=1}^{\infty} [H(\Gamma_j, V_j) \mathbf{1}_{\{U_j \leq t\}} - t a_j], \quad 0 \leq t \leq T \quad (12)$$

where $\{\Gamma_j\}$ are the arrival times in a Poisson process with rate one, $\{V_j\}$ are i.i.d. random vectors taking values in some Euclidean space, $\{U_j\}$ are i.i.d. uniform on $[0, T]$ variables and the sequences $\{\Gamma_j\}$, $\{V_j\}$, and $\{U_j\}$ are independent of each other. Here H is a jointly measurable real-valued function such that $r \mapsto |H(r, v)|$ is nonincreasing and $a_j \in \mathbb{R}$. The choice of H and V is not unique; they have to satisfy the condition

$$\int_{-\infty}^{\infty} f(x) \nu(dx) = T^{-1} \int_0^{\infty} \mathbb{E} f(H(r, V)) dr \quad (13)$$

for any nonnegative Borel function f with $f(0) = 0$. See [10] for a historical survey on methods of selection of (H, V) , the relation of equation (12) to the Lévy–Itô representation in equation (3), and further details on representation in equation (12). In the following examples the time domain $[0, T]$ will be fixed and $\{\Gamma_j\}$, $\{U_j\}$ will be as above. Then H, V will be specified. $\{V_j\}$ are assumed to be independent of $\{\Gamma_j\}$ and $\{U_j\}$.

Example 4: Gamma process We use the notation of Example 1. Let $\{V_j\}$ be an i.i.d. sequence of exponential random variables with parameter λ . Then

$$X(t) = \sum_{j=1}^{\infty} e^{-\Gamma_j/(\alpha T)} V_j \mathbf{1}_{\{U_j \leq t\}}, \quad 0 \leq t \leq T \quad (14)$$

4 Lévy Processes, Simulation of

represents a gamma process with parameters α, λ . Such representation for gamma random variables is due to Bondesson [11], an extension to gamma processes is immediate. Since $e^{-\Gamma_j/(\alpha T)} V_j = O(e^{-j/(\alpha T)} \log j)$, the series converges very fast and its generation is a good alternative to the simulation of a discrete skeleton discussed in Example 1.

Example 5: Stable process Here we will only present the symmetric case. Representations for non-symmetric cases can be found in [7] and [8]. Let $\{V_j\}$ be an i.i.d. sequence such that $\mathbb{P}(V_j = \pm 1) = 1/2$. Then

$$X(t) = \theta T^{1/\alpha} c_\alpha \sum_{j=1}^{\infty} V_j \Gamma_j^{-1/\alpha} \mathbf{1}_{\{U_j \leq t\}}, \quad 0 \leq t \leq T \quad (15)$$

represents a stable process with the cumulant function $\psi(u) = \exp(-\theta^\alpha |u|^\alpha)$, where $0 < \alpha < 2$, $c_\alpha = |\Gamma(1 - \alpha) \cos \frac{\pi\alpha}{2}|^{-1/\alpha}$ when $\alpha \neq 1$, and $c_1 = \pi/2$. This is the so-called LePage representation.

Example 6: Tempered stable process A Lévy process $\{X(t)\}$ with a generating triplet $(a, 0, \nu)$ is said to be *tempered α -stable* if $\nu(dx) = |x|^{-\alpha-1} g(x) dx$, where $(-1)^n D^n g(x) \geq 0$ for $x > 0$, $D^n g(x) \geq 0$ for $x < 0$, $n = 0, 1, \dots$ and $\lim_{|x| \rightarrow \infty} g(x) = 0$. A typical example is $g(x) = \kappa_1 e^{-\lambda_1 x}$ for $x > 0$ and $g(x) = \kappa_2 e^{\lambda_2 x}$ for $x < 0$. In this case $\{X(t)\}$ is also called a *CGMY process* ($\kappa_1, \kappa_2 \geq 0$, $\lambda_1, \lambda_2 > 0$). Series representations of tempered α -stable processes are attributable to Rosiński [12]. Here we will present a special case as an example.

Suppose that $\nu(dx) = \kappa |x|^{-\alpha-1} e^{-\lambda|x|} dx$. Let $\{\epsilon_j\}$, $\{\eta_j\}$, and $\{\xi_j\}$ be sequences of i.i.d. random variables such that $\mathbb{P}(\epsilon_j = \pm 1) = 1/2$, $\eta_j \sim \text{Exponential}(\lambda)$, and $\xi_j \sim \text{Uniform}(0, 1)$. All random sequences $\{\Gamma_j\}$, $\{U_j\}$, $\{\epsilon_j\}$, $\{\eta_j\}$, and $\{\xi_j\}$ are assumed to be independent of each other. Then

$$X(t) = \sum_{j=1}^{\infty} \epsilon_j \left(\left(\frac{\alpha \Gamma_j}{2\kappa T} \right)^{-1/\alpha} \wedge \eta_j \xi_j^{1/\alpha} \right) \mathbf{1}_{\{U_j \leq t\}}, \quad 0 \leq t \leq T \quad (16)$$

represents a symmetric tempered α -stable process with Lévy measure as above.

One can couple sample paths of an α -stable and a tempered α -stable process via equations (15)–(16). As an illustration we show in Figure 1 a simulation of both processes having $\alpha = 1.3$, θ, κ satisfying $\theta c_\alpha = (2\kappa)^{1/\alpha} = 1$, $\lambda = 1$, and $T = 1$ (see Examples 5 and 6). We used continuous lines instead of dots to draw (discontinuous!) sample paths for a better clarity of the plot. It is quite instructive to see how the tempering occurs and diminishes high variability of a stable sample path.

Poisson and Gaussian Approximations

Any Lévy process $\{X(t)\}$ with the generating triplet $(a, 0, \nu)$ can be decomposed into a sum of two independent Lévy processes

$$X(t) = X^\epsilon(t) + X_\epsilon(t) \quad (17)$$

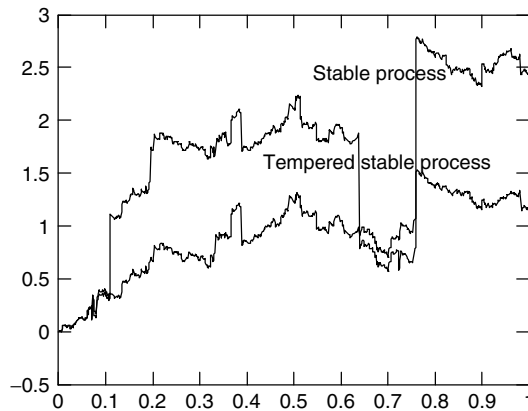


Figure 1 Trajectories of coupled stable and tempered stable processes

where $\{X^\epsilon(t)\}$ is a compound Poisson process with a drift and the distribution of jumps proportional to $\nu^\epsilon = \nu_{|\{x|>\epsilon\}}$ and the process $\{X_\epsilon(t)\}$ has mean zero and the Lévy measure $\nu_\epsilon = \nu_{|\{x|\leq\epsilon\}}$; $\epsilon > 0$ is fixed. We refer to $\{X_\epsilon(t)\}$ as a *small-jump part* of $\{X(t)\}$. When ϵ is small one can consider

$$X(t) \approx X^\epsilon(t) \quad (18)$$

and simulate the right-hand side as an approximation to $\{X(t)\}$. However, when the intensity of small jumps is high (as for a stable process with $\alpha > 1$), this method may require simulating of an enormous number of jumps to obtain a reasonable accuracy of the approximation. In that case it was proposed to replace X_ϵ in equation (17) by a Brownian motion with a variance $\sigma_\epsilon^2 = \text{Var}(X_\epsilon(1))$. Asmussen and Rosiński [13] rigorously discussed this approximation and showed that $\sigma_\epsilon^{-1}X_\epsilon$ converges in distribution to W , where W is a standard Brownian motion, if and only if $\sigma_{c\sigma_\epsilon \wedge \epsilon} \sim \sigma_\epsilon$ for each $c > 0$ as $\epsilon \downarrow 0$. A close sufficient condition for the latter one is $\sigma_\epsilon/\epsilon \rightarrow \infty$ as $\epsilon \downarrow 0$. Using this Gaussian approximation of X_ϵ we get

$$X(t) \approx X^\epsilon(t) + \sigma_\epsilon W(t) \quad (19)$$

where $\{W(t)\}$ is assumed to be independent of $\{X^\epsilon(t)\}$.

To select ϵ experimentally, one needs a consistent procedure for generating $\{X^\epsilon(t)\}$ as $\epsilon \downarrow 0$. This is possible when a series representation (equation 12) is available. One can choose, for example,

$$X^\epsilon(t) = \sum_{j: \Gamma_j \leq \epsilon^{-1}} [H(\Gamma_j, V_j)\mathbf{1}_{\{U_j \leq t\}} - ta_j] \quad (20)$$

The resulting decomposition $X = X^\epsilon + X_\epsilon$ is not exactly the same as in equation (17) (see for instance Example 6) but the terms of it are independent and the same criterion as above for Gaussian approximation of X_ϵ applies (see [14]).

Stable and tempered stable processes admit Gaussian approximation of small jumps (*cf.* Examples 5 and 6) but a gamma process does not. As discussed in Example 4, a (compound) Poisson approximation of a gamma process is very fast in this case, consequently the small-jump part can be neglected.

Lévy processes are rarely available. This essentially rules out an approximation by a random walk from a discrete skeleton. However, in the multidimensional case one can still use Poisson and Gaussian approximations combined with series representations (see [14]) or use Lévy copulas (see [3] and references therein).

References

- [1] Bertoin, J. (1996). *Lévy Processes*, Cambridge University Press.
- [2] Sato, K. (1999). *Lévy Processes and Infinitely Divisible Distributions*, Cambridge University Press, Cambridge.
- [3] Cont, R. & Tankov, P. (2004). *Financial Modelling with Jump Processes*, Chapman & Hall/CRC, Boca Raton.
- [4] Asmussen, S. & Glynn, P.W. (2007). *Stochastic Simulation: Algorithms and Analysis*, Springer.
- [5] Devroye, L. (1986). *Non-Uniform Random Variate Generation*, Springer, New York.
- [6] Chambers, J.M., Mallows, C.L. & Stuck, B.W. (1976). A method for simulating stable random variables, *Journal of the American Statistical Association* **71**, 340–344.
- [7] Janicki, A. & Weron, A. (1994). *Simulation and Chaotic Behavior of α -Stable Stochastic Processes*, Marcel Dekker, New York.
- [8] Samorodnitsky, G. & Taqqu, M.S. (1994). *Non-Gaussian Stable Processes*, Chapman & Hall.
- [9] Damien, P., Laud, P.W. & Smith, A.F.M. (1995). Approximate random variate generation from infinitely divisible distributions with applications to Bayesian inference, *Journal of the Royal Statistical Society* **B57**, 547–563.
- [10] Rosiński J. (2001). Series representations of Lévy processes from the perspective of point processes, in *Lévy Processes*, O.E. Barndorff-Nielsen, T. Mikosch & S.I. Resnick, eds, Birkhäuser, Boston, pp. 401–415.
- [11] Bondesson, L. (1982). On simulation from infinitely divisible distributions, *Advances in Applied Probability* **14**, 855–869.
- [12] Rosiński, J. (2007). Tempering stable processes, *Stochastic Processes and their Applications* **117**, 677–707.
- [13] Asmussen, S. & Rosiński, J. (2001). Approximations of small jumps of Lévy processes with a view towards simulation, *Journal of Applied Probability* **38**, 482–493.
- [14] Cohen, S. & Rosiński, J. (2007). Gaussian approximation of multivariate Lévy processes with applications to simulation of tempered stable processes, *Bernoulli* **13**, 195–210.

JAN ROSIŃSKI

Multidimensional Extensions

Contrary to the one-dimensional case, closed formulas for simulation of increments of multidimensional

Life Distributions, Simulation of

Sampling Random Values from a Probability Distribution

Let X be a random variable (rv) obeying a cumulative distribution function (cdf) $F_X(x)$ defined as follows:

$$\text{Prob}\{X \leq x\} = F_X(x); F_X(-\infty) = 0; F_X(\infty) = 1 \quad (1)$$

In the following, if the rv X obeys a cdf $F_X(x)$ we shall write $X \sim F_X(x)$.

From the definition it follows that $F_X(x)$ is a nondecreasing function and we further assume that it is continuous and differentiable as well. The corresponding **probability density function** (pdf) is then

$$f_X(x) = \frac{dF_X(x)}{dx}; f_X(x) \geq 0; \int_{-\infty}^{\infty} f_X(x) dx = 1 \quad (2)$$

We now aim at sampling numbers from the cdf $F_X(x)$. This amounts to obtaining a sequence of $N \gg 1$ values $\{X\} \equiv \{X_1, X_2, \dots, X_N\}$ that satisfy the following conditions:

1. The number n of sampled values within an interval $\Delta x \ll X_{\max} - X_{\min}$, where $X_{\min}$ and $X_{\max}$ are the minimum and maximum values in $\{X\}$, should be such that

$$\frac{n}{N} \approx \int_{\Delta x} f_X(x) dx \quad (3)$$

In other words, we require that the histogram of the sampled data should approximate $f_X(x)$.

2. The X_i values should be uncorrelated.
3. If the sequence $\{X\}$ is periodic, the period after which the numbers start repeating should be as large as possible.

Among all the distributions, the uniform distribution in the interval $[0, 1)$, denoted as $U[0, 1)$ or more simply $U(0, 1)$, plays a role of fundamental importance since sampling from this distribution allows obtaining rv's obeying any other distribution.

Sampling from the Uniform Distribution $U[0, 1)$

The cdf and pdf of the distribution $U[0, 1)$ are as follows:

$$U_R(r) = r; u_R(r) = 1 \quad 0 \leq r < 1 \quad (4)$$

The generation of random numbers R uniformly distributed in $[0, 1)$ has represented, and still represents, an important subject of active research. In the beginning, the outcomes of intrinsically random phenomena were used (e.g. throwing a coin or dice, spinning the roulette, counting of radioactive sources of constant intensity, etc.), but soon it was realized that, apart from the disuniformity due to imperfections of the mechanisms of generation or detection, the frequency of data thus obtained was too low and the sequences could not be reproduced, so that it was difficult to find and fix the errors in the computer codes in which the sampled random numbers were used.

To overcome these difficulties, the next idea was to fill tables of random numbers to store in the computers (in 1955, RAND corporation published a table with 10^6 numbers), but the access to the computer memory decreased the calculation speed and, above all, the sequences that had been memorized were always too short with respect to the growing necessities.

Finally, in 1956, von Neumann proposed to have the computer directly generate the "random" numbers by means of an appropriate function $g(\cdot)$, which should allow one to find the next number R_{k+1} from the preceding one R_k , i.e. $R_{k+1} = g(R_k)$.

The sequence thus generated is inevitably periodic: in the course of the sequence, when a number is obtained that had been obtained before, the subsequence between these two numbers repeats itself cyclically, i.e. the sequence enters a loop. Furthermore, the sequence itself can be reproduced so that it is obviously not "random", but rather deterministic. However, if the function $g(\cdot)$ is chosen correctly, it can be said to have a pseudorandom character if it satisfies a number of randomness tests.

In particular, von Neumann proposed to obtain R_{k+1} by taking the central digits of the square of R_k . For example, for a computer with a four digit word, if $R_k = 4567$, then $R_k^2 = 20857489$ and $R_{k+1} = 8574$, $R_{k+2} = 5134$, and so on. This function turns out to be lengthy to calculate and to give rise to rather short periods; furthermore, if one obtains

2 Life Distributions, Simulation of

$R_k = 0000$, then all the following numbers are also zero.

Presently, the most commonly used methods for generating sequences $\{R\}$ of numbers from a uniform distribution are inspired by the Monte Carlo roulette game [1–6].

In a real roulette game, the ball, thrown with high initial speed, performs a large number of revolutions around the wheel and finally it comes to rest within one of the numbered compartments. In an ideal machine nobody would doubt that the final compartment, or its associated number, is actually uniformly sampled among all the possible compartments or numbers. In the domain of the real numbers, within the interval $[0, 1)$, the game could be modeled by throwing a point on the positive x axis very far from the origin, utilizing a method having an intrinsic dispersion much larger than unity. The difference between the value so obtained and the largest integer smaller than this value may then be reasonably assumed as sampled from $U[0, 1)$. Obviously this statement is true if a *suitable* method is utilized for throwing the point. In a computer, the above procedure is performed by means of a mixed congruential relationship of the kind

$$R_{k+1} = (aR_k + c) \bmod m \quad (5)$$

In words, the new number R_{k+1} is the remainder, modulo m (a positive integer), of an affine transform of the old R_k , with nonnegative integer coefficients a and c . The above expression in some way resembles the uniform sampling in the roulette game, $aR_k + c$ playing the role of the distance traveled by the ball and m that of the wheel circumference. The sequence so obtained is made up of numbers $R_k < m$ and is periodic with period $p < m$. For example, if we choose $R_0 = a = c = 5$ and $m = 7$, the sequence is $\{5, 2, 1, 3, 6, 0, 5, \dots\}$, with a period $p = 6$. The sequences generated with the method described above are actually deterministic so that the sampled numbers are more appropriately called *pseudorandom* numbers. However, the constants a , c , and m may be selected so that:

- the sequence satisfies essentially all randomness tests;
- the period p is very large.

Clearly the numbers generated by the above procedure are always smaller than m so that, when divided by m , they lie in the interval $[0, 1)$.

It is important to note that for the generation of pseudorandom numbers $U[0, 1)$ many computer codes do not make use of machine subroutines, but use congruential subroutines which are part of the program itself. Thus, for example, it is possible that an excellent program executed on a machine with a word of length different from the one it was written for, gives absurd results. In this case, it should not be concluded that the program is “garbage”, but it would be sufficient to appropriately modify the subroutine that generates the random numbers.

Sampling by the Inverse Transform Method: Continuous Distributions

Let $X \in (-\infty, +\infty)$ be an rv with cdf $F_X(x)$ and pdf $f_X(x)$, i.e.

$$F_X(x) = \int_{-\infty}^x f_X(x') dx' = \Pr\{X \leq x\} \quad (6)$$

Since $F_X(x)$ is a nondecreasing function, for any $y \in [0, 1)$, its inverse may be defined as

$$F_X^{-1}(y) = \inf\{x : F_X(x) \geq y\} \quad (7)$$

With this definition, we take into account the possibility that in some interval $(x_s, x_d]$ $F_X(x)$ is constant (and $f_X(x)$ zero), that is

$$F_X(x) = \gamma \quad \text{for } x_s \leq x \leq x_d \quad (8)$$

In this case, from the definition in equation (1) it follows that corresponding to the value γ , the minimum value x_s is assigned to the inverse function $F_X^{-1}(\gamma)$. This is actually as if $F_X(x)$ were not defined in $(x_s, x_d]$; however, this does not represent a disadvantage, since values in this interval can never be sampled because the pdf $f_X(x)$ is zero in that interval. Thus, in the following, we will suppose that the intervals $(x_s, x_d]$ (open to the left and closed to the right), in which $F_X(x)$ is constant, are excluded from the definition domain of the rv X . By so doing, the $F_X(x)$ will always be increasing (instead of nondecreasing).

We now show that it is always possible to obtain values $X \sim F_X(x)$ starting from values R sampled

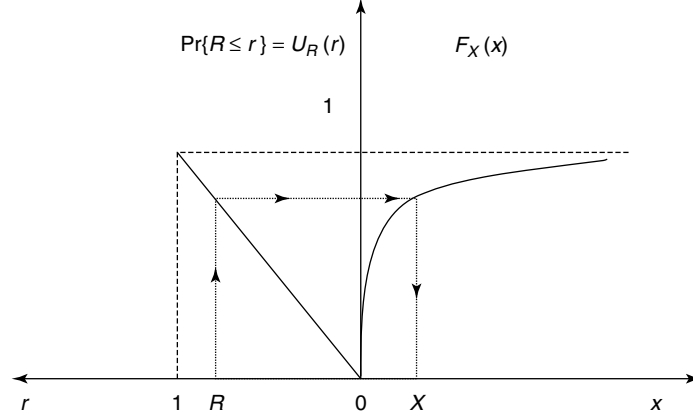


Figure 1 Inverse transform method: continuous distributions $R \sim U_R(r) = r$ in $[0, 1) \Rightarrow X \sim F_X(x)$ [©Lilole-Verlag GmbH, 2002.]

from the uniform distribution $U_R[0, 1)$ [1–6]. In fact, if R is uniformly distributed in $[0, 1)$, we have

$$\Pr\{R \leq r\} = U_R(r) = r \quad (9)$$

Corresponding to a *number* R extracted from a $U_R(r)$, we calculate the *number* $X = F_X^{-1}(R)$ and wonder about its distribution. As it can be seen in Figure 1, for the variable X we have

$$\Pr\{X \leq x\} = \Pr\{F_X^{-1}(R) \leq x\} \quad (10)$$

Because F_X is an increasing function, by applying F_X to the arguments at the rhs of equation (6), the inequality is conserved and from equation (8) we have

$$\Pr\{X \leq x\} = \Pr\{R \leq F_X(x)\} = F_X(x) \quad (11)$$

It follows that $X = F_X^{-1}(R)$ is extracted from $F_X(x)$. Furthermore, because $F_X(x) = r$,

$$\Pr\{X \leq x\} = \Pr\{R \leq r\} \quad (12)$$

In terms of the cdf,

$$U_R(R) = F_X(X) \quad \text{and} \quad R = \int_{-\infty}^X f_X(x') dx' \quad (13)$$

This is the fundamental relationship of the inverse transform method, which for any R value sampled from the uniform distribution U_R $[0, 1)$ gives the corresponding X value sampled from the $F_X(x)$ distribution (Figure 1). However, it often occurs that

the cdf $F_X(x)$ is noninvertible analytically, so that from equation (13) it is not possible to find $X \sim F_X(x)$ as a function of $R \sim U[0, 1)$. An approximate procedure that is often employed in these cases is the interpolation of $F_X(x)$ with a polygonal function and the inversion of equation (13) using the polygonal. Clearly, the precision of this procedure increases with the number of points of $F_X(x)$ through which the polygonal passes. The calculation of the polygonal is performed as follows:

- If the interval of variation of x is infinite, it is approximated by the finite interval (x_a, x_b) in which the values of the pdf $f_X(x)$ are sensibly different from zero: for example in case of the univariate **s-normal distribution** $N(\mu, \sigma^2)$ with mean value μ and variance σ^2 , this interval may be chosen as $(\mu - 5\sigma, \mu + 5\sigma)$.
- The interval $(0, 1)$ in which both $F_X(x)$ and $U_R(r)$ are defined is divided into n equal subintervals of length $1/n$ and the points $x_0 = x_a, x_1, x_2, \dots, x_n = x_b$ such that $F_X(x_i) = i/n$, ($i = 0, 1, \dots, n$) are found, e.g. by a numerical procedure.

At this point the sampling may start: for each R sampled from the distribution $U_R[0, 1)$ we compute the integer $i^* = \text{Int}(R \cdot n)$ and then obtain the corresponding X value by interpolating between the points $(x_{i^*}, (i^*/n))$ and $(x_{i^*+1}, (i^* + 1/n))$. For example, in case of a linear interpolation we have

$$X = x_{i^*} + (x_{i^*+1} - x_{i^*})(R \cdot n - i^*) \quad (14)$$

4 Life Distributions, Simulation of

For a fixed number n of points x_i upon which the interpolation is applied, the described procedure can be improved by interpolating with arcs of parabolas in place of line segments. The arcs can be obtained by imposing continuity conditions of the function and its derivatives at the points x_i (cubic **splines**). The expression of X as a function of R is in this case more precise, but more burdensome to calculate. Currently, given the ease with which it is possible to increase the RAM of the computers, to increase the precision it is possibly preferable to increase the number n of points and to use the polygonal interpolation: as a rule of thumb, a good choice is often $n = 500$.

Examples of Applications of the Sampling Methods

Uniform Distribution in the Interval (a, b)

An rv X is uniformly distributed in the interval (a, b) if

$$\begin{aligned} F_X(x) &= \frac{x-a}{b-a} & f_X(x) &= \frac{1}{b-a} & \text{for } a \leq x \leq b \\ &= 0 & &= 0 & \text{for } x < a \\ &= 1 & &= 0 & \text{for } x > b \end{aligned} \quad (15)$$

Substitution in equation (13) and solving with respect to X yields

$$X = a + (b - a)R \quad (16)$$

Exponential Distribution

An rv $X \in (0, \infty)$ is said to be exponentially distributed if its cdf $F_X(X)$ and pdf $f_X(X)$ are

$$\begin{aligned} F_X(x) &= 1 - e^{-\int_0^x \lambda(u) du}; \\ &= 0 \\ f_X(x) &= \lambda(x) e^{-\int_0^x \lambda(u) du} & \text{for } 0 \leq x \leq \infty \\ &= 0 & \text{otherwise} \end{aligned} \quad (17)$$

where $\lambda(\cdot)$ is the transition rate.

In the following, we shall refer to the exponential distribution of the random variable T , representing the *time to failure* of a component, with cdf $F_T(t)$

and pdf $f_T(t)$. In this case, the transition rate $\lambda(t)$ is called *hazard function* or *failure rate*.

With regard to sampling a realization t of the exponentially distributed random variable T , this can be obtained by solving equation (13), i.e.

$$\int_0^t \lambda(u) du = -\log(1 - R) \quad (18)$$

where the realization $R \sim U[0, 1)$.

Let us first consider the time homogeneous case, i.e. with constant λ . This situation corresponds to the *useful life* for which the component has been designed to operate and during which failures occur at random with no influence from the usage time. Correspondingly, equation (17) becomes

$$F(T) = 1 - e^{-\lambda t}; \quad f(T) = \lambda e^{-\lambda t} \quad (19)$$

Realizations of the associated exponentially distributed rv T are easily obtained from the inverse transform method. The sampling of a given number $N \gg 1$ of realizations is performed by repeating N times the following procedure:

- sample a realization of $R \sim U[0, 1)$
- compute $t = -(1/\lambda) \log(1 - R)$.

Weibull Distribution

A generalization of the above case in the time domain consists in assuming that the probability density of the occurrence of an event (e.g. the failure of a component), is time dependent. A case, commonly considered in practice is that in which the transition rate is of the kind

$$\lambda(t) = \beta \alpha t^{\alpha-1} \quad (20)$$

with $\beta > 0, \alpha > 0$. The corresponding distribution is called *Weibull's distribution* and was proposed in the 1950s by Weibull in the course of his studies on the strength of materials. The cdf and pdf of the Weibull distribution are as follows:

$$F(T) = 1 - e^{-\beta t^\alpha}, \quad f(T) = \alpha \beta t^{\alpha-1} e^{-\beta t^\alpha} \quad (21)$$

In the particular case of $\alpha = 1$, the Weibull distribution reduces to the exponential distribution with constant transition rate $\lambda = \beta$.

In practice, a realization t of the random variable T is sampled from the Weibull distribution through the following steps:

- sampling of a realization of the random variable $R \sim U[0, 1)$
- computation of $t = (-1/\beta) \ln(1 - R))^{1/\alpha}$.

References

- [1] Cashwell, E.D. & Everett, C.J. (1959). *A Practical Manual on the Monte Carlo Method for Random Walk Problems*, Pergamon Press, New York.
- [2] Rubinstein, R.Y. (1981). *Simulation and the Monte Carlo Method*, John Wiley & Sons.
- [3] Kalos, M.H. & Whitlock, P.A. (1986). *Monte Carlo methods, Vol. I: Basics*, John Wiley & Sons.
- [4] Lux, I. & Koblinger, L. (1991). *Monte Carlo Particle Transport Methods: Neutron and Photon Calculations*, CRC Press.
- [5] Dubi, A. (1999). *Monte Carlo Applications in Systems Engineering*, John Wiley & Sons.
- [6] Marseguerra, M. & Zio, E. (2002). *Basics of the Monte Carlo Method with Application to System Reliability*, LiLoLe-Verlag GmbH.

MARZIO MARSEGUERRA AND ENRICO ZIO

Simulation of the Beta-Stacy Process

Introduction

In Bayesian nonparametric inference, the parameter space consists of the class of all probability measures and the aim is to estimate a random distribution function or some of its functionals of statistical relevance. A common tool for constructing random distribution functions is represented by suitable transformations of stochastic processes. In particular, survival analysis is concerned with the estimation of lifetime distributions and, hence, one needs to define a process $\{F(t), t \in \mathbb{R}_+\}$ such that $F(0) = 0$ almost surely (a.s.), F is a.s. nondecreasing, a.s. right continuous, and $\lim_{t \rightarrow \infty} F(t) = 1$ a.s. A large section of the Bayesian nonparametric literature has considered the problem of *right-censored* data, which is a particular type of incomplete information: instead of observing the actual lifetime T^0 , one observes $T = \min(T^0, C)$ for some censoring time C taken independent of T^0 . The aim is to construct and make inference on random lifetime distributions which are suitable for application to this type of data. Typically, two requirements are in order: a mathematically tractable form for the posterior distribution of F given the observations and a simulation algorithm for drawing samples from the prior and the posterior distribution. As for the first point, a desirable property is the conjugacy of F with respect to right-censored data, which means that the posterior distribution belongs to the same class of random distribution functions as the prior. Conjugacy implies that, if we know how to sample from the prior, then the same type of algorithm can be used for the posterior, thus remarkably alleviating the computational burden.

This chapter is concerned with a class of priors termed *neutral-to-the-right* (NTR) processes, introduced by [1]. Other classes of stochastic processes for survival analysis are discussed in other articles (see **Survival Analysis, Nonparametric; Polya Trees and Their Use in Reliability and Survival Analysis**). NTR processes include the Dirichlet process (see [2]) as a special case and have conjugacy to right-censored observations as a distinctive feature. If we set Z to be the cumulative hazard

$Z(t) = -\log(1 - F(t))$, then an appealing characterization of NTR processes is that F is NTR if and only if Z has independent increments with $Z(0) = 0$ a.s., is nondecreasing a.s., right continuous a.s., and $\lim_{t \rightarrow \infty} Z(t) = \infty$ a.s. (see [3]). For such a process Z , there exist at most countably many fixed points of discontinuity $t_1, t_2, \dots$ with corresponding jumps $J_1, J_2, \dots$ in $\mathbb{R}_+$ having densities $f_{t_1}, f_{t_2}, \dots$. Then the process $Z_c(t) = Z(t) - \sum_{t_j \leq t} J_j$ has Lévy–Khinchine representation (see [4]):

$$-\log E \exp(-\phi Z_c(t)) = \int_0^\infty (1 - e^{-\phi z}) \nu_t(dz) \quad (1)$$

for every $\phi \geq 0$, where $\{\nu_t : t \geq 0\}$ is a family of Lévy measures on $\mathbb{R}_+$ satisfying

- L.1** $\nu_0 \equiv 0$ and $\int_0^\infty \min(1, z) \nu_t(dz) < \infty$ holds true for every $t > 0$;
- L.2** for every measurable B , $\nu_t(B)$ is nondecreasing and continuous in t .

The distribution of F is then completely characterized by the set of jump times $M = \{t_1, t_2, \dots\}$, the corresponding densities of the jumps $f_{t_1}, f_{t_2}, \dots$ and the Lévy measure $\nu_t(\cdot)$.

We consider a particular NTR process, namely the *beta-Stacy* process, introduced in [5], whose corresponding Lévy measure has the form

$$\nu_t(dz) = \frac{1}{1 - e^{-z}} \int_0^t e^{-z\beta(s)} \alpha(ds) dz \quad (2)$$

where $\beta(\cdot)$ is a nonnegative function, $\alpha(\cdot)$ is an absolutely continuous measure, and $\int_0^\infty \alpha(ds)/\beta(s) = \infty$. It is possible to give formulas for the beta-Stacy process in the presence of fixed points of discontinuities by means of a noncontinuous measure $\alpha(\cdot)$ (see [5]); however, for ease of exposition we do not consider this case and take $M = \emptyset$.

In the next section, we describe the properties of the beta-Stacy process and develop the posterior distribution given the right-censored data. In the section titled “Simulation of the Posterior Process”, we review methods to simulate from a general NTR process and, then, tailor them to the beta-Stacy case and conclude with some final remarks.

Property and Posterior of the Beta-Stacy

The beta-Stacy process is well suited for applications in survival analysis since it enjoys the conjugacy property and encompasses various NTR processes proposed in the literature. In fact, the Dirichlet process is a special case with $\alpha(\mathbb{R}_+) < \infty$ and $\beta(t) = \alpha[t, \infty)$, whereas the simple homogeneous process (see [6]) is obtained when $\beta(\cdot)$ is constant. Moreover, the beta-Stacy process is closely related to the celebrated beta-process (see [7]) used for modeling cumulative hazards: indeed the Lévy measure of the beta-process can be obtained from the one of the beta-Stacy in equation (2) via the transformation $v_t \circ g^{-1}(dz)$, where $g(z) = 1 - e^{-z}$.

The prior parameters $\alpha(\cdot)$ and $\beta(\cdot)$ have a straightforward interpretation in terms of the mean and the variance of F : if we set $\mu(t) = -\log[E(1 - F(t))]$ and $\lambda(t) = -\log[E(1 - F(t))^2]$, then

$$\begin{aligned} d\mu(t) &= \alpha(ds)/\beta(s), \quad \text{and} \\ d\lambda(t)/d\mu(t) &= 2 - [1 + \beta(t)]^{-1} \end{aligned} \quad (3)$$

This allows us to explicit prior opinions about the location and the uncertainty in terms of the Lévy measure, see [5]. As for the functionals of the beta-Stacy process, it is worth noting that the random mean $\int_0^\infty t dF(t)$ is studied in [8]: indeed, the random mean can be represented as the exponential functional of the corresponding cumulative hazard, $\int_0^\infty \exp[-Z(t)] dt$, and has the interpretation of the random expected lifetime.

We now provide the description of the posterior distribution of the beta-Stacy process, where we use a star as superscript to denote posterior quantities. Consider a sample $T_1^0, \dots, T_n^0$ of exchangeable observations directed by the beta-Stacy process, that is they are independent and identically distributed (i.i.d.) given F . Assume that $(T_1, \delta_1), \dots, (T_n, \delta_n)$ is observed, where $T_i = \min(T_i^0, C_i)$, $\delta_i = I\{T_i^0 \leq C_i\}$ and $C_1, \dots, C_n$ are the censoring times. Define the counting process $N(t)$, $t \in \mathbb{R}_+$, and the left-continuous at-risk process $Y(t)$, $t \in \mathbb{R}_+$, by

$$\begin{aligned} N(t) &= \sum_{i=1}^n I\{T_i \leq t \text{ and } \delta_i = 1\}, \\ Y(t) &= \sum_{i=1}^n I\{T_i \geq t\} \end{aligned} \quad (4)$$

In particular, $N\{t\}$ is the number of observed T_i^0 's at time point t . Given the data, the posterior distribution of $F(t)$, denoted by $[F(t) | \text{data}]$, is a beta-Stacy process with fixed jumps at $M^* = \{t_i : \delta_i = 1, i = 1, \dots, n\}$ with density given by

$$\begin{aligned} f_{t_i}^*(z) &\propto (1 - e^{-z})^{N\{t_i\}-1} \\ &\times e^{-z[\beta(t_i) + Y(t_i) - N\{t_i\}]} I\{z \in \mathbb{R}_+\} \end{aligned} \quad (5)$$

Alternatively, $f_{t_i}^*(z)$ can be described as the density of a jump random variable $J_{t_i}^*$ satisfying

$$1 - \exp\{-J_{t_i}^*\} \sim \text{beta}(N\{t_i\}, \beta(t_i) + Y(t_i) - N\{t_i\}) \quad (6)$$

where we write $\text{beta}(a, b)$ for a beta random variable of parameter a and b . The continuous component of the posterior process has Lévy measure:

$$v_t^*(dz) = \frac{1}{1 - e^{-z}} \int_0^t e^{-z[\beta(s) + Y(s)]} \alpha(ds) dz \quad (7)$$

which is still of the form in equation (2) with updated parameter $\beta^*(\cdot) = \beta(\cdot) + Y(\cdot)$. If we use the notation $\alpha^*(dt) = \alpha(dt) + dN(t)$, the Bayes estimate of F can be written as:

$$\begin{aligned} \widehat{F}(t) &= E(F(t) | \text{data}) \\ &= 1 - \prod_{[0, t]} \left(1 - \frac{\alpha^*(ds)}{\beta^*(s) + \alpha^*\{s\}} \right) \end{aligned} \quad (8)$$

where $\prod$ stands for the product integral, see [9]. For further details, see [5].

Simulation of the Posterior Process

In this section we review simulation schemes proposed in the literature for NTR processes, with a view toward simulating the posterior distribution of the beta-Stacy process. For the case of the Dirichlet process, see **Dirichlet Process, Simulation of**.

We consider a prior NTR process F with no fixed points of discontinuity. Furthermore, we consider a process for which $v_t(\mathbb{R}_+) = \infty$ for every $t > 0$, which is tantamount to requiring the process to jump infinitely often on any time interval. This clearly happens for the beta-Stacy process. Assume that the Lévy measure admits density, that is

$$v_t(dz) = dz \int_0^t k(z, s) ds \quad (9)$$

and recall that $k(z, s) ds = (1 - e^{-z})^{-1} e^{-z\beta(s)} \alpha(ds)$ for the beta-Stacy process.

In the posterior $Z_c^*(t)$ has updated Lévy density

$$k^*(s, z) = e^{-zY(s)} k(s, z) \quad (10)$$

while there are fixed jumps $J_{t_i}^*$ at noncensored observations with density given by

$$f_{t_i}^*(z) \propto (1 - e^{-z})^{N\{t_i\}} e^{-z[Y(t_i) - N\{t_i\}]} k(z, t) \quad (11)$$

see [3]. It is easy to show that this leads to the formulas for the posterior beta-Stacy process as described in the previous section.

Since simulating an NTR process F is equivalent to simulating the cumulative hazard process Z , sampling from $[F(t) | \text{data}]$ consists of two tasks: (a) simulate from an independent-increment process $Z_c^*(t)$ with Lévy density $k^*(z, s)$ and (b) draw random jump variates $\{J_{t_i}^* : \delta_i = 1, i = 1, \dots, n\}$ from densities given in equation (11). Note that (a) and (b) can be done separately because of the independent-increment property.

Simulation of the Continuous Component

Simulation methods based on the Lévy measure have been the object of major developments (see **Lévy Processes, Simulation of**). In the present setting, they have been used essentially in two ways: generation of the infinitely divisible increments of $Z_c^*(t)$ for a given time partition or a series representation of the whole process $\{Z_c^*(t), t \in \mathbb{R}_+\}$.

As for the first approach, we consider the *Damien, Laud, and Smith algorithm* (DLS95) in [10] and the *Walker and Damien algorithm* (WD98) in [11]. Methods that follow the second approach are the *Ferguson and Klass algorithm* (FK74) in [12] and the *Wolpert and Ickstadt algorithm* (WI98) in [13].

The idea behind DLS95 and WD98 is that, to simulate from the posterior distribution of an interval probability $[F(b) - F(a) | \text{data}]$, one ultimately needs to sample from the jumps $\{J_{t_i}^* : \delta_i = 1, t_i \in [a, b]\}$ and from the increments of $Z_c^*(t)$ for a given partition of the interval $[a, b]$. Let Z_Δ denote the increment of $Z_c^*(t)$ in a given time interval. The random variable Z_Δ is infinitely divisible with log-Laplace transform given by

$$-\log E[\exp(-\phi Z_\Delta)] = \int_0^\infty (1 - e^{-\phi z}) \nu_\Delta(dz) \quad (12)$$

where $\nu_\Delta(dz) = dz \int_\Delta k^*(s, z) ds$. It is a well-known result that each infinitely divisible distribution can be represented as a stochastic integral with respect to a Poisson process, in the present setting

$$Z_\Delta = \int_0^\infty z dP(z) \quad (13)$$

where $P(\cdot)$ stands for a **Poisson process** with intensity measure $\nu_\Delta(dz)$. Note that $\nu_\Delta(B)$, for $B \in \mathcal{B}(\mathbb{R}_+)$, may be interpreted as the Poisson arrival rate of increments for the process $Z_c^*(t)$ in the interval Δ of sizes in B . Since $\nu_\Delta(dz)$ is not a finite measure, in general it is not possible to simulate exactly from the correct distribution of $P(\cdot)$: $P(\cdot)$ has infinitely many jumps in any neighborhood of zero. The two algorithms obviate this by resorting to different approximation techniques, but both require simulation from a jump-size distribution recovered from the measure $\nu_\Delta(dz)$.

The DLS95 method generates a finite number m of jump levels $\{z_i\}$ from the distribution $\gamma(dz) \propto h(z) \nu_\Delta(dz)$, where $h(z)$ is a nonnegative function such that $\gamma = \int_0^\infty h(z) \nu_\Delta(dz) < \infty$. Note that, by property (L. 1), it is sufficient that $h(z) \leq \min(1, z)$. Then, for $i = 1, \dots, m$, take a Poisson variate q_i with rate $\lambda_i = \gamma/[m h(z_i)]$ and approximate Z_Δ by $\sum_{i=1}^m z_i q_i$. This entails an approximation of $P(\cdot)$ with another Poisson process concentrated on the singletons $\{z_i\}$ with intensity measure $\frac{1}{m} \sum_i \delta_{z_i}(dz) \lambda_i$. For details see Damien *et al.* [10], who implemented it for the gamma-process, the simple homogeneous and the Dirichlet process, whereas the case of the beta-process was considered in [14]. We next give the algorithm for the beta-Stacy process, following [5]. In the sequel let $\text{Pois}(\lambda)$ denote a Poisson random variable of rate λ .

DLS95 Algorithm

- Take $h(z) = 1 - e^{-z}$ so that $\gamma = \int_0^\infty \int_\Delta e^{-z[\beta(s) + Y(s)]} \alpha(ds) dz$.
- Generate m independent draws $\{z_i\}$ from the distribution

$$G(dz) = \gamma^{-1} \int_\Delta e^{-z[\beta(s) + Y(s)]} \alpha(ds) dz \quad (14)$$

- For $i = 1, \dots, m$, generate q_i from $Q_i \sim \text{Pois}(\gamma/m(1 - e^{-z_i}))$.
- Set $Z_\Delta \approx \sum_{i=1}^m z_i q_i$. □

4 Simulation of the Beta-Stacy Process

The WD98 method consists in cutting off the part of small jumps from the Poisson integral in equation (13), following an idea that first appeared in [15, Section 2]. For $\epsilon > 0$ such that $\lambda_\epsilon = \nu_\Delta(\epsilon, \infty) < \infty$, consider $Z_\Delta = X_\epsilon + Y_\epsilon$ where X_ϵ is the sum of the jump times of $P(\cdot)$ in (ϵ, ∞) and Y_ϵ is the sum of the jump times in $(0, \epsilon]$. Then, X_ϵ has a.s. a finite number of jumps that can be simulated exactly by generating jump levels from a density proportional to $\frac{d}{dz} \nu_\Delta((\epsilon, z])$ for $z \in (\epsilon, \infty)$. Finally, approximate Z_Δ with X_ϵ , treating Y_ϵ as variable error. In fact, Y_ϵ can be taken as small as needed for small ϵ . The algorithm for the beta-Stacy process is as follows:

WD98 Algorithm

- Take $\lambda_\epsilon = \int_\epsilon^\infty \int_\Delta e^{-z[\beta(s)+Y(s)]} \alpha(ds)$ and $G_\epsilon(z) = \lambda_\epsilon^{-1} \int_\epsilon^z \int_\Delta e^{-z[\beta(s)+Y(s)]} \alpha(ds)$.
- Generate q from $Q \sim \text{Pois}(\lambda_\epsilon)$.
- Generate q independent draws $\{z_i\}$ from c.d.f. $G_\epsilon(z)$.
- Set $Z(\Delta) \approx \sum_{i=1}^q z_i$. $\square$

For details, see [11], which shows how to simulate the z_i 's variates by using a Gibbs sampler with auxiliary variables, and how to calculate the normalizing constant λ_ϵ .

Consider now the second class of simulation methods, based on series representations. Exploiting the fact that $Z_c^*(t)$ is a pure jump process, algorithms FK72 and WI98 are based on representing $Z_c^*(t)$ as the sum

$$Z_c^*(t) = \sum_i J_i I\{U_i \leq t\} \quad (15)$$

where, for $i = 1, 2, \dots$, J_i is a random jump level and U_i is a random location. This idea has been used to simulate a whole trajectory of the process by [12, 13 and 16]. See also, [17], who reviewed the different methods of representing the Poisson process of the jumps underlining the Lévy–Khinchine formula in equation (1). Since for $\nu_\infty(\mathbb{R}_+) = \infty$ the sum in equation (15) has infinitely many terms, in practice one has to truncate the series and approximate the process with a finite number of jumps. The error of approximation is determined by both the number of considered jumps and their sizes. Both algorithms aim at simulating $\{J_i\}$ and $\{U_i\}$ via inversion of the Lévy measure, but they differ on whether to draw first the jump sizes or the jump locations.

The FK72 algorithm is as follows. Define the jumps $\{J_i\}$ by

$$\text{Prob}(J_1 \leq z_1) = \exp(-\nu_\infty^*[z_1, \infty)) \quad (16)$$

$$\begin{aligned} \text{Prob}(J_i \leq z_i \mid J_{i-1} \leq z_{i-1}) \\ = \exp(-\nu_\infty^*[z_i, \infty) + \nu_\infty^*[z_{i-1}, \infty)) \end{aligned} \quad (17)$$

where $z_i \leq z_{i-1}$. This can be done by generating successive jumps $\tau_1, \tau_2, \dots$ from a Poisson process of unit rate (partial sums of independent standard exponential variates), then J_i is taken as the solution z of the equation: $\tau_i = \nu_\infty^*[z, \infty)$, that is via inversion of the Lévy measure $\nu_\infty^*[\cdot, \infty)$. Then the process $Z_c^*(t)$ is represented as the series:

$$Z_c^*(t) = \sum_i J_i I\{\omega_i \leq n_t(J_i)\} \quad (18)$$

where $\omega_1, \omega_2, \dots$ are i.i.d. random numbers from the uniform distribution on $[0, 1]$ and $n_t(z) = \int_0^t k^*(z, s) ds / \int_0^\infty k^*(z, s) ds$. Note that, by property (L.2), $n_t(z)$ is nondecreasing and, thus, the jumps are generated in decreasing order; see [16] for more details. For the special case of the beta-Stacy process, the algorithm can be detailed as follows:

FK72 Algorithm

- Write $M(z) = \nu_\infty^*[z, \infty)$ (see equation (7)) and note that

$$M(z) = \int_0^\infty B_{e^{-z}}(\beta(s) + Y(s), 0) \alpha(ds) \quad (19)$$

where $B_x(a, b) = \int_0^x s^{a-1} (1-s)^{b-1} ds$ denotes the incomplete beta-function.

- Generate m successive jump times $\{\tau_i\}$ of a Poisson process of unit rate.
- Perform a **Monte Carlo** integration of $M(z)$: generate K independent draws $\{s_j\}$ from a distribution function $\alpha(ds)/\alpha(\mathbb{R}_+)$, then set

$$M(z) \approx K^{-1} \sum_{j \leq K} \alpha(\mathbb{R}_+) B_{e^{-z}}(\beta(s_j) + Y(s_j), 0) \quad (20)$$

- For $i = 1, \dots, m$, generate J_i as the value z solution of $\tau_i = M(z)$.
- Generate m independent draws $\{\omega_i\}$ from the uniform distribution on $[0, 1]$ and set, for

$i = 1, \dots, m$, U_i as the solution s of the equation

$$\omega_i = \frac{\int_0^s e^{-z[\beta(u)+Y(u)]} \alpha(du)}{\int_0^\infty e^{-z[\beta(u)+Y(u)]} \alpha(du)} \quad (21)$$

- Set $Z_c^*(t) \approx \sum_{i=1}^m J_i I\{U_i \leq t\}$. $\square$

The WI98 algorithm aims at generating jump locations first. Assume that we may write the Lévy measure in the form

$$\nu_t^*(dz) = dz \int_0^t k^*(z, s) dz ds = dz \int_0^t \eta(z, s) \Pi(ds) \quad (22)$$

for $\Pi(ds)$ a probability measure on $\mathbb{R}_+$ from which samples may be drawn, and so $\eta(z, s) = ds k^*(z, s) / \Pi(ds)$. This form of $\nu_t^*(\cdot)$ actually covers all the Lévy processes appearing in the Bayesian literature. For example, for the posterior beta-Stacy process we may take $\Pi(ds) = \alpha(ds) / \alpha(\mathbb{R}_+)$ (see the Monte Carlo step above) and $\eta(z, s) = \alpha(\mathbb{R}_+) (1 - e^{-z})^{-1} \exp(-z[\beta(s) + Y(s)])$. The algorithm proceeds by generating m random locations $\{U_i\}$ from $\Pi(ds)$. Then, for $i = 1, \dots, m$, J_i is set as the solution z of the equation $\tau_i = \int_z^\infty \eta(x, U_i) dx$, where $\tau_1, \dots, \tau_m$ are successive jump times of a standard Poisson process at unit rate. The approximation error decreases with the number of jumps considered, but the jumps are not generated in decreasing order, unless one takes $\Pi(ds)$ constant. For more details, see [18], where the algorithm is adapted to sample from the beta-process, the extended gamma-process (see [19]), and the simple homogeneous process. See also [18] for a comparison of WI98 with DLS95, and [16] for a comparative discussion with FK72. The details for the WI98 algorithm in the beta-Stacy case are as follows:

WI98 Algorithm

- Use the incomplete beta-function to write

$$\int_\epsilon^\infty \eta(z, s) dz = \alpha(\mathbb{R}_+) B_{e^{-\epsilon}}(\beta(s) + Y(s), 0) \quad (23)$$

- Generate m independent draws $\{U_i\}$ from the distribution $\alpha(ds) / \alpha(\mathbb{R}_+)$.
- Generate m successive jump times $\{\tau_i\}$ of a Poisson process of unit rate.

- For $i = 1, \dots, m$, set J_i as the solution z of the equation

$$\tau_i = \alpha(\mathbb{R}_+) B_{e^{-\tau_i}}(\beta(S_i^c) + Y(S_i^c), 0) \quad (24)$$

- Set $Z_c(t) \approx \sum_{i=1}^m J_i^c I\{S_i^c \leq t\}$. $\square$

Simulation of the Jump Component

There exist a few methods for sampling from densities known up to a constant of proportionality like the density in equation (11), following the emergence of Markov chain Monte Carlo methods (see **Markov Chain Monte Carlo, Introduction**). Indeed, it is possible to use techniques for generating random variates as in [10, Appendix A], consisting in a Gibbs sampling with introduction of auxiliary variables, see [11, Section 3.1]. Note, however, that this is not required for the beta-Stacy process, since the jumps of $[F(t) | \text{data}]$ follow a beta-distribution of equation (6). In practice one has to generate beta-variates $\{e_i\}$, with $e_i \sim \text{beta}(N\{t_i\}, \beta(t_i) + Y(t_i) - N\{t_i\})$, for each noncensored observation and, then, set $J_{t_i}^* = 1 - \log(1 - e_i)$.

Conclusions

We have described four different algorithms for simulating from the beta-Stacy process. Note that such algorithms can be adapted to any NTR process and that the Lévy–Khinchine formula constitutes the theoretical background for these contributions. Another method for simulating the increment of a nondecreasing, independent-increment process is in [15], where a different characterization of infinitely divisible distributions *via* shot noise processes is exploited. See also [16] and [17] for series representations of Lévy processes based on this method, and [20] for an application to the extended gamma-process.

References

- [1] Doksum, K. (1974). Tailfree and neutral random probabilities and their posterior distributions, *Annals of Probability* **2**, 183–201.
- [2] Ferguson, T.S. (1973). A Bayesian analysis of some non-parametric problems, *Annals of Statistics* **1**, 209–230.
- [3] Ferguson, T.S. (1974). Prior distributions on spaces of probability measures, *Annals of Statistics* **2**, 615–629.
- [4] Sato, K. (1999). *Lévy Processes and Infinitely Divisible Distributions*, Cambridge University Press.

- [5] Walker, S.G. & Muliere, P. (1997). Beta-Stacy processes and a generalization of the Polya urn scheme, *Annals of Statistics* **25**, 1762–1780.
- [6] Ferguson, T.S. & Phadia, E.G. (1979). Bayesian non-parametric estimation based on censored data, *Annals of Statistics* **7**, 163–186.
- [7] Hjort, N.L. (1990). Non-parametric Bayes estimators based on beta processes in models for life history data, *Annals of Statistics* **18**, 1259–1294.
- [8] Epifani, I., Lijoi, A. & Prünster, I. (2003). Exponential functionals and means of neutral-to-the-right priors, *Biometrika* **90**, 791–808.
- [9] Gill, R.D. & Johansen, S. (1990). A survey of product-integration with a view toward application to survival analysis, *Annals of Statistics* **18**, 1501–1555.
- [10] Damien, P., Laud, P.W. & Smith, A.F.M. (1995). Random variate generation approximating infinitely divisible distributions with application to Bayesian inference, *Journal of the Royal Statistical Society. Series B* **57**, 547–564.
- [11] Walker, S.G. & Damien, P. (1998). A full Bayesian nonparametric analysis involving a neutral to the right process, *Scandinavian Journal of Statistics* **25**, 669–680.
- [12] Ferguson, T.S. & Klass, M.J. (1972). A representation of independent increment processes without Gaussian components, *Annals of Mathematical Statistics* **43**, 1634–1643.
- [13] Wolpert, R.L. & Ickstadt, K. (1998). Gamma/Poisson random field models for spatial statistics, *Biometrika* **85**, 251–267.
- [14] Damien, P., Laud, P.W. & Smith, A.F.M. (1996). Implementation of Bayesian non-parametric inference based on beta processes, *Scandinavian Journal of Statistics* **23**, 27–36.
- [15] Bondesson, L. (1982). On simulation from infinitely divisible distributions, *Advances in Applied Probability* **14**, 855–869.
- [16] Walker, S.G. & Damien, P. (2000). Representation of Lévy processes without Gaussian components, *Biometrika* **87**, 477–483.
- [17] Rosiński J. (2001). Series representation of Lévy processes from the perspective of point processes, in *Lévy Processes – Theory and Applications*, O.E. Barndorff-Nielsen, T. Mikosch & S.I. Resnick, eds, Birkhauser, Boston, pp. 401–415.
- [18] Wolpert R.L. & Ickstadt K. (1998). Simulation of Lévy random fields, in *Practical Nonparametric and Semiparametric Bayesian Statistics, (Lecture Notes in Statistics 133)*, D. Dey, P. Muller & D. Sinha, eds, Springer, New York, 227–242.
- [19] Dykstra R.L. & Laud P.W. (1981). A Bayesian non-parametric approach to reliability. *Annals of Statistics* **9**, 356–367.
- [20] Laud P.W., Smith A.F.M. & Damien P. (1996). Monte Carlo methods for approximating a posterior hazard rate process. *Statistics and Computing* **6**, 77–83.

Related Articles

Dirichlet Process, Simulation of; Lévy Processes, Simulation of; Markov Chain Monte Carlo, Introduction; Polya Trees and Their Use in Reliability and Survival Analysis Survival Analysis, Nonparametric.

PIERPAOLO DE BLASI

Dirichlet Process, Simulation of

The Dirichlet Process and Mixtures

Bayesian nonparametric **prior distributions** have become a valuable tool within the context of hierarchical modeling (*see Survival Analysis, Nonparametric; Polya Trees and Their Use in Reliability and Survival Analysis*). In contrast to traditional empirical distribution- or rank-based methods, nonparametric modeling in the Bayesian paradigm consists of constructing prior distributions on the space of all possible distribution functions or some large subset of it. It entails putting a prior distribution on the *distribution* of data without use of a family indexed by a finite-dimensional parameter as in parametric Bayesian models. This difference allows the data distribution to escape unwanted constraints of parametric families, such as unimodality. The resulting flexibility is especially useful when modeling mixture distributions, because the composition of the mixture does not have to be specified deterministically.

The most popular nonparametric prior distribution, the Dirichlet process, was formally introduced in the work of Ferguson [1, 2]. If we consider an arbitrary distribution function, G , on the positive real numbers, and any finite partition of $(0, \infty)$ given by $0 < t_1 < t_2 < \dots < t_{k-1} < \infty$, then the Dirichlet process prior results in the representation of the uncertainty in G via the Dirichlet distribution as $Y = (Y_1, \dots, Y_k) \sim \mathcal{D}(\pi_1, \dots, \pi_k)$ where

$$Y_1 = G(t_1), Y_i = G(t_i) - G(t_{i-1}), \\ i = 2, \dots, k-1, Y_k = 1 - G(t_{k-1}) \quad (1)$$

$\pi_i = \alpha(G_0(t_i) - G_0(t_{i-1}))$, $\alpha > 0$, and G_0 is a distribution function on $(0, \infty)$. The development in [1] is much more general, defining the Dirichlet process for distributions on arbitrary probability spaces, thus including the usual multivariate case of finite-dimensional Euclidean space with the Borel σ algebra.

While Ferguson's work provided the theoretical foundation for modeling using the Dirichlet process, we will discuss a particular extension – the Dirichlet process mixture (DPM) – that has proved useful in

applications. The difference between the two can be seen readily within the context of a simple hierarchical model for data $\mathbf{Y} = (Y_1, \dots, Y_n)$, written as

$$Y_i | \theta_i \stackrel{\text{ind}}{\sim} L(Y_i | \theta_i) \\ \theta_i | G \stackrel{\text{ind}}{\sim} G \\ G | \alpha, G_0 \sim \mathcal{DP}(\alpha G_0) \quad (2)$$

Here $L(Y_i | \theta_i)$ represents the likelihood of an observed data point conditional on a parameter θ_i , with unknown distribution G , on which we place a Dirichlet process prior involving the two parameters α and G_0 . The latter is a centering distribution for the process, while α is a scalar variability parameter controlling how much one allows G to deviate from G_0 . The likelihood $L(Y_i | \theta_i)$ has a parametric form, for example $\text{Normal}(\mu, \sigma^2)$ where $\theta = (\mu, \sigma^2)$. In this case, G is a bivariate distribution on $(-\infty, \infty) \times (0, \infty)$. Thus, the model is a DPM of parametric distributions in the likelihood.^a While the DPM model is contained in the work of Lo [5], its computational feasibility and usefulness was initially demonstrated in a paper [6] by Escobar.

The DPM model has been extended in many ways, including the use of a prior distribution for α and taking G_0 to be a member of a parametric family, $G_0(\cdot | \kappa)$, with a random parameter κ . For the sake of simplicity, we will consider the case where α and κ are fixed. Incorporating prior distributions on these parameters is straightforward (for example, see [7] for placing a prior on α and [8] for placing one on κ).

To make inferences from data, we need to address the posterior distribution conditional on Y . The objects of inference here are $\theta_1, \dots, \theta_n$ as well as G , the distribution governing future values of θ . From a simulation perspective, inference for G can be obtained via samples from the conditional distribution of θ_{n+1} . In the next sections we therefore concentrate on sampling $\theta_1, \dots, \theta_n, \theta_{n+1} | \mathbf{Y}$.

When G_0 is a member of a parametric family that is conjugate to the likelihood in $L(Y_i | \theta_i)$, computation is simplified. In this case we refer to the DPM model as being *conjugate*. The complementary case is called a *nonconjugate DPM model*.

Simulation in Conjugate DPM Models

The first computational algorithm in DPM models originated in the work of Escobar[6], and Escobar

2 Dirichlet Process, Simulation of

and West [9], within the context of normal mixture models and density estimation. The basic theory behind the algorithm relates to the fact that the infinite-dimensional parameter G in equation (2) can be integrated out and the model written in terms of the marginal prior distribution for θ . Results in [10–12] imply that the distribution of θ is almost surely discrete and is given by

$$\theta_1 \sim G_0, \quad \theta_{k+1} | \theta_1, \dots, \theta_k \sim \frac{\alpha}{\alpha + k} G_0 + \frac{1}{\alpha + k} \sum_{i=1}^k \delta_{\theta_i}(\theta_{k+1}), \quad k = 1, \dots \quad (3)$$

where $\delta_{\theta_i}(\cdot)$ denotes a unit point mass at θ_i . This means that given $(\theta_1, \dots, \theta_k)$, θ_{k+1} is either a distinct value (with probability $\alpha/(\alpha + k)$) drawn from G_0 , or one randomly drawn from the existing θ_i ($i = 1, \dots, k$). This prior can then be applied to carry out posterior inference by iteratively sampling θ_i within a Gibbs loop (*see Markov Chain Monte Carlo, Introduction*) via the full conditionals

$$\theta_i | \theta_{-i}, \mathbf{Y} \sim q_0 G_i(\theta_i) + \sum_{j=1, j \neq i}^n q_j \delta_{\theta_j}(\theta_i) \quad (4)$$

where θ_{-i} is $(\theta_1, \dots, \theta_{i-1}, \theta_{i+1}, \dots, \theta_n)$, $q_j = L(y_j | \theta_j)$, $G_i(\theta_i)$ is the parametric posterior distribution conditional on the i th observation given by $dG_i(\theta_i) \propto L(y_i | \theta_i) dG_0(\theta_i)$, and $q_0 = \int L(Y_i | \theta_i) \times dG_0(\theta_i)$. If $L(Y_i | \theta_i)$ and G_0 are conjugate, the integral in q_0 can be derived explicitly.

Many of the other algorithms developed for conjugate DPM models are very similar in spirit to Escobar's initial contribution. The differences lie in exploiting the inherent structure of the θ_i 's due to the discrete nature of the Dirichlet process. Rather than dealing directly with each θ_i individually, the current state of the $\theta = (\theta_1, \dots, \theta_n)$ can be represented in terms of a *configuration* [13].

Definition Let $K \leq n$ denote the number of distinct values in θ . Denote by $\phi = (\phi_1, \dots, \phi_K)$ these distinct values in some arbitrary order. Groups of elements of θ made by the ϕ_j 's are called *clusters*. Then define a vector of cluster membership pointers $s = (s_1, \dots, s_n)$ by $s_i = j \Leftrightarrow \phi_j = \theta_i$. A configuration c of θ is the pair (ϕ, s) . In addition, let n_j equal the number of θ_i 's equal to ϕ_j . Since the indices in ϕ

are arbitrary, there are $k!$ configurations representing a given θ .

By separating distinct values and membership pointers, this definition allows one to implement the sampling in equation (4) more efficiently by using the conditionals $\phi | s, Y$ and $s | \phi, Y$ or by integrating out the ϕ . We note here that $\phi_j | s, Y$ are independent random variates, each respectively from the posterior distribution obtained under prior G_0 and likelihood terms $L(Y_i | \theta_i)$ contributed by i such that $s_i = j$. Such algorithms based on configurations were utilized by Neal [14], MacEachern [15], and Bush and MacEachern [16]. We describe in some detail a more general and recent version in the section titled “Nonconjugate DPM Model Simulation”.

Split–Merge MCMC Techniques

Although the above algorithms produce **Markov chain** Monte Carlo (MCMC) chains that converge to the correct posterior distribution, they can run into difficulties when the true mixture components are not well separated. An example would be a mixture of normal distributions where the differences in component means are small relative to each component's variance. In this situation, the above algorithms can mix poorly because they update the cluster membership pointers in s sequentially. In order to separate two close components, the algorithm would need to move single observations into a new cluster. However, these new clusters can represent intermediate states with low associated probability [17]. In response to this difficulty, Jain and Neal [18] and Green and Richardson [19] have developed split–merge techniques whereby groups of observations are allowed to be updated simultaneously, thereby skipping low-probability intermediate states.

Jain and Neal begin by selecting at random two elements of θ whose membership pointers we denote as s_i and s_j . If $s_i \neq s_j$ then propose merging the two clusters into one. If $s_i = s_j$, then propose splitting that cluster into two. (In actual implementation, Jain and Neal advocate using a restricted Gibbs sampling scan in order to propose more likely splits of the single component.) For either of these situations, the decision to update the configuration with either the merge or split proposal is then based on a Metropolis–Hastings update [20–22]. If we denote the current

configuration as $c = (\phi, s)$, and the proposed state as $c^* = (\phi^*, s^*)$, then the chain would move to s^* with probability,

$$\alpha(s^*, s) = \min \left[1, \frac{q(s|s^*) \pi(s^*) p(y|s^*)}{q(s^*|s) \pi(s) p(y|s)} \right] \quad (5)$$

where $q(s^*|s)$ denotes the proposal probability of moving from s to s^* and $\pi(s)$ represents the prior probability for the entire configuration s . Expressions for $\pi(\cdot)$ and $q(\cdot, \cdot)$ are available in [18, 23].

The split–merge algorithm has been implemented in the conjugate situation as well as in what Jain and Neal term as *conditionally conjugate* DPM models (see [23] for details). The primary difficulty lies in evaluating the likelihood in the Metropolis update, namely,

$$p(y|s) = \prod_{i=1}^n \int L(y_i|\theta) H_{i,s_i}(\theta) \quad (6)$$

Here H_{i,s_i} is the posterior distribution for theta under the prior G_0 and all observations y_k such that $k \in \{k : k < i, s_k = s_i\}$. If the last set is empty, H_{i,s_i} simply equals G_0 . Unless the above integral is analytically tractable, as in conjugate models, the likelihood evaluation can quickly become nontrivial.

The split–merge algorithm of Green and Richardson originates from a reversible jump framework. The general sampling scheme proceeds in the same manner as with Jain and Neal’s algorithm where split and merge proposals are proposed and accepted with appropriate probabilities in a Metropolis–Hastings update. However, Green and Richardson construct their proposals based on general reversible jump guidelines, which we outline below.

Consider the DPM model in equation (2) where θ_i is p dimensional. As before, let $\phi = (\phi_1, \dots, \phi_K)$ represent the collection of distinct values of θ . During a split proposal we need to assign elements of a randomly chosen cluster $g_j = \{i : s_i = j\}$ to two clusters generically denoted g_{j-} and g_{j+} as well as choose corresponding ϕ_{j-} and ϕ_{j+} . In a reverse manner, a merge move combines two randomly chosen clusters and their parameter values into $g\phi_j$.

The general reversible jump scheme calls for the following steps.

1. Methods for combining ϕ_{j-} and ϕ_{j+} into ϕ_j during a proposed merge and generating ϕ_{j-} and

ϕ_{j+} from ϕ_j during a split. This is accomplished via p conservation conditions

$$m_l(\phi_j) = w_- m_l(\phi_{j-}) + w_+ m_l(\phi_{j+}), \quad l = 1, \dots, p \quad (7)$$

where the vector function $m : \mathcal{R}^p \rightarrow \mathcal{R}^p$ is invertible and $w_- + w_+ = 1$. Typically, m_l are **moments** of the distribution $L(y|\phi)$. For the merge proposal the above definition suffices, taking $w \sim \beta(n_{j-} + \omega, n_{j+} + \omega)$, where ω represents a fixed tuning parameter. For splitting, one uses auxiliary variables $u = (u_1, \dots, u_p)$ and a bijection between (ϕ_j, u) and (ϕ_{j-}, ϕ_{j+}) to arrive at ϕ_{j-} and ϕ_{j+} , with $w_- \sim U(0, 1)$.

2. Selecting and constructing split/merge clusters. First draw two membership pointers, s_i and s_k at random.

(a) If $s_i = s_k = j$ propose a split by

$$P(s_i = j-) = \frac{w_- L(Y_i|\phi_{j-})}{w_- L(Y_i|\phi_{j-}) + w_+ L(Y_i|\phi_{j+})}, \quad \text{for each } i \ni s_i = j \quad (8)$$

- (b) If $s_i \neq s_k$, merge clusters g_{s_i} and g_{s_k} into one with an associated parameter obtained as in Step 1.

3. Calculate the Metropolis–Hastings acceptance probability for the proposed new configuration c^*
 - (a) For split proposals, $\alpha(c^*, c) = \min(1, A)$.
 - (b) For merge proposals, $\alpha(c^*, c) = \min(1, A^{-1})$.

Here A is problem specific and depends, for example, on $L(y|\phi)$ and the Jacobian of the bijection in Step 1. For the mixture of normal likelihoods, Green and Richardson [19] provide an expression for A .

The split–merge techniques of Jain and Neal and Green and Richardson should be viewed as enhancements rather than replacements for the sequential approaches. In fact, Jain and Neal advocate using a combination of both procedures, especially in the case of complicated mixtures. The split–merge algorithm enables large component movements, which can then be fine-tuned by following with a sequential update. This combination is straightforward (although potentially computationally demanding) within a Gibbs sampling scheme, as one can simply

iterate between split–merge and sequential updates, perhaps randomly with some predetermined probabilities.

Nonconjugate DPM Model Simulation

In the case of nonconjugate DPM models, MacEachern and Müller [13] propose an algorithm that avoids calculating q_0 , instead replacing it by likelihood evaluations. They begin by extending ϕ to n components $(\phi_1, \dots, \phi_n)$ and defining two distinct sets,

$$\begin{aligned}\phi &= \{\phi_F, \phi_E\}, & \phi_F &= \{\phi_1, \dots, \phi_K\}, \\ \phi_E &= \{\phi_{K+1}, \dots, \phi_n\}\end{aligned}\quad (9)$$

Here ϕ_F can be thought of as full clusters, meaning that each θ_i is currently equal to one of the ϕ_j for $j = 1, 2, \dots, K$ whereas ϕ_E represents empty clusters with no current members among the θ_i . For updating the configuration, while maintaining the above structure of $\phi = (\phi_F, \phi_E)$, the authors then introduce a *no-gaps* constraint on the labeling of clusters, meaning that for $j = 1, 2, \dots, K$, they require $n_j \geq 1$. Given this constraint, the algorithm involves the following steps.

1. For $i = 1, 2, \dots, n$, sample $(s_i | s_{-i}, \phi_F)$.
 - (a) Case 1: If $n_{s_i} = 1$, leave s_i unchanged with probability $(k-1)/k$. Otherwise, relabel the current cluster labeled s_i as cluster K . Then assign θ_i to one of the clusters $j = 1, 2, \dots, K$ via

$$\begin{aligned}P(s_i = j | s_{-i}, \phi_F, y_i) \\ = c_K^{-1} (n_j^-)^{\delta(1 \leq j < K)} \left(\frac{\alpha}{K}\right)^{\delta(j=K)} L(y_i | \phi_j)\end{aligned}\quad (10)$$

where n_j^- are cluster sizes without θ_i , i.e., $n_j - I(s_i = j)$, and $c_K = \sum_{j=1}^K P(s_i = j | s_{-i}, \phi_F, y_i)$. Note that unless s_i is placed in the K th cluster, the configuration loses a distinct value, meaning $K = K - 1$.

- (b) Case 2: If $n_{s_i} > 1$ then draw ϕ_{K+1} from G_0 and sample s_i with probabilities $(j = 1, 2, \dots, K + 1)$,

$$P(s_i = j | s_{-i}, \phi_F, y_i) = c_{K+1}^{-1} (n_j^-)^{\delta(1 \leq j \leq K)}$$

$$\times \left(\frac{\alpha}{K+1}\right)^{\delta(j=K+1)} L(y_i | \phi_j) \quad (11)$$

If s_i is put into cluster $K + 1$, this adds a component to the configuration, implying that $K = K + 1$.

2. Sample $(\phi | s)$. For $\phi_j \in \phi_F$, we now have K independent parametric models, $y_i | \phi \sim L(y_i | \phi_j) \forall i \ni s_i = j$ and $\phi_j \sim G_0$, $j = 1, \dots, K$, each with sample size n_j . Thus ϕ_j can be updated from the corresponding posteriors. However, when $L(y_i | \theta_i)$ and $G_0(\theta_i)$ are nonconjugate, this update could require sampling from non-standard densities, therefore necessitating general sampling approaches like adaptive-rejection Sampling [24, 25] or the Metropolis–Hastings algorithm (*see Hierarchical Markov Chain Monte Carlo (MCMC) for Bayesian System Reliability*).

Neal [26] proposed an alternative to the no-gaps algorithm that also avoids, in the cluster updating step, the need for integration in the nonconjugate setting. He uses the conditional prior for s_i as a proposal density in constructing a Metropolis–Hastings update for s_i . Conditional on the current configuration, a new cluster membership, s_i^* can be proposed as

$$P(s_i^* = j | s_{-i}) = \frac{n_j^-}{n - 1 + \alpha}, \quad \forall j \ni s_k = j \text{ for some } k = 1, \dots, n \quad (12)$$

$$P(s_i^* \neq s_k \forall k \neq i | s_{-i}) = \frac{\alpha}{n - 1 + \alpha} \quad (13)$$

Note that in the latter case, a new $\phi_{s_i^*}$ needs to be generated from G_0 . Given this proposal density, the acceptance probability reduces to a convenient ratio of likelihoods:

$$\alpha(s_i^*, s_i) = \min \left[1, \frac{L(Y_i | \phi_{s_i^*})}{L(Y_i | \phi_{s_i})} \right] \quad (14)$$

The final step in the algorithm is exactly Step 2 in the no-gaps algorithm.

Another interesting and simple algorithm for the nonconjugate setting, based on an entirely different approach that uses a latent variable formulation, is due to Walker and Damien [27]. They rewrite the

target conditional posterior distribution in equation (4) as

$$p(\theta_i | \phi^-, \mathbf{Y}) \propto L(Y_i | \theta_i) \left\{ \alpha G_0(\theta_i) + \sum_{j=1, j \neq i}^n \delta_{\theta_j}(\theta_i) \right\} \quad (15)$$

and then introduce latent variables u_i , which are conditionally distributed $\text{Uniform}(0, L(Y_i | \theta_i))$. With $u = (u_1, \dots, u_n)$ this leads to the joint conditional distribution

$$p(\theta, u | \mathbf{Y}) \propto \prod_{i=1}^n \left\{ \delta(u_i < L(Y_i | \theta_i)) \left(\alpha G_0(\theta_i) + \sum_{j < i} \delta_{\theta_j}(\theta_i) \right) \right\} \quad (16)$$

Clearly, the introduction of u still leads to the same conditional distribution as in equation (15). The update for θ involves the following two steps. First sample all $u_i | \theta_i$ independently from uniform densities on $(0, L(Y_i | \theta_i))$, and then sample from the conditional for each θ_i available in the convenient form

$$p(\theta_i | u_i, \mathbf{Y}) \propto \alpha G_0(\theta_i) \delta(L(Y_i | \theta_i) > u_i) + \sum_{L(Y_i | \theta_j) > u_i} \delta_{\theta_j}(\theta_i) \quad (17)$$

Walker and Damien [27] provide two examples of the method.

Extensions

The Dirichlet process has been generalized in many different directions. These include neutral-to-the-right processes [28], Polya trees [2, 4, 29], Lévy-driven **Markov processes** [30], and others. From the computational standpoint, generalizations that have more commonalities with the basic techniques developed for the Dirichlet process are stick-breaking priors [31], species sampling models [32], and use of the weighted Chinese restaurant process in [32]. Dirichlet process techniques have also been incorporated in and adapted to other classes of models, in particular, multiplicative intensity models [33]. The continued popularity of methods stemming from the Dirichlet

process appears to be driven by its flexibility and computational feasibility.

End Notes

^aDirichlet process mixtures (the focus of this chapter) are distinct from mixtures of Dirichlet processes (MDP), as introduced in the work of Antoniak [3]. MDPs begin with G as distribution for data, and express uncertainty in G by mixing Dirichlet processes with different centering measures, making G_0 itself random. Hanson and Johnson [4] elaborate the difference between MDP and DPM. Unfortunately, the terminology has not yet become standardized and often appears in the literature in conflicting ways.

References

- [1] Ferguson, T. (1973). A Bayesian analysis of some nonparametric problems, *The Annals of Statistics* **1**(2), 209–230.
- [2] Ferguson, T. (1974). Prior distributions on spaces of probability measures, *The Annals of Statistics* **2**(4), 615–629.
- [3] Antoniak, C. (1974). Mixtures of Dirichlet processes with applications to Bayesian nonparametric problems, *The Annals of Statistics* **2**(6), 1152–1174.
- [4] Hanson, T. & Johnson, W. (2002). Modeling regression error with a mixture of Polya trees, *Journal of the American Statistical Association* **97**(460), 1020–1033.
- [5] Lo, A.Y. (1984). On a class of Bayesian nonparametric estimates: I. Density estimates, *The Annals of Statistics* **12**(2), 351–357.
- [6] Escobar, M. (1994). Estimating normal means with a Dirichlet process prior, *Journal of the American Statistical Association* **89**(425), 268–277.
- [7] West, M. (1992). Hyperparameter estimation in Dirichlet process mixture models, Technical Report 03, Duke University-Institute of Statistics and Decision Sciences.
- [8] MacEachern, S. & Müller, P. (2000). Efficient MCMC schemes for robust model extensions using encompassing Dirichlet process mixture models, in *Robust Bayesian Analysis*, D.R. Insua & F. Ruggeri, eds, Springer-Verlag, pp. 295–315.
- [9] Escobar, M. & West, M. (1995). Bayesian density estimation and inference using mixtures, *Journal of the American Statistical Association* **90**(430), 577–588.
- [10] Blackwell, D. (1973). Discreteness of Ferguson selections, *The Annals of Statistics* **1**(2), 356–358.
- [11] Blackwell, D. & MacQueen, J.B. (1973). Ferguson distributions via pólya urn schemes, *The Annals of Statistics* **1**(2), 353–355.

- [12] Sethuraman, J. (1994). A constructive definition of Dirichlet priors, *Statistica Sinica* **4**, 639–650.
- [13] MacEachern, S. & Müller, P. (1998). Estimating mixture of Dirichlet process models, *Journal of Computational and Graphical Statistics* **7**(2), 223–238.
- [14] Neal, R.M. (1992). Bayesian mixture modeling, in *Maximum Entropy and Bayesian Methods*, C.R. Smith, G.J. Erickson & P.O. Neudorfer, eds, Kluwer Academic Publishers, Dordrecht, pp. 197–211.
- [15] MacEachern, S.N. (1994). Estimating normal means with a conjugate-style Dirichlet process prior, *Communications in Statistics: Simulation and Computation* **23**, 727–741.
- [16] Bush, C.A. & MacEachern, S.N. (1996). A semiparametric Bayesian model for randomised block designs, *Biometrika* **83**(2), 275–285.
- [17] Celeux, G., Hurn, M. & Robert, C.P. (2000). Computational and inferential difficulties with mixture posterior distributions, *Journal of the American Statistical Association* **95**, 957–970.
- [18] Jain, S. & Neal, R. (2004). A split-merge Markov chain Monte Carlo procedure for the Dirichlet process mixture model, *Journal of Computational and Graphical Statistics* **13**, 158–182.
- [19] Green, P.J. & Richardson, S. (2001). Modelling heterogeneity with and without the Dirichlet process, *Scandinavian Journal of Statistics* **28**, 355–375.
- [20] Metropolis, M., Rosenbluth, A., Rosenbluth, M., Teller, H. & Teller, E. (1953). Equation of state calculations by fast computing machines, *Journal of Chemical Physics* **21**, 1087–1092.
- [21] Hastings, W.K. (1970). Monte Carlo sampling methods using Markov chains and their applications, *Biometrika* **57**, 97–109.
- [22] Chib, S. & Greenberg, E. (1995). Understanding the Metropolis-Hastings algorithm, *The American Statistician* **49**, 327–335.
- [23] Jain, S. & Neal, R. (2005). Splitting and merging components of a nonconjugate Dirichlet process mixture model, Technical Report 0507, University of Toronto.
- [24] Gilks, W. & Wild, P. (1992). Adaptive rejection sampling for Gibbs sampling, *Applied Statistics* **41**(2), 337–348.
- [25] Gilks, W. & Tan, K.K.C. (1995). Adaptive rejection metropolis sampling within Gibbs sampling, *Applied Statistics* **44**(4), 455–472.
- [26] Neal, R.M. (2000). Markov chain sampling methods for Dirichlet process mixture models, *Journal of Computational and Graphical Statistics* **9**, 249–265.
- [27] Walker, S.G. & Damien, P. (1998). Methods for Bayesian nonparametric inference involving stochastic processes, in *Practical Nonparametric and Semiparametric Bayesian Statistics*, D. Dey, P. Müller & D. Sinha, eds, Springer-Verlag, New York, pp. 243–254.
- [28] Doksum, K. (1974). Tailfree and neutral random probabilities and their posterior distributions, *The Annals of Probability* **2**, 183–201.
- [29] Lavine, M. (1992). Some aspects of Polya tree distributions for statistical modelling, *The Annals of Statistics* **20**, 1222–1235.
- [30] Nieto-Barajas, L.E. & Walker, S.G. (2004). Bayesian nonparametric survival analysis via Lévy driven Markov processes, *Statistica Sinica* **14**(4), 1127–1146.
- [31] Ishwaran, H. & James, L.F. (2001). Gibbs sampling methods for stick-breaking priors, *Journal of the American Statistical Association* **96**(2), 161–173.
- [32] Ishwaran, H. & James, L.F. (2003). Generalized weighted Chinese restaurant processes for species sampling mixture models, *Statistica Sinica* **13**, 1211–1235.
- [33] Ishwaran, H. & James, L.F. (2004). Computational methods for multiplicative intensity models using weighted gamma processes: proportional hazards, marked point processes, and panel count data, *Journal of the American Statistical Association* **99**(465), 175–190.

Related Articles

Convergence and Mixing in Markov Chain Monte Carlo; Hierarchical Markov Chain Monte Carlo (MCMC) for Bayesian System Reliability; Markov Chain Monte Carlo, Introduction; Polya Trees and Their Use in Reliability and Survival Analysis; Survival Analysis, Nonparametric; Simulation of the Beta-Stacy Process.

PURUSHOTTAM W. LAUD AND NICHOLAS M. PAJEWSKI

Smoothed Function Estimation for Censored Data

Introduction

Consider a nonnegative random variable X with *distribution function* (d.f.) $F(\cdot)$ and *density function* $f(\cdot)$. In the context of reliability and survival analysis, among many associated functions of interest are the *survival function* (s.f.) $S(t) = 1 - F(t)$, the *hazard function* $h(t) = f(t)/S(t) = -(d/dt) \log S(t)$, the *cumulative hazard function* (chf)

$$H(t) = \int_0^t h(y) dy = -\log S(t), \quad t \geq 0 \quad (1)$$

and the *mean residual life*

$$m(t) = E(T - t | T > t) = \frac{\left(\int_t^\infty S(y) dy \right)}{S(t)} \quad (2)$$

Since the knowledge of the density function or the d.f. provides information about other functions, the basic strategy for finding their smooth estimators is to find a smooth estimator of the d.f. and use it to find smooth estimators of other derived functions. The setup considered here is that of right censoring in which instead of observing the random sample $X_1, \dots, X_n$, from the distribution F , we observe the censored lifetimes $Z_i = \min(X_i, Y_i)$ and the indicator variables $\delta_i = I(X_i \leq Y_i)$, where $Y_1, Y_2, \dots, Y_n$ represents a sequence of independent censoring times with d.f. G . In this case we have the nonparametric maximum-likelihood estimator of $S(t)$ as given by Kaplan and Meier [1], also known as the *product-limit* (PL) estimator. It may be written as (see Bhattacharjee and Sen [2])

$$\hat{S}_n(t) = \begin{cases} 1 & \text{for } 0 \leq t < Z_{1:n}^* \\ \prod_{i \leq j} \left(\frac{n - k_i}{n - k_i + 1} \right) & \text{for } Z_{j:n}^* \leq t < Z_{j+1:n}^*; 1 \leq j \leq M_n \end{cases} \quad (3)$$

where $Z_{i:n}^*$, $i = 1, 2, \dots, M_n$, denote the M_n ordered *failure points*, i.e., $Z_{j:n}^* = Z_{k_j:n}$ ($Z_{0:n}^* = 0$, $Z_{M_n+1:n}^* =$

∞), with $Z_{i:n}$, $i = 1, 2, \dots, n$ denoting the ordered statistics of the sample $Z_1, \dots, Z_n$. This estimator has jump discontinuities at the failure points and therefore is not suitable for producing estimates of $f(t)$ and $h(t)$, though estimators of $H(t)$ and $m(t)$ may be obtained by replacing $S(t)$ by $\hat{S}_n(t)$ in equations (1) and (2).

On the other hand, Aalen [3] exploited counting-process martingale theory to propose and study the following *unbiased* estimator of $H(t)$, called the *Nelson–Aalen estimator*:

$$\hat{H}_n(t) = \int_0^t \frac{dN_{1n}(s)}{R_n(s)} = \sum_{i: Z_{i:n} \leq t} \frac{\delta_{(i:n)}}{n - i + 1} \quad (4)$$

where $N_{1n}(t) = \sum_{i=1}^n \delta_i I(Z_i \leq t)$ and $R_n(t) = \sum_{i=1}^n I(Z_i \geq t)$ are the *failure* (i.e., uncensored) and the *number-at-risk* counting processes, respectively, and $\delta_{(i:n)}$ denotes the value of δ_i corresponding to the ordered statistic $Z_{i:n}$.

Notably, however, none of the above estimators produce *smooth* estimators of $H(t)$ and $m(t)$. Thus, the main emphasis in the literature on smooth function estimation so far has been on *density* or *hazard-rate* estimation, that is, on smoothing either $S_n(t)$ or $H_n(t)$. This has been influenced by advances and popularity of the *kernel density estimator* Rosenblatt [4] and Parzen [5] as discussed below for some of the procedures under each category. Recall that a kernel estimator for uncensored data is given by

$$\hat{f}_n(x) = (\omega_n)^{-1} \int_0^\infty k\left(\frac{x-s}{\omega_n}\right) d\hat{F}_n(s) \quad (5)$$

where $\hat{F}_n(\cdot)$ is the empirical distribution function (edf), $\omega_n(>0)$, known as the *bandwidth*, chosen so that $\omega_n \rightarrow 0$ but $n\omega_n \rightarrow \infty$, as $n \rightarrow \infty$; $k(\cdot)$ is termed the *kernel function* and is typically assumed to be a symmetric probability density function (pdf) with zero mean and unit variance.

Density Estimation: Smoothing of PL Estimator

Blum and Susarla [6] proposed and studied an adaptation of the kernel estimator to censored data, by considering an equivalent problem for the uncensored observations. Lo *et al.* [7] proved the consistency and asymptotic normality of the kernel density estimators using an improved version of an almost-sure representation of a *modified* $\hat{F}_n(x) = 1 - \hat{S}_n(x)$ originally

2 Smoothed Function Estimation for Censored Data

due to Lo and Singh [8], of the form

$$\hat{F}_n(x) = F(x) + n^{-1} \sum_{i=1}^n \xi_i(x) + O\left(\frac{\log n}{n}\right) \quad (6)$$

Here $\xi_i(x) = S(x)\zeta(X_i, \delta_i, x)$, where

$$\zeta(z, \delta, x) = -g(z \wedge x) + [\bar{L}(x)]^{-1} I\{z \leq x, \delta = 1\} \quad (7)$$

with $g(y)$, $\bar{L}(x)$, and $L_1(x)$ defined as

$$g(y) = \int_0^y [\bar{L}(s)]^{-2} dL_1(s) \quad (8)$$

$\bar{L}(x) = (1 - F(x))(1 - G(x))$ and $L_1(x) = P(X_i \leq x, \delta_i = 1)$. A smooth estimator of the hazard function is then given by $\hat{h}_n(x) = \hat{f}_n(x)/(1 - \hat{F}_n(x))$, consistency and other large sample properties of which also follow (see Lo *et al.* [7]).

Hazard Estimation: Smoothing of Nelson–Aalen Estimator

Ramlau-Hansen [9], Yandell [10], and Tanner and Wong [11] all considered kernel smoothing of the Nelson–Aalen estimator, providing a smooth estimator of $h(t)$ as given by

$$\begin{aligned} \hat{h}_n(t) &= (\omega_n)^{-1} \int_0^\infty k\left(\frac{t-s}{\omega_n}\right) d\hat{H}_n(s) \\ &= (\omega_n)^{-1} \sum_{i=1}^n \frac{\delta_{(i:n)}}{n-i+1} k\left(\frac{t-Z_{i:n}}{\omega_n}\right) \end{aligned} \quad (9)$$

While Ramlau-Hansen [9] used martingale methods to obtain pointwise asymptotic normality, and Yandell [10] obtained asymptotic distribution of *maximum absolute deviation* via a strong approximation, they restricted the observations to a compact set. Tanner and Wong [11], however, avoid this restriction in an interesting way – by tackling the estimator directly via *Hájék projection*.

A list of available results concerning the above smooth density estimator and other variants along with hazard rate is contained in the survey article by Padgett and McNichols [12].

Bandwidth Selection and Other Methods

Choice of the bandwidth, $\omega_n > 0$, is the most crucial issue in any kind of smoothing. Typically, ω_n has an

optimal rate of $O(n^{-1/5})$. Marron and Padgett [13] show that *least-squares cross-validation* produces asymptotically optimal bandwidths in the case of density estimation. Patil [14] proves the same for hazard-rate estimation. Further, Patil [15] obtains the convergence rate of these bandwidths, as well as that of the corresponding *integrated squared error*, in a unified framework that covers all the four cases: density or hazard rate, uncensored or censored data.

As for alternative estimation methods, Antoniadis *et al.* [16] study wavelet-based hazard estimators, which have more recently been generalized by Brunel and Comte [17], who consider general projection estimator methods that include trigonometric series, polynomials, and wavelets. These authors prove optimality of the estimators in an asymptotic minimax sense.

Kernel estimators suffer from the following two kinds of *boundary bias*. First, as observed by Silverman [18] and others, they may provide positive mass outside the support of the target function. Even otherwise, as remarked in Wand *et al.* [19], the estimator in equation (5) “works well for densities that are not far from Gaussian in shape”; however, “it can perform very poorly when the shape is far from Gaussian”, especially near the boundaries. Second, it may fail to estimate discontinuity at boundaries, such as in the case of exponential distribution.

Different approaches proposed to deal with these problems in the uncensored data case can also be adapted to the censored data case. Below we mention some of the prominent methods: the transformation methods of Wand *et al.* [19], Ruppert and Wand [20], and Marron and Ruppert [21]; the method of Bagai and Prakasa Rao [22] based on a kernel k with support $\mathbb{R}^+$, which, however, fails to use the whole data; the interesting, asymmetric-kernel estimators of Chen [23] using gamma kernels, and of Scaillet [24], using inverse Gaussian and reciprocal inverse Gaussian kernels, whose variances, however, blow up at $x = 0$. These last methods have been recently adapted by Bouezmarni and Rombouts [25] for density estimation in the censored-data case, but the problems mentioned above still persist in this case.

Chaubey and Sen [26], on the other hand, proposed a density estimator as the derivative of a smooth version of the edf, by adapting Hille’s smoothing lemma (see [27]) in a stochastic setup.

This, in contrast to the proposal of Bagai and Prakasa Rao [22], uses the whole data and easily generalizes to the censored-data case, as studied in Chaubey and Sen [28] for estimation of density, survival, hazard, and chfs, and in Chaubey and Sen [29] for mean residual life function. It does not suffer from the problems mentioned above. In the next section we demonstrate the use of Hille's lemma for obtaining smooth estimators.

The Chaubey–Sen Estimator

Here we shall consider, as in Chaubey and Sen [28], the following modification of $\hat{S}_n(t)$ as in Efron [30]:

$$S_n(t) = \begin{cases} \hat{S}_n(t) & \text{for } t < Z_{n:n} \\ 0 & \text{for } t \geq Z_{n:n} \end{cases} \quad (10)$$

Further, we state Hille's lemma which underlies the construction of all the estimators in this section:

Lemma (Feller [27, Lemma 1, Section VII.1]) *Let $u(t)$ be a bounded, continuous function on $\mathbf{R}^+$. Then*

$$e^{-\lambda t} \sum_{k \geq 0} u\left(\frac{k}{\lambda}\right) \frac{(\lambda t)^k}{k!} \rightarrow u(t), \text{ as } \lambda \rightarrow \infty \quad (11)$$

uniformly in any finite interval $\mathbf{J}$ contained in $\mathbf{R}^+$.

Substituting Efron's $S_n(t)$ [30] for $u(t)$ in the above lemma, Chaubey and Sen [28] obtained the following smooth survival function estimator:

$$\tilde{S}_n(t) = \sum_{k=0}^N S_n\left(\frac{k}{\lambda_n}\right) p_k(\lambda_n t) \quad (12)$$

where $p_k(\mu)$, $k = 0, 1, 2, \dots$ denote the Poisson probabilities, $p_k(\mu) = e^{-\mu} \mu^k / k!$, $k = 0, 1, 2, \dots$, $N = [\lambda_n Z_{n:n}]$, and λ_n is suitably chosen (see Chaubey and Sen [31] for a discussion on the choice of λ_n).

This is an extension of the idea of Chaubey and Sen [26, 32] to the case of randomly censored data. Chaubey and Sen [28] used truncated Poisson weights constructed in a suitable manner. It was later recognized for the uncensored case that (see Chaubey and Sen [33]) the truncated weights may not be suitable in forming the associated estimator of the *mean residual life* function. Moreover, the untruncated weights are asymptotically equivalent to truncated weights.

Hence, we exclude the discussion of the truncated weights here and present smooth estimators based on the untruncated weights. Chaubey and Sen [28] show that the absolute error between smooth and nonsmooth estimators of the s.f. over any compact set is of the order $O(n^{-3/4}(\log n)^{1+\delta})$, for some $\delta > 0$, under the existence of the first moment of F .

Noting that $\tilde{S}_n(t)$ defined above is smooth, bounded, and continuously differentiable with respect to $t \in \mathbf{R}^+$, we obtain the following estimators of the associated functionals:

$$\begin{aligned} \tilde{F}_n(t) &= 1 - \tilde{S}_n(t), \quad \tilde{f}_n(t) = \lambda_n \sum_{k=0}^N p_k(\lambda_n t) \\ &\times \left[S_n\left(\frac{k}{\lambda_n}\right) - S_n\left(\frac{k+1}{\lambda_n}\right) \right] \end{aligned} \quad (13)$$

$$\tilde{h}_n(t) = \frac{\tilde{f}_n(t)}{\tilde{S}_n(t)}, \quad \tilde{m}_n(t) = \frac{1}{\tilde{S}_n(t)} \int_t^\infty \tilde{S}_n(y) dy \quad (14)$$

Further, an alternative hazard-rate estimator is obtained by smoothing the Nelson–Aalen estimator:

$$\tilde{H}_n(t) = \sum_{k=0}^N H_n\left(\frac{k}{\lambda_n}\right) p_k(\lambda_n t) \quad (15)$$

$$\begin{aligned} \tilde{h}_n^{\text{NA}}(t) &= \left(\frac{d}{dt}\right) \tilde{H}_n(t) = \lambda_n \sum_{k=0}^N p_k(\lambda_n t) \\ &\times \left[H_n\left(\frac{k+1}{\lambda_n}\right) - H_n\left(\frac{k}{\lambda_n}\right) \right] \end{aligned} \quad (16)$$

where $H_n(\cdot)$ is as in Equation (4). The computational form of the smooth mean residual life estimator is provided below (see Chaubey and Sen [28]):

$$\tilde{m}_n(t) = \frac{\sum_{k=0}^N P_k(\lambda_n t) S_n\left(\frac{k}{\lambda_n}\right)}{\lambda_n \sum_{k=0}^N p_k(t \lambda_n) S_n\left(\frac{k}{\lambda_n}\right)} \quad (17)$$

where $N = [\lambda_n Z_{n:n}]$ and, $P_k(\mu)$ denotes the cumulative probability for the Poisson distribution with mean μ , namely, $P_k(\mu) = \sum_{r=0}^k p_r(\mu) = \sum_{r=0}^k e^{-\mu} \times$

4 Smoothed Function Estimation for Censored Data

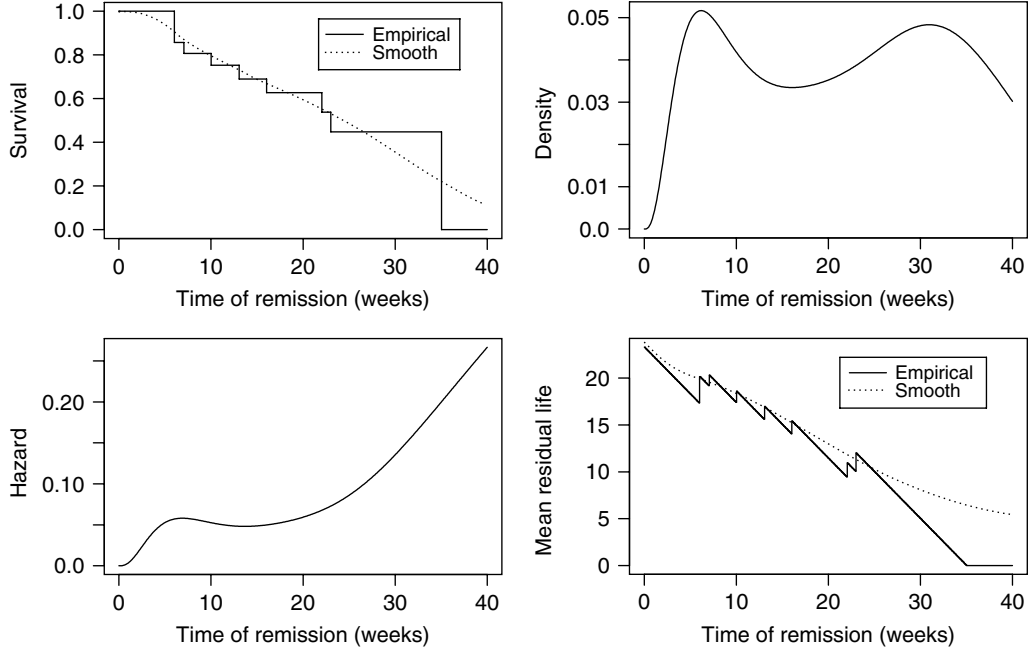


Figure 1 Kaplan–Meier survival function estimates and the corresponding Chaubey–Sen smooth estimates for survival, density, hazard, and mean residual life for remission duration data

$(\mu^r/r!)$. The asymptotics of $\tilde{m}_n(t)$ are obtained in Chaubey and Sen [29]. The smoothing procedure discussed here is also found useful in smooth isotonic estimation of density, hazard, and mean residual life functions by Chaubey and Sen [34].

The choice of λ_n is important in the study of the asymptotic properties of the estimators $\tilde{S}_n$, $\tilde{f}_n$, $\tilde{h}_n$, and $\tilde{H}_n$. Chaubey and Sen [28] recommend the choice $\lambda_n = n/Z_{n:n}$, for a practical purpose, which is appropriate for the censoring and lifetimes having infinite support and finite first **moments** (in this case $N = n$). However, through further numerical investigations (see Chaubey and Sen [31]), the choice $\lambda_n = (n - 1)/Z_{n:n}$ is found more suitable. This is the choice used in the numerical illustration below.

A Numerical Example

Here we illustrate the smooth estimators using the data in Gehan [35], in which the drug 6-mercaptopurine (6-MP) was compared to a placebo with respect to the ability to maintain remission in acute leukemia patients. We have only considered the

data for the 6-MP drug where 21 subjects were used and the maximum censored time was 35 weeks. The remission times (in weeks) and the corresponding PL survival probabilities are given in Table 1.

The value of λ_n chosen is $20/35 = 0.5714$. The plots of the smooth survival function, the density function, the hazard function, and the mean residual life function are presented in Figure 1. We see that, whereas the smooth survival curve stays very close to the Kaplan–Meier survival curve for the whole observed range, the smoothed mean residual life values seem to stay above the empirical mean residual life values, especially near the largest observation. The reason for this may lie in the fact that the empirical mean residual life estimator may not be defined beyond the largest failure.

Table 1 Remission times and Kaplan–Meier (KM) survival estimates for 6-MP study

Remission time(t_j)	6	7	10	13	16	22	23
$\hat{S}_n(t_j)$	0.857	0.807	0.753	0.690	0.627	0.538	0.448

Acknowledgments

The first two authors acknowledge partial funding of this research by their research grants from NSERC of Canada.

References

- [1] Kaplan, E.L. & Meier, P. (1958). Nonparametric estimation from incomplete observations, *Journal of the American Statistical Association* **53**, 457–481.
- [2] Bhattacharjee, M.C. & Sen, P.K. (1995). Kolmogorov-Smirnov type tests for NB(W)UE alternatives under censoring schemes, in *Analysis of Censored Data, IMS Lecture-Monograph Series*, H.L. Koul & J.V. Deshpande, eds, Institute of Mathematical Statistics, Hayward, California, Vol. 27, pp. 25–38.
- [3] Aalen, O. (1978). Nonparametric inference for a family of counting processes, *Annals of Statistics* **6**, 701–726.
- [4] Rosenblatt, M. (1956). Remarks on some nonparametric estimates of density functions, *Annals of Mathematical Statistics* **27**, 832–837.
- [5] Parzen, E. (1962). On estimation of probability density and mode, *Annals of Mathematical Statistics* **33**, 1065–1070.
- [6] Blum, J.R. & Susarla, V. (1980). Maximal deviation theory of density and failure rate function estimates based on censored data, in *Multivariate Analysis V*, P.R. Krishnaiah, ed, North Holland, Amsterdam, pp. 213–222.
- [7] Lo, S.-H., Mack, Y.P. & Wang, J.-L. (1989). Density and hazard rate estimation for censored data via strong representation of the Kaplan-Meier estimator, *Probability Theory and Related Fields* **80**, 461–473.
- [8] Lo, S.-H. & Singh, K. (1986). The product-limit estimator and the bootstrap: some asymptotic representations, *Probability Theory and Related Fields* **71**, 455–465.
- [9] Ramlau-Hansen, H. (1983). Smoothing counting process intensities by means of kernel functions, *Annals of Statistics* **11**, 453–466.
- [10] Yandell, B.S. (1983). Nonparametric inference for rates with censored survival data, *Annals of Statistics* **11**, 1119–1135.
- [11] Tanner, M.A. & Wong, W.H. (1983). The estimation of the hazard function from randomly censored data by the kernel method, *Annals of Statistics* **11**, 989–993.
- [12] Padgett, W.J. & McNichols, D.T. (1984). Nonparametric density estimation from censored data, *Communications in Statistics – Theory and Methods* **A13**, 1581–1611.
- [13] Marron, J.S. & Padgett, W.J. (1987). Asymptotically optimal bandwidth selection for kernel density estimators from randomly right-censored samples, *Annals of Statistics* **15**, 1520–1535.
- [14] Patil, P.N. (1993a). Bandwidth choice for nonparametric hazard rate estimation, *Journal of Statistical Planning and Inference* **35**, 15–30.
- [15] Patil, P.N. (1993b). On the least squares cross-validation bandwidth in hazard rate estimation, *Annals of Statistics* **21**, 1792–1810.
- [16] Antoniadis, A., Grégoire, G. & Nason, G. (1999). Density and hazard rate estimation for right-censored data by using wavelet methods, *Journal of the Royal Statistical Society. Series B* **61**, 63–84.
- [17] Brunel, E. & Comte, F. (2005). Penalized contrast estimation of density and hazard rate with censored data, *Sankhyā* **67**, 441–475.
- [18] Silverman, B.W. (1986). *Density Estimation for Statistics and Data Analysis*, Chapman and Hall, London.
- [19] Wand, M.P., Marron, J.S. & Ruppert, D. (1991). Transformations in density estimation, *Journal of the American Statistical Association* **86**, 343–361.
- [20] Ruppert, D. & Wand, M.P. (1992). Correcting for kurtosis in density estimation, *Australian Journal of Statistics* **34**, 19–29.
- [21] Marron, J.S. & Ruppert, D. (1994). Transformations to reduce boundary bias in kernel density estimation, *Journal of the Royal Statistical Society. Series B* **56**, 653–671.
- [22] Bagai, I. & Prakasa Rao, B.L.S. (1996). Kernel type density estimates for positive valued random variables, *Sankhyā* **A57**, 56–67.
- [23] Chen, S.X. (2000). Probability density function estimation using Gamma kernels, *Annals of Institute of Statistical Mathematics* **52**, 471–480.
- [24] Scaillet, O. (2004). Density estimation using inverse Gaussian and reciprocal inverse Gaussian kernels, *Journal of Nonparametric Statistics* **16**, 217–226.
- [25] Bouezmarni, T. & Rombouts, J.V. (2006). *Density and Hazard Rate Estimation for Censored and α -mixing Data Using Gamma Kernels*, Cahiers de recherche IEA-06-16, HEC, Institut d'Économie Appliquée, Montreal.
- [26] Chaubey, Y.P. & Sen, P.K. (1996). On smooth estimation of survival and density functions, *Statistics and Decisions* **14**, 1–22.
- [27] Feller, W. (1966). *An Introduction to Probability Theory and its Applications*, John Wiley & Sons, New York-London-Sydney, Vol. II.
- [28] Chaubey, Y.P. & Sen, P.K. (1998). On smooth functional estimation under random censorship, in *Frontiers in Reliability 4, Series on Quality, Reliability and Engineering Statistics*, A.P. Basu, S.K. Basu & S. Mukhopadhyay, eds, World Scientific, Singapore, pp. 83–97.
- [29] Chaubey, Y.P. & Sen, A. (2007). Smooth estimation of mean residual life under random censoring, in *Beyond Parametrics in Interdisciplinary Research: Festschrift to P K Sen*, N. Balakrishnan, E. Peña & M.J. Silvapulle, eds, IMS Lecture Notes and Monograph Series IMS, Hayward.
- [30] Efron, B. (1967). The two sample problem with censored data, *Proceedings of Fifth Berkeley Symposium on Mathematical Statistics and Probability* **48**, 831–853.
- [31] Chaubey, Y.P. & Sen, P.K. (2000). Tail-behavior of survival function and their smooth estimators, in *Perspectives in Statistical Sciences*, A.K. Basu, J.K. Ghosh,

- P.K. Sen & B.K. Sinha, eds, Oxford, New Delhi, pp. 86–101.
- [32] Chaubey, Y.P. & Sen, P.K. (1997). On smooth estimation of hazard and cumulative hazard functions, in *Frontiers in Probability*, S.P. Mukherjee, S.K. Basu & B.K. Sinha, eds, Narosa, New Delhi, pp. 91–99.
- [33] Chaubey, Y.P. & Sen, P.K. (1999). On smooth estimation of mean residual life, *Journal of Statistical Planning and Inference* **75**, 223–236.
- [34] Chaubey, Y.P. & Sen, P.K. (2002). Smooth isotonic estimation of density, hazard and MRL functions, *Calcutta Statistical Association Bulletin* **52**, 205–208.
- [35] Gehan, E.A. (1965). A generalized Wilcoxon test for comparing arbitrarily singly-censored samples, *Biometrika* **52**, 203–223.

Further Reading

Wand, M.P. & Jones, M.C. (1995). *Kernel Smoothing*, Chapman and Hall, London.

Related Articles

Density and Failure Rate Estimation; Intensity Functions for Nonhomogeneous Poisson Processes; Probability Density Function (PDF).

YOGENDRA P. CHAUBEY, ARUSHARKA SEN
AND PRANAB K. SEN

Convergence and Mixing in Markov Chain Monte Carlo

Convergence and Mixing of Markov Chain Monte Carlo

The aim of Markov chain **Monte Carlo** (MCMC) simulation is to sample a target distribution of interest, π , or to summarize it, e.g., compute mean, variances, or probabilities with respect to π . A Markov chain having π as its unique invariant and limiting distribution is constructed and simulated until convergence to π . For a detailed description of various algorithms available to construct such Markov chains (Metropolis–Hastings, **Gibbs sampler**, etc.) and for conditions that guarantee uniqueness of the stationary distribution and convergence to it, see **Markov Chain Monte Carlo, Introduction**, whose notation is adopted here as much as possible. In the sequel, finite state spaces are referred to, to keep things simple, but the considerations made here extend to general state spaces. A Markov chain will be identified with the corresponding transition matrix (in finite state spaces) or kernel (in general state spaces). Given that the time available for simulation is limited, researchers search for strategies to improve the performance of the MCMC sampler they implement. Performance can be enhanced in two main ways: either by increasing the speed of convergence to stationarity or by reducing the asymptotic variance of the resulting MCMC estimators of, say, $\mu_\pi(f) = \int f(x)\pi(dx) = E_\pi f(X)$. The latter is often referred to as *improved mixing* (clever exploration of the state space) or *efficiency* (since asymptotically, MCMC estimators are unbiased, the **mean squared error** coincides with the variance). By asymptotic variance, we mean the limit as n (the simulation time) tends to infinity, of n times the variance of the MCMC estimator, $\hat{\mu}_n(f) = \frac{1}{n} \sum_{i=1}^n f(X_i)$, computed on a π -stationary chain updated using the transition matrix Q :

$$\begin{aligned} v(f, Q) &= \lim_{n \rightarrow \infty} n \text{Var}_\pi(\hat{\mu}_n) \\ &= \text{Var}_\pi f(X) \left[1 + 2 \sum_{j=1}^{\infty} \rho_j \right] \end{aligned} \quad (1)$$

where $\rho_j = \text{Cor}_\pi[f(X_0), f(X_j)]$ is the lag- j **autocorrelation** along the simulated π -stationary chain. The quantity $\tau = 1 + 2 \sum_{j=1}^{\infty} \rho_j$ is known as the *integrated autocorrelation time* and can be interpreted as the number of correlated samples with the same variance-reducing power as one independent sample. The integrated autocorrelation time can be estimated from the sample path, using Sokal's adaptive truncated periodogram estimator, $\hat{\tau}$ [1].

Speed of convergence is a decreasing function of the eigenvalues in absolute value of the Markov chain transition matrix (in particular what matters the most is the rate of convergence of the Markov chain, which coincides with the second largest eigenvalue in absolute value, the smaller the better), while efficiency is a function, again decreasing, of the plain eigenvalues. The two goals of speeding up convergence and reducing variance, can thus become conflicting and a good practice would be to use a transition matrix having good properties in terms of convergence for the first part of the simulation, i.e. during the burn-in phase and then switching to a Markov chain with good mixing properties. Still, the two goals are strictly connected in that pushing the spectrum of the transition matrix (which, we recall, is always between -1 and $+1$) toward the negative side, which can be achieved by inducing negative **correlation** along the Markov chain sample path, is a hard task. Thus, asymptotic variance reduction is, most of the time, in real applications, achieved by trying to reduce the autocorrelation along the sample path to zero.

In the sequel, we will review some of the current most popular practices to improve the performance of MCMC samplers giving the relevant references to the literature or pointing the reader to web sites with annotated bibliography of the relevant topics covered. One of the reference web sites for MCMC purposes is <http://www.statslab.cam.ac.uk/~MCMC> where preprints of many papers related to MCMC, referring to both theoretical and applied aspects, are posted.

Proper Scaling of the Proposal

For an **independence Metropolis–Hastings** algorithm, the closer the proposal is to the target, the faster is convergence to stationarity. In the limit, if the proposal coincides with the target, then MCMC

becomes **Monte Carlo simulation** and convergence happens in a single time step. In this case, the asymptotic variance of the (MC)MC estimator reduces to: $v(f, Q) = \text{Var}_\pi f(X)$. Unfortunately, in complicated realistic settings, it is quite unlikely that independent and identically distributed (i.i.d.) samples from the target can be obtained (*see Monte Carlo Methods, Univariate and Multivariate* for cases where direct Monte Carlo is available for multivariate distributions).

For a **random-walk Metropolis–Hastings** algorithm, guidelines to properly scale the target are given in [2, 3] for continuous state spaces. If the target is univariate normal, a random-walk Metropolis algorithm with proposal $x' = x_t + u$, where $u \sim N(0, \sigma^2)$, is optimally scaled when $\sigma = 2.38$, which corresponds to an acceptance rate of 44%. In general, if the scale parameter of the proposal is too small, nearly all proposals will be accepted, but the Markov chain will move around the state space very slowly. On the other hand, a scale parameter that is too large will result in an extremely poor acceptance rate and therefore in a highly inefficient algorithm. The authors in [2] suggest that the acceptance rate is a useful reference to be monitored when tuning the proposal in more general settings also. When the dimension of the target, d , increases, the optimal tuning parameter in a random-walk Metropolis–Hastings algorithm, becomes $\sigma = 2.4^2/d$, and the authors in [3] recommend calibrating the acceptance rate to about 25% to minimize the sample path autocorrelation. This golden rule seems to hold in a large variety of settings.

Metropolized Gibbs and Peskun Ordering

One of the advantages of the Gibbs sampler over Metropolis–Hastings type algorithms is that no tuning is needed since the proposal distribution coincides with the full-conditional. The difficulty with the Gibbs sampler arises when, in the target distribution, some of the coordinates are highly correlated. In these cases, a Gibbs sampler where each coordinate is updated separately (which is what is typically done in practice), results in very slow mixing in that the Markov chain takes a long time to traverse the state space. Unfortunately, not knowing much about the target, it is hard to guess which are the correlated coordinates that should be updated jointly

to improve efficiency; thus, a good practice is to monitor the estimated correlation as the simulation proceeds and change the sampling updating scheme if a problem is recognized. Having to sample from a highly correlated target creates mixing problems for the general Metropolis–Hastings algorithm also and the suggestion above would work here as well. As an alternative, one could try to reparameterize the target by transforming the original variables into less correlated ones. Having run the sampler on the transformed space, the original variables can be recovered simply by transforming back (assuming a unique inverse transformation exists). Examples of reparameterizations for a variety of linear and nonlinear models can be found in [4].

One interesting variation of the Gibbs sampler is the “Metropolized Gibbs sampler” [5]. The idea is to exclude, when updating each coordinate X_i , an immediate draw of X_i itself, thus ensuring that the Markov chain does not preserve the same position over successive points in time. The proposal distribution is not the full-conditional, $\pi(X_i|X_{-i})$ anymore, where $X_{-i} = (X_j, j \neq i)$, but a value Y_i , different from X_i is drawn with probability $\pi(Y_i|X_{-i})/[1 - \pi(X_i|X_{-i})]$. By doing so, one needs to insert a reject–accept step to correct for the possible bias induced by not using the full-conditional distribution as the proposal. The proper acceptance probability is given in [5]. The Metropolized Gibbs sampler is better than the regular Gibbs sampler on discrete state spaces (or when at least one of the components in the target is discrete) in terms of efficiency, i.e. better in the Peskun sense [6]. The intuition behind Peskun (partial) ordering is that whenever the same position is retained over time, the sampler becomes less efficient in that the state space is not explored. In general state spaces, the same position is retained over time whenever a proposal is rejected. Thus, in principle, we could improve the Metropolis–Hastings and the **Reversible jump** [7] algorithms in the Peskun efficiency ordering by decreasing the rejection frequency of the proposed moves. An application of this intuition appears in [8] and [9] where a delaying rejection mechanism is introduced to improve the Metropolis–Hastings and the Reversible jumps algorithms (respectively) on a sweep-by-sweep basis. The delaying rejection strategy is illustrated in more detail in the section titled “Adaptive MCMC and Particle Filters”.

Auxiliary Variables

For better mixing or faster convergence to stationarity, auxiliary or latent variables can be introduced in an MCMC simulation.

The idea of introducing auxiliary variables in MCMC sampling arose in statistical physics [10, 11], and was brought into the mainstream statistical literature by Besag and Green in [12].

Auxiliary variable techniques exploit the general principle that a complicated posterior distribution may become more tractable if it is embedded in a higher-dimensional state space. This is the same principle that drives data augmentation and expectation maximization (EM) algorithms (see **Expectation Maximization Algorithm** and [13]). Just to give an example of this principle at work, consider a Bayesian missing data problem. In this setting the “complete data” model (obtained by augmenting the observed data model with the missing data), typically has a nice analytical form and all the posterior computations can be carried out in closed form.

Once the higher-dimensional sample is obtained either by MCMC, importance sampling, or other simulation technique, a projection back on the original state space is obtained simply by discarding the auxiliary variables.

The simple *slice sampler* is probably the most popular auxiliary variable MCMC method, where a single auxiliary variable, γ , is introduced, and the joint distribution of the original variable of interest of interest, X and γ is specified by taking the marginal distribution of X to be the target and specifying the distribution of $\gamma|X$ to be uniform on the interval $[0, \pi(X)]$:

$$\begin{aligned}\pi(\gamma, X) &= \pi(X)\pi(\gamma|X) \\ &= \pi(X) \frac{1}{\pi(X)} \mathbf{1}_{\{0 \leq \gamma \leq \pi(X)\}}(\gamma)\end{aligned}\quad (2)$$

where $\mathbf{1}_A(\gamma)$ is the indicator function of the set A . The joint distribution over the enlarged state space is then explored *via* Gibbs sampler. This amounts to updating $\gamma|X$ by generating a uniform random variable over the interval specified above and then sampling $X|\gamma$ from a uniform distribution over the set $\{X : \pi(X) > \gamma\}$. While the first updating step is straightforward, the second one might be difficult since it requires the identification of the level sets of the target distribution. For a clever way to perform

this second update see [14], where an adaptive rejection mechanism is proposed for this aim.

Easy-to-check conditions for uniform and geometric ergodicity of the slice sampler are given in [15] and [16], respectively.

The simple slice sampler can be extended by introducing more than one auxiliary variable and this is of particular interest in a Bayesian context. The clever introduction of multiple auxiliary variables in a variety of Bayesian nonconjugate and hierarchical models is illustrated in [17].

In [15] it is proved that an independence Metropolis–Hastings algorithm can be turned into a corresponding slice sampler that produces MCMC estimators with a smaller asymptotic variance. In other words, the independence Metropolis–Hastings is dominated in the Peskun ordering [6] by the corresponding slice sampler. One can also easily devise over relaxed versions of slice sampling [14], which sometimes greatly improve the speed of convergence to stationarity by suppressing random-walk behavior.

Furthermore, in [18], the slice sampler is proved to be stochastically monotone with respect to the ordering given on the state space by the target: $X \prec Y$ if $\pi(X) \leq \pi(Y)$. This is the condition that is needed to turn the slice sampler into a perfect simulation algorithm (see **Perfect Sampling**). Another version of perfect slice sampler for mixture models is given in [19].

A different, interesting way of enlarging the state space by introducing auxiliary variables is given by *hybrid Monte Carlo* algorithms [20, 21], where a momentum vector (with the same dimension as the original variable of interest, X) is sampled from a Gaussian distribution and then a deterministic time-reversible and volume-preserving algorithm is run on the enlarged state space for some steps to obtain the proposed move, which is then either retained or rejected, based on some acceptance probability that preserves reversibility over the enlarged state space.

A final example of an auxiliary variable introduced for better mixing is the temperature parameter in the *simulated tempering* strategy [22]. The idea is to run, in series, m MCMC samplers with different stationary distributions and occasionally interchange the position of two of the chains. The m stationary distributions are obtained from the original target *via* incremental heating. Higher temperatures will lead to flatter stationary distributions that are thus easier to explore. A sampler converging to a heated target will

4 Convergence and Mixing in MCMC

be fast mixing and, *via* swap moves among adjacent chains, colder samplers will take advantage of this and inherit the good mixing property.

Langevin Diffusion MCMC Algorithms

If the gradient of the posterior distribution is available, it can be used to set up a Langevin algorithm. The simplest form of this type of algorithms uses as proposal density:

$$q(X, X') = \frac{1}{(2\pi\sigma^2)^{d/2}} \exp \left\{ -\frac{\|X' - X - \sigma^2 \nabla \log \pi(x)/2\|^2}{2\sigma^2} \right\} \quad (3)$$

for some suitable scaling parameter σ . Roberts and Rosenthal in [23] show that the choice of the scaling parameter should lead to an acceptance rate of 0.574 to achieve optimal convergence rates in a special case where the components of the target are uncorrelated. The only difference between a Langevin diffusion MCMC and a regular Metropolis–Hastings algorithm is that the additional information contained in the gradient (which, in difficult settings, could be replaced by numerical derivatives) is used to push the mean of the proposal toward high posterior regions. The caveat is that care should be taken in avoiding overshooting that causes the Markov chain to jump from one tail region of the posterior to another. Theoretical study of convergence properties of Langevin algorithm is given in [24–26].

Adaptive MCMC and Particle Filters

Recent research lines in MCMC aim at finding almost automatic algorithms that need little off-line pilot tuning of the proposal. To this scope, a few adaptive MCMC strategies have been recently proposed in the MCMC literature. There are mainly two ways of learning from the past history of the simulation to better design and calibrate the proposal. One can either preserve the Markovian structure of the sampler and this is obtained by performing only partial adaptation as in the delayed rejection strategy [9, 27], or by adapting at regeneration times as in [28]. Alternatively, one can adapt to the whole

history, thus obtaining a non-Markovian sampler, but then one has to prove the ergodicity of the resulting algorithm from first principles as in [29–31].

In the sequel, we will focus on two adaptive strategies, namely, a local and a global one. The intuition behind the *delayed rejection strategy* is that, upon rejection of a candidate move, instead of advancing simulation time and retaining the same position, a second-stage proposal distribution can be used to generate a new candidate move. The acceptance probability of this second-stage move is computed to preserve reversibility with respect to the target [32]. The rationale of this second-stage proposal is twofold: by trying harder to move away from the current position, better exploration of the state space is achieved. This improves the mixing properties of the sampler as can be formally proved by showing that the resulting algorithm dominates, in terms of Peskun ordering [6], the basic Metropolis–Hastings sampler. The other reason that makes the delayed rejection strategy interesting is that the second-stage proposal can be different from the first-stage one and can depend, not only on the current position of the Markov chain, but also on the first-stage proposal. This allows partial adaptation and local calibration of the proposal.

In the *adaptive Metropolis* algorithm, a Gaussian proposal distribution is centered at the current position and its **covariance** matrix is calibrated using the whole sample path of the Markov chain. After an initial nonadaptation period of length t_0 , the proposal covariance matrix is obtained as

$$C_t = s_d \text{Cov}(X_0, \dots, X_{t-1}) + s_d \varepsilon I_d, \quad t > t_0 \quad (4)$$

where s_d is a parameter that depends only on the dimension d of the state space where π is defined, $\varepsilon > 0$ is a constant that we may choose to be very small and I_d denotes the d -dimensional identity matrix. In [30] ergodicity of the resulting algorithm is proved if the target is bounded and supported on a bounded set. The adaptive Metropolis and other more flexible adaptive strategies are proved to be ergodic under less stringent hypothesis in [33]. A very powerful adaptive strategy that combines global and local adaptation is the DRAM algorithm (Delayed Rejection in Adaptive Metropolis). See [31], where the idea is tested on high-dimensional real applications.

Many important applications involve real-time data that are collected in a sequential manner. Examples include signal processing problems, target tracking, on-line medical monitoring, computer vision, financial applications, speech recognition, and communications, among others. When new data become available, the likelihood and therefore also the posterior distribution are updated and thus the target distribution of the MCMC simulation evolves as we analyze it.

In this setting, there is little time to run, until convergence, a new simulation between two successive data inputs. Still we can rely on the fact that the target distribution should not change much between two successive time intervals and thus, instead of running a totally new MCMC sampler, one could reweight the current, realized Markov chain path. Furthermore, instead of running a single chain, we can work with a set of “particles” that approximate the target at any given point in time. Particle filter algorithms are constructed by combining steps that involve updating the weights, resampling with probabilities proportional to the weights to eliminate unrepresentative particles, and changing the particle position by MCMC moves, among others.

Interested researchers can refer to the web site of sequential Monte Carlo methods and particle filtering: <http://www.sigproc.eng.cam.ac.uk/smc/> and to the book by Doucet *et al.* [34], for further reading on the topic.

To conclude, it is emphasized that particle filters can be helpful also in static contexts where standard MCMC algorithms show poor mixing. For example, given static target distribution can be turned into a dynamic one, as suggested in [35], by splitting the available data into batches that are used to recursively update the target that thus artificially evolves in time.

Importance Sampling and Population Monte Carlo

A direct competitor of MCMC is importance sampling. An importance distribution is selected, say g , and samples are generated from it. Since the quantity $\mu_\pi(f)$ can be rewritten as

$$\mu_\pi(f) = \int f(x) \frac{\pi(x)}{g(x)} g(x) dx \quad (5)$$

an asymptotically unbiased estimator of μ is

$$\hat{\mu}_n^{(2)}(f) = \frac{1}{n} \sum_{i=1}^n f(X_i) \frac{\pi(X_i)}{g(X_i)}, \quad (6)$$

where X_i are i.i.d. from g . The quantities $w_i = \frac{\pi(X_i)}{g(X_i)}$ are the importance weights. The importance distribution, g , plays the same role as the proposal distribution, q , in the Metropolis–Hastings sampler and the weights, w , play a role similar to the one of the acceptance probabilities. The importance sampling estimator has finite variance only if

$$E_\pi \left[f^2(X) \frac{\pi(X)}{g(X)} \right] < \infty \quad (7)$$

For further details on the implementation and the performance of importance sampling, *see Monte Carlo Methods, Univariate and Multivariate*.

An interesting combination of importance sampling techniques, adaptive Metropolis–Hastings algorithms, and particle filter ideas gives rise to the so-called population Monte Carlo methodology, the name being inspired from a survey by Iba [36]. The idea is to produce a population of points (or particles) of size N , say, associated with importance weights, and to update both the position and the weight of those points repeatedly over iterations. More precisely, at any iteration t of the population Monte Carlo algorithm, N particles are generated from an importance distribution that is determined based on the behavior of the $N \times (t-1)$ previous particles with respect to the target distribution. (The iterations are completely artificial in that their number can be chosen beforehand or during the simulation experiment and that stopping the experiment at any iteration does not modify the validity of the algorithm.) The new particles thus generated are then reweighted and possibly resampled by multinomial sampling or more clever schemes ([37, 38]). The major appeal of population Monte Carlo is that the importance distribution can be constructed in a completely open way based on all particles at any time in the past while preserving the fundamental importance sampling property. This idea sounds very similar to what is done in the adaptive MCMC strategies but with much more **degrees of freedom** because there is no need to worry about ergodicity and burn-in issues. As a result, importance sampling is much more amenable to adaptivity than MCMC because its asymptotics are only in N (as regular importance sampling schemes), not in t (as

MCMC methods). While the first paper on population Monte Carlo is [39], more details and updates are provided in both [40] and [41], the latter establishing improvement over the iterations.

Conclusions

Despite the fact that properly scaling the proposal distribution in a Metropolis–Hastings type algorithm is a hard task, guidelines exist in the literature, and adaptive MCMC strategies can be quite helpful as they permit almost automatic on-line tuning. Another difficult task in MCMC simulation is the definition of the length of the burn-in phase. Fortunately, free software for convergence diagnostics can, at least, tell us when convergence is not achieved. Also, when available, perfect simulation avoids the burn-in issues. For a detailed description of perfect simulation see **Perfect Sampling**, here we recall that both the independence sampler [42] and the slice sampler [18, 19], can be turned into perfect samplers under some regularity conditions, which are the same conditions that guarantee their uniform ergodicity. Another way to avoid the burn-in issue and, at the same time, to take advantage of adaptation in designing the instrumental distribution is *via* population Monte Carlo.

One of the main advantages of MCMC over other deterministic approximation methods is that it does not suffer the curse of dimensionality.

It is commonly agreed that “finding the ideal proposal chain is an art” [43] and this chapter gives some guidelines on where to look for it and what it looks like. Finally, techniques to reduce the amount of computation per iteration are also important in reducing run times and thus improving efficiency.

Also, even if a nonoptimal sampler has been run, the obtained samples could still be used in a more efficient way than simply taking the ergodic average. In particular, it is recalled that postprocessing the chain, by subsampling or Rao–Blackwellization can be quite useful.

When subsampling, typically every $\hat{\tau}$ observation is retained, thus reducing the covariance term in the expression of the asymptotic variance and building the MCMC estimators on almost-independent samples.

Rao–Blackwellization relies on the well-known principle that conditioning reduces the variance. The

idea is to replace $f(X_i)$ in $\hat{\mu}_n(f)$ by a conditional expectation, $E_\pi[f(X_i)|h(X_i)]$, for some function h , or to condition on the previous value of the chain thus using $E[f(X_i)|X_{i-1} = x]$ instead [44–47].

To conclude, we point out that there remains much to be done toward automation of MCMC simulations. Still there is a caveat to the blind use of MCMC without first trying to analytically integrate out some of the parameters, whenever possible, or checking that the distribution of interest is proper.

References

- [1] Sokal, A.D. (1998). Monte Carlo methods in statistical mechanics: foundations and new algorithms, *Cours de Troisième Cycle de la Physique en Suisse Romande*, Lausanne.
- [2] Gelman, A., Gilks, W.R. & Roberts, G.O. (1996). Efficient Metropolis jumping rules, in *Bayesian Statistics*, J.O. Berger, J.M. Bernardo, A.P. Dawid, D.V. Lindley & A.F.M. Smith, eds, Oxford University Press, Oxford, Number 5, pp. 599–608.
- [3] Roberts, G.O., Gelman, A. & Gilks, W.R. (1997). Weak convergence and optimal scaling of random walk Metropolis algorithms, *Annals of Applied Probability* **7**, 110–120.
- [4] Gilks, W.R. & Roberts, G.O. (1996). *Strategies for Improving MCMC*, Chapman, London, pp. 89–114.
- [5] Liu, J.S. (1996). Peskun’s Theorem and a Modified Discrete-State Gibbs Sampler, *Biometrika* **83**(3), 681–682.
- [6] Peskun, P.H. (1973). Optimum Monte Carlo sampling using Markov chains, *Biometrika* **60**, 607–612.
- [7] Green, P.J. (1995). Reversible jump Markov chain Monte Carlo computation and Bayesian model determination, *Biometrika* **82**, 711–732.
- [8] Tierney, L. & Mira, A. (1999). Some adaptive Monte Carlo methods for Bayesian inference, *Statistics in Medicine* **18**, 2507–2515.
- [9] Green, P.J. & Mira, A. (2001). Delayed rejection in Reversible jump Metropolis-Hastings, *Biometrika* **88**, 1035–1053.
- [10] Swendsen, R.H. & Wang, J.S. (1987). Nonuniversal critical dynamics in Monte Carlo simulations, *Physical Review Letters* **58**, 86–88.
- [11] Edwards, R.G. & Sokal, A.D. (1988). Generalization of the Fortuin-Kasteleyn-Swendsen-Wang representation and Monte Carlo algorithm, *Physical Review D* **38**, 2009–2012.
- [12] Besag, J. & Green, P.J. (1993). Spatial statistics and Bayesian computation (with discussion), *Journal of the Royal Statistical Society. Series B* **55**, 25–38.
- [13] Meng, X.L. & van Dyk, D. (1997). The EM algorithm—an old folk-song sung to a new tune (with discussion), *Journal of the Royal Statistical Society. Series B* **59**, 511–568.

-
- [14] Neal, R.M. (2003). Slice sampling (with discussion), *Annals of Statistics* **31**, 705–767.
- [15] Mira, A. & Tierney, L. (2002). Efficiency and convergence properties of slice samplers, *Scandinavian Journal of Statistics* **29**(1), 1–12.
- [16] Roberts, G.O. & Rosenthal, J.S. (1999). Convergence of slice sampler Markov chains, *Journal of the Royal Statistical Society. Series B* **61**, 643–660.
- [17] Damien, P., Wakefield, J. & Walker, S. (1999). Gibbs sampling for Bayesian non-conjugate and hierarchical models by using auxiliary variables, *Journal of the Royal Statistical Society. Series B* **61**(2), 331–344.
- [18] Mira, A., Møller, J. & Roberts, G.O. (2001). Perfect slice samplers, *Journal of the Royal Statistical Society. Series B* **63**, 583–606.
- [19] Casella, G., Mengersen, K.L., Robert, C.P. & Titterton, D.M. (2002). Perfect slice samplers for mixtures of distributions, *Journal of the Royal Statistical Society. Series B* **64**(4), 777–790.
- [20] Duane, S., Kennedy, A.D., Pendleton, B.J. & Roweth, D. (1987). Hybrid Monte Carlo, *Physics Letters* **B195**, 216–222.
- [21] Neal, R.M. (1994). An improved acceptance procedure for the hybrid Monte Carlo algorithm, *Journal of Computational Physics* **111**, 194–203.
- [22] Geyer, C.J. & Thompson, E.A. (1992). Constrained Monte Carlo maximum likelihood for dependent data (with discussion), *Journal of the Royal Statistical Society. Series B* **54**, 657–699.
- [23] Roberts, G.O. & Rosenthal, J.S. (1998). Markov chain Monte Carlo: some practical implications of theoretical results (with discussion), *Canadian Journal of Statistics* **26**, 5–32.
- [24] Roberts, G.O. & Tweedie, R.L. (1996). Exponential convergence of Langevin diffusions and their discrete approximations, *Bernoulli* **2**, 341–364.
- [25] Stramer, O. & Tweedie, R.L. (1999). Langevin-type models i: diffusions with given stationary distributions, and their discretizations, *Methodology and Computing in Applied Probability* **1**, 283–306.
- [26] Stramer, O. & Tweedie, R.L. (1999). Langevin-type models ii: Self-targeting candidates for Hastings-Metropolis algorithms, *Methodology and Computing in Applied Probability* **1**, 307–328.
- [27] Tierney, L. & Mira, A. (1998). Some adaptive Monte Carlo methods for Bayesian inference, *Statistics in Medicine* **18**, 2507–2515.
- [28] Gilks, W.R., Roberts, G.O. & Sahu, S.K. (1998). Adaptive Markov chain Monte Carlo, *Journal of the American Statistical Association* **93**, 1045–1054.
- [29] Haario, H., Saksman, E. & Tamminen, J. (1999). Adaptive proposal distribution for random walk Metropolis algorithm, *Computational Statistics* **14**(3), 375–395.
- [30] Haario, H., Saksman, E. & Tamminen, J. (2001). An adaptive Metropolis algorithm, *Bernoulli* **7**(2), 223–242.
- [31] Haario, H., Laine, M., Mira, A. & Saksman, E. (2006). DRAM: efficient adaptive MCMC, *Statistics and Computing* **16**, 339–354.
- [32] Mira, A. (2001). On Metropolis-Hastings algorithms with delayed rejection, *Metron* **LIX**(3–4), 231–241.
- [33] Andrieu, C. & Moulines, E. (2006). On the ergodicity properties of some adaptive MCMC algorithms, *Annals of Applied Probability* **16**(3), 1–8.
- [34] Doucet, A., de Freitas, N. & Gordon, N. (2001). *Sequential Monte Carlo Methods in Practice*, Springer-Verlag.
- [35] Chopin, N. (2002). A sequential particle filter method for static models, *Biometrika* **89**, 539–552.
- [36] Iba, Y. (2000). Population-based Monte Carlo algorithms, *Transactions of the Japanese Society for Artificial Intelligence* **16**(2), 279–286.
- [37] Kitagawa, G. (1996). Monte Carlo filter and smoother for non-Gaussian non-linear state space models, *Journal of Computational and Graphical Statistics*, **5**, 1–25.
- [38] Carpenter, J., Clifford, P. & Fernhead, P. (1999). Building robust simulation-based filters for evolving datasets, *IEE Proceedings – Radar, Sonar Navigation* **146**, 2–4.
- [39] Cappé, O., Guillin, A., Marin, J.M. & Robert, C.P. (2004). Population Monte Carlo, *Journal of Computational and Graphical Statistics* **13**, 907–929.
- [40] Robert, C.P. & Casella, G. (2004). *Monte Carlo Statistical Methods*, 2nd Edition, Springer, New York.
- [41] Douc, R., Guillin, A., Marin, J.-M. & Robert, C.P. (2007). Convergence of adaptive mixtures of importance sampling schemes, *Annals of Statistics* (To appear).
- [42] Corcoran, J. & Tweedie, R.T. (2002). Perfect sampling from independent metropolis-hastings chains, *Journal of Statistical Planning and Inference* **104**(2), 297–314.
- [43] Liu, J.S. (2001). *Monte Carlo Strategies in Scientific Computing*, Springer-Verlag, New York.
- [44] Gelfand, A.E. & Smith, A.F.M. (1990). Sampling based approaches to calculating marginal densities, *Journal of the American Statistical Association* **85**, 398–409.
- [45] Liu, J.S., Wong, W.H. & Kong, A. (1994). Covariance structure of the Gibbs sampler with applications to the comparisons of estimators and sampling schemes, *Biometrika* **81**, 27–40.
- [46] Casella, G. & Robert, C.P. (1996). Rao-Blackwellisation of sampling schemes, *Biometrika* **83**(1), 81–94.
- [47] McKeague, I.W. & Wefelmeyer, W. (2000). Markov chain Monte Carlo and Rao-Blackwellisation, *Journal of Statistical Planning and Inference* **85**, 171–182.

ANTONIETTA MIRA

Optimal Reliability Design-Modeling

Introduction

Reliability has become an even greater concern in recent years because high-tech industrial processes, with increasing levels of sophistication, comprise most engineering systems today. Based on enhancing **component reliability** and providing **redundancy**, while considering the trade-off between system performance and resources, optimal reliability design that aims to determine an optimal system-level configuration has long been an important topic in reliability engineering (*see Reliability Optimization*). Since 1960, many publications have addressed this problem using different system structures, performance measures, optimization techniques, and options for reliability improvement.

References 1, 2 and 3 provide good overview of the early work in system reliability optimization. Tillman *et al.* [3] were the first to classify articles by system structure, problem type, and solution methods. Also described and analyzed in [4] are the advantages and shortcomings of various optimization techniques. It was during the 1970s that various heuristics were developed to solve complex system reliability problems in cases where the traditional parametric optimization techniques were insufficient. In their 2000 report, Kuo and Prasad [1] summarize the developments in optimization techniques, along with recent optimization methods such as metaheuristics. This chapter discusses the contributions made to the literature since the publication of [1]. Most of the recent work in this area is devoted to:

- multistate system (MSS) optimization;
- percentile life as a system performance measure;
- multiobjective optimization;
- active and cold-standby redundancy;
- optimization techniques.

Based on their system performance, reliability systems can be classified as binary-state systems and MSS. A binary-state system and its components may exist in only two possible states—either working or failed. Binary-state system reliability models have played very important roles in practical

reliability engineering. To satisfactorily describe the performance of a complex system we may need more than two levels of satisfaction, for example, excellent, average, and poor [5]. For this reason, MSS reliability models were proposed in the 1970s and since 1998 a large part of the work devoted to MSS optimal design has emerged (*see Multistate Reliability Theory*). The primary task of MSS optimization is to define the relationship between component states and system states.

Measures of system performance are basically of four kinds: reliability, availability, mean time-to-failure (TTF), and percentile life (*see System Availability*). Reliability has been widely used and thoroughly studied as the primary performance measure for nonmaintained systems. For a maintained system, however, availability, which describes the percentage of time the system really functions, should be considered instead of reliability. Availability is most commonly employed as the performance measure for renewable MSS. Meanwhile, percentile life is preferred to reliability and mean TTF when the system mission time is indeterminate, as in most practical cases.

Some important design principles for improving system performance are summarized in [1]. This chapter primarily reviews articles that address either the provision of redundant components in parallel or the combination of structural redundancy with the enhancement of component reliability (*see Reliability of Redundant Systems; Reliability Allocation*). These are called *redundancy allocation problems* and *reliability-redundancy allocation problems*, respectively. Redundancy allocation problems are well documented in [1] and [5], which employ a special case of reliability-redundancy allocation problems without exploring the alternatives of component-combined improvement. Recently, a lot of the effort in optimal reliability design has been placed on general reliability-redundancy allocation problems, rather than on redundancy allocation problems.

In practice, two redundancy schemes are available: active and cold-standby. Cold-standby redundancy provides higher reliability, but it is hard to implement because of the difficulty of failure detection. Reliability design problems have generally been formulated considering active redundancy; however, an actual optimal design may include active redundancy or cold-standby redundancy or both.

2 Optimal Reliability Design-Modeling

However, any effort for improvement usually requires resources. Quite often it is hard for a single-objective system to adequately describe a real optimal design problem. For this reason, multiobjective system design problem always deserves a lot of attention.

This chapter describes the state-of-the-art of optimal reliability design. Emphasizing the foci mentioned above, it includes four main problem formulations in optimal **reliability allocation**; describes advances related to those four types of optimization problems; and also provides conclusions and a discussion of future challenges related to reliability optimization problems.

Problem Formulations

Among the diversified problems in optimal reliability design, the following four basic formulations are widely covered.

Problem 1 (P_1):

$$\begin{aligned} & \max R_S = f(x) \\ & \text{s.t.} \\ & g_i(x) \leq b_i, \quad \text{for } i = 1, \dots, m \\ & x \in X \\ & \text{or} \\ & \min C_S = f(x) \\ & \text{s.t.} \\ & R_S \geq R_0 \\ & g_i(x) \leq b_i, \quad \text{for } i = 1, \dots, m \\ & x \in X \end{aligned}$$

Problem 1 formulates the traditional reliability-redundancy allocation problem with either reliability or cost as the objective function. Its solution includes two parts: the component choices and their corresponding optimal redundancy levels.

Problem 2 (P_2):

$$\begin{aligned} & \max t_{\alpha, x} \\ & \text{s.t.} \\ & g_i(t_{\alpha, x}; x) \leq b_i, \quad \text{for } i = 1, \dots, m \\ & x \in X \end{aligned}$$

Problem 2 uses percentile life as the system performance measure instead of reliability. Percentile life is preferred especially when the system mission time is indeterminate. However, it is hard to find a closed analytical form of percentile life in decision variables.

Problem 3 (P_3):

$$\begin{aligned} & \max E(x, T, W^*) \\ & \text{s.t.} \\ & g_i(x) \leq b_i, \quad \text{for } i = 1, \dots, m \\ & x \in X \\ & \text{or} \\ & \min C_s(x) \\ & \text{s.t.} \\ & E(x, T, W^*) \geq E_0 \\ & g_i(x) \leq b_i, \quad \text{for } i = 1, \dots, m \\ & x \in X \end{aligned}$$

Problem 3 represents MSS optimization problems. Here, E is used as a measure of the entire **system availability** to satisfy the custom demand represented by a cumulative demand curve with a known T and W^* .

Problem 4 (P_4):

$$\begin{aligned} & \max z = [f_1(x), f_2(x), \dots, f_S(x)] \\ & \text{s.t.} \\ & g_i(x) \leq b_i, \quad \text{for } i = 1, \dots, m \\ & x \in X \end{aligned}$$

For multiobjective optimization, as formulated by Problem 4, a Pareto optimal set, which includes

Table 1 Reference classification by problem formulations

P_1	[6], [7], [8], [9], [10], [11], [12], [13], [14], [15], [16], [17], [18], [19], [20], [21], [22], [23], [24], [25], [26], [27], [28], [29], [30], [31], [32], [33], [4], [34], [35], [36], [37], [38]
P_2	[39], [40], [41], [42], [37], [38]
P_3	[43], [44], [45], [46], [47], [48], [49], [50], [51], [52], [53], [54], [55], [56], [57], [58], [59], [60], [22], [61], [62], [63], [64]
P_4	[65], [66], [67], [68], [69], [70], [71], [72], [73], [37], [38]

all of the best possible trade-offs between given objectives, rather than a single optimal solution, is usually identified.

In the above formulations, the resource constraints may be linear or nonlinear or both.

The articles, classified by problem formulations, is summarized in Table 1.

Brief Reviews of Advances in P_1 – P_4

Active and Cold-Standby Redundancy (P_1)

P_1 is generally limited to active redundancy. A new optimal system configuration is obtained when active and cold-standby redundancies are both involved in the design. A cold-standby redundant component does not fail before it is put into operation by the action of switching, whereas the failure pattern of an active redundant component does not depend on whether the component is idle or in operation. Cold-standby redundancy can provide higher reliability, but it is hard to implement because of the difficulties involved in failure detection and switching.

In reference 10, optimal solutions to reliability-redundancy allocation problems are determined for nonrepairable systems designed with multiple **k -out-of- n** subsystems in series. The individual subsystems may use either active or cold-standby redundancy, or they may require no redundancy. Assuming an exponentially distributed component TTF with rate λ_{ij} , the failure process of subsystem i with cold-standby redundancy can be described by a **Poisson process** with rate $\lambda_{ij}k_i$, while the subsystem reliability with active redundancy is computed by standard binomial techniques. For series-**parallel systems** with only cold-standby redundancy, reference 11 employs the more flexible and realistic Erlang distributed component TTF and $\rho_i(t)$ is introduced to describe the reliability of the imperfect detection/switching mechanism for each subsystem. Reference 12 directly extends this earlier work by introducing the choice of redundancy strategies as an additional decision variable. For each of these methods, however, no mixture of component types or redundancy strategies is allowed within any of the subsystems.

In addition, reference 74 investigates the problem of where to allocate a spare in a k -out-of- n : F system of dependent components through minimal standby redundancy; and reference 75 studies the allocation

of one active redundancy when it differs based on the component with which it is to be allocated. Reference 28 considers the problem of optimally allocating a fixed number of s -identical multifunctional spares for a deterministic or stochastic mission time. In spite of some sufficiency conditions for optimality, the proposed algorithm can be easily implemented even for large systems.

Percentile Life Optimization (P_2)

Many diversified models and solution methods, where reliability is used as the system performance measure, have been proposed and developed since the 1960s. However, this is not an appropriate choice when mission time cannot be clearly specified or a system is intended for use as long as it functions. Average life is also not reliable, especially when the implications of failure are critical or the variance in the system life is high. Percentile life is considered to be a more appropriate measure, since it incorporates system designer and user risk. When using percentile life as the objective function, the main difficulty is its mathematical inconvenience, because it is hard to find a closed analytical form of percentile life in the decision variables.

Reference 39 solves redundancy allocation problems for series-parallel systems where the objective is to maximize a lower percentile of the system TTF distribution. Component TTF has a Weibull distribution with known deterministic parameters. Later in the literature, reference 40 addresses similar problems where Weibull shape parameters are accurately estimated but scale parameters are random variables following a uniform distribution.

Reference 42 develops a lexicographic search methodology that is, the first to provide exact optimal redundancy allocations for percentile life optimization problems. This algorithm is general for any continuous increasing lifetime distribution.

Three important results are presented in [41], which describe the general relationships between reliability and percentile life maximizing problems.

- S_2 equals S_1 given $\alpha_t = 1 - R_s(t, x_t^*)$, where $x_t^* \in S_1$;
- S_1 equals S_2 given t_{α, x_{α}^*} , where $x_{\alpha}^* \in S_2$;
- let $\psi(t)$ be the optimal objective value of P_1 . For a fixed α , $t_{\alpha} = \inf\{t \geq 0 : \Psi(t) \leq 1 - \alpha\}$ is the optimal objective value of P_2 .

MSS Optimization (P_3)

MSS is defined as a system that can unambiguously perform its task at different performance levels, depending on the state of its components, which can be characterized by nominal performance rate, availability and cost. Based on their physical nature, MSS can be classified into two important types: Type I MSS (e.g., power systems) and Type II MSS (e.g., computer systems), which use capacity and operation time as their performance measures, respectively.

In the article, an important optimization strategy, combining universal generating function (UGF) and genetic algorithm (GA), has been well developed and widely applied to reliability optimization problems of renewable MSS. In this strategy, there are two main tasks:

- According to the system structure and the system physical nature, obtain the system UGF from the component UGFs.
- Find an effective decoding and encoding technique to improve the efficiency of the GA.

Reference 43 first uses a UGF approach to evaluate the availability of a series-parallel MSS with relatively small computational resources. The essential property of the U -transform enables the total U -function for a MSS to be obtained by simple algebraic operations involving individual component U -functions. Later, reference 76 combines importance and sensitivity analysis and reference 77 extends this UGF approach to MSS with dependent elements. Table 2 summarizes the application of UGF to some typical MSS structures in optimal reliability design.

Table 2 Application of UGF approach

Series-parallel system	[43], [45], [76], [49], [51], [54], [57] [58], [59], [62]
Bridge system	[44], [46], [78], [61]
LMSSWS	[53], [79], [60]
WVS	[48], [80]
ATN	[52], [81]
LMCCS	[55], [56]
ACCN	[82]

LMSSWS: linear multi-state sliding-window system; WVS: weighted voting system; ATN: acyclic transmission network; LMCCS: linear multi-state consecutively connected system; ACCN: acyclic consecutively connected network

With this UGF and GA strategy, reference 61 solves the structure optimization of a MSS with time redundancy; reference 47 considers a MSS consisting of two parts: RGS including a number of resource generating subsystems and MPS including elements that consume a fixed amount of resources to perform their tasks; references 49, 54, 78 considers multistate series-parallel systems with two failure modes – open mode and closed mode, and references 46, 50, 55, 57–59 use survivability, instead of reliability to describe the ability of a MSS to tolerate both internal failures and external attacks.

In addition to a GA, reference 62 presents an ant colony method that combines with a UGF technique to find an optimal series-parallel power-structure configuration.

Besides this primary UGF approach, a few other methods have been proposed for MSS reliability optimization problems. Reference 64 develops a heuristic algorithm RAMC for a Type I multistate series-parallel system, while reference 83 presents a novel continuous-state system mode approximating the objective reliability function by a **neural network** approach.

Multiobjective Optimization (P_4)

In the previous discussion, all problems were single-objective problems. Rarely does a single-objective problem with several hard constraints adequately represent a real problem for which an optimal design is required. When designing a reliable system, as formulated by P_4 , it is always desirable to simultaneously optimize several opposing design objectives such as reliability, cost, even volume, and weight. For this reason, a recently proposed multiobjective system design problem deserves a lot of attention. The objectives of this problem are to maximize the system reliability estimates and minimize their associated variance while considering the uncertainty of the component reliability estimations. A Pareto optimal set, which includes all of the best possible trade-offs between the given objectives, rather than a single optimal solution, is usually identified for multiobjective optimization problems.

When considering complex systems, the reliability optimization problem has been modeled as a fuzzy multiobjective optimization problem in reference 69, considering the influence of various kinds of aggregators.

With a weighting technique, references 70, 71 solve multiobjective reliability-redundancy allocation problems using similar linear membership functions for both objectives and constraints; reference 67 transfers P_4 into a single-objective optimization problem and proposes a GA-based approach whose parameters can be adjusted with the experimental plan technique; and reference 65 develops a multiobjective GA to obtain an optimal system configuration and **inspection policy** by considering every target as a separate objective.

P_4 is considered for series-parallel systems, $RB/1/1$, and bridge systems in [66] with multiple objectives to maximize the system reliability while minimizing its associated variance when the component reliability estimates are treated as random variables. For series-parallel systems, component reliabilities of the same type are considered to be dependent since they usually share the same reliability estimate from a pooled data set. The system variance is straightforwardly expressed as a function in the higher **moments** of the component unreliability estimates [84]. For $RB/1/1$, the hardware components are considered identical and statistically independent, while even independently developed software versions are found to have related faults as presented by the parameters Prv and $Pall$. Pareto optimal solutions are found by solving a series of weighted objective problems with incrementally varied weights. It is worth noting that significantly different designs are obtained when the formulation incorporates estimation uncertainty or when the component reliability estimates are treated as statistically dependent. Similarly, [68] uses a multiobjective GA to select an optimal network design that balances the dual objectives of high reliability and low uncertainty in its estimation. But the latter exploits **Monte Carlo simulation** as the objective function evaluation engine (see also **Monte Carlo Methods, Univariate and Multivariate**).

Besides GA, reference 72 illustrates the application of the ant colony optimization algorithm to solve both continuous function and combinatorial optimization problems, while reference 73 tests five simulated annealing-based multiobjective algorithms—suppapitnarm multi-objective simulated annealing (SMOSA), ulungu multi-objective simulated annealing (UMOSA), pareto simulated annealing (PSA), pareto domination based multi-objective

simulated annealing (PDMOSA) and weight based multi-objective simulated annealing (WMOSA).

Conclusions and Discussions

We have reviewed the recent research on optimal reliability design. Many publications have addressed this problem using different system structures, performance measures, problem formulations and optimization techniques.

The systems considered here mainly include series-parallel systems, k -out-of- n ; G systems, bridge networks, n -version programming architecture, recovery block architecture and other unspecified coherent systems (see **Coherent Systems; Parallel, Series, and Series-Parallel Systems; k -out-of- n Systems**). The recently introduced N -version programming (NVP) and recovery block (RB) belong to the category of fault tolerant architecture, which usually considers both software and hardware.

Reliability is still employed as a system performance measure in most cases, but percentile life does provide a new perspective on optimal design without the requirement of a specified mission time. Availability is primarily used as the performance measure of renewable MSS whose optimal design has been emphasized and well developed in the past 10 years. Optimal design problems are generally formulated to maximize an appropriate system performance measure under resource constraints, and more realistic problems involving multiobjective programming are also being considered.

Optimal reliability design has attracted many researchers, who have produced hundreds of publications since 1960. Because of the increasing complexity of practical engineering systems and the critical importance of reliability in these complex systems, this still seems to be a very fruitful area for future research.

Compared to traditional binary-state systems, there are still many unsolved topics in MSS optimal design, which include the following:

- using percentile life as a system performance measure;
- involving cold-standby redundancy;
- nonrenewable MSS optimal design.

The research dealing with the understanding and application of reliability at the nano-level has also

demonstrated its attraction and vitality in recent years. Optimal system design that considers reliability within the uniqueness of nano-systems has seldom been reported in the literature. It deserves a lot more attention in the future. In addition, uncertainty and component dependency will be critical areas to consider in future research on optimal reliability design.

Acknowledgment

This work was partly supported by NSF Projects DMI-0429176. It is republished with the permission of *IEEE Transactions on Systems, Man, Cybernetics, Part A: Systems and Humans*.

Notations

x_j	number of components at subsystem j
r_j	component reliability at subsystem j
n	number of subsystems in the system
m	number of resources
x	$(x_1, \dots, x_n, r_1, \dots, r_n)$
$g_i(x)$	total amount of resource i required for x
R_S	system reliability
C_S	total system cost
R_0	a specified minimum R_S
C_0	a specified minimum C_S
α	system user's risk level, $0 < \alpha < 1$
$t_{\alpha,x}$	system percentile life, $t_{\alpha,x} = \inf\{t \geq 0 : R_S \leq 1 - \alpha\}$
E_S	generalized MSS availability index
E_0	a specified minimum E_S

References

- [1] Kuo, W. & Prasad, V.R. (2000). An annotated overview of system-reliability optimization, *IEEE Transactions on Reliability* **49**, 487–493.
- [2] Misra, K.B. (1986). On optimal reliability design: a review, *System Science* **12**, 5–30.
- [3] Tillman, F.A., Hwang, C.L. & Kuo, W. (1977). Optimization techniques for system reliability with redundancy- a review, *IEEE Transactions on Reliability* **R-26**, 148–152.
- [4] Taguchi, T. & Yokota, T. (1999). Optimal design problem of system reliability with interval coefficient using improved genetic algorithms, *Computers and Industrial Engineering* **37**, 145–149.
- [5] Kuo, W., Prasad, V.R., Tillman, F.A. & Hwang, C.L. (2001). *Optimal Reliability Design: Fundamentals and Applications*, Cambridge University Press.
- [6] Altıparmak, F., Dengiz, B. & Smith, A.E. (2003). Optimal design of reliable computer networks: a comparison of meta heuristics, *Journal of Heuristics* **9**, 471–487.
- [7] Amari, S.V., Dugan, J.B. & Misra, R.B. (1999). Optimal reliability of systems subject to imperfect fault-coverage, *IEEE Transactions on Reliability* **48**, 275–284.
- [8] Amari, S.V., Pham, H. & Dill, G. (2004). Optimal design of k-out-of-n: G subsystems subjected to imperfect fault-coverage, *IEEE Transactions on Reliability* **53**, 567–575.
- [9] Chen, T.C. & You, P.S. (2005). Immune algorithms-based approach for redundant reliability problems with multiple component choices, *Computers in Industry* **56**, 195–205.
- [10] Coit, D.W. & Liu, J. (2000). System reliability optimization with k-out-of-n subsystems, *International Journal of Reliability, Quality and Safety Engineering* **7**, 129–142.
- [11] Coit, D.W. (2001). Cold-standby redundancy optimization for nonrepairable systems, *IIE Transactions* **33**, 471–478.
- [12] Coit, D.W. (2003). Maximization of system reliability with a choice of redundancy strategies, *IIE Transactions* **35**, 535–543.
- [13] Djerdjour, M. & Rekab, K. (2001). A branch and bound algorithm for designing reliable systems at a minimum cost, *Applied Mathematics and Computation* **118**, 247–259.
- [14] Elegbede, C., Chu, C., Adjallah, K. & Yalaoui, F. (2003). Reliability allocation through cost minimization, *IEEE Transactions on Reliability* **52**, 106–111.
- [15] Ha, C. & Kuo, W. (2005). Multi-path approach for reliability-redundancy allocation using a scaling method, *Journal of Heuristics* **11**, 201–237.
- [16] Hsieh, Y.C., Chen, T.C. & Bricker, D.L. (1998). Genetic algorithm for reliability design problems, *Microelectronics Reliability* **38**, 1599–1605.
- [17] Hsieh, Y.C. (2002). A linear approximation for redundant reliability problems with multiple component choices, *Computers and Industrial Engineering* **44**, 91–103.
- [18] Kulturel-Konak, S., Smith, A.E. & Coit, D.W. (2003). Efficiently solving the redundancy allocation problem using tabu search, *IIE Transactions* **35**, 515–526.
- [19] Lee, C.Y., Gen, M. & Kuo, W. (2002). Reliability optimization design using hybridized genetic algorithm with a neural network technique, *IEICE Transactions on Fundamentals of Electronics, Communications and Computer Sciences* **E84-A**, 627–637.
- [20] Lee, C.Y., Yun, Y. & Gen, M. (2002). Reliability optimization design using hybrid NN-GA with fuzzy logic controller, *IEICE Transactions on Fundamentals of Electronics, Communications and Computer Sciences* **E85-A**, 432–447.

- [21] Lee, C.Y., Gen, M. & Tsujimura, Y. (2002). Reliability optimization design for coherent systems by hybrid GA with fuzzy logic controller and local search, *IEICE Transactions on Fundamentals of Electronics, Communications and Computer Sciences* **E85-A**, 880–891.
- [22] Levitin, G. (2005). Optimal structure of fault-tolerant software systems, *Reliability Engineering and System Safety* **89**, 286–295.
- [23] Liang, Y.C. & Smith, A.E. (2004). An ant colony optimization algorithm for the redundancy allocation problem, *IEEE Transactions on Reliability* **53**, 417–423.
- [24] Lin, F. & Kuo, W. (2002). Reliability importance and invariant optimal allocation, *Journal of Heuristics* **8**, 155–172.
- [25] Nahas, N. & Nourelfath, M. (2005). Ant system for reliability optimization of a series system with multiple-choice and budget constraints, *Reliability Engineering and System Safety* **87**, 1–12.
- [26] Kevin, Y.K.N.G. & Sancho, N.G.F. (2001). A hybrid dynamic programming/depth-first search algorithm with an application to redundancy allocation, *IIE Transactions* **33**, 1047–1058.
- [27] Prasad, V.R. & Raghavachari, M. (1998). Optimal allocation of interchangeable components in a series-parallel system, *IEEE Transactions on Reliability* **47**, 255–260.
- [28] Prasad, V.R., Kuo, W. & Kim, K.O. (2000). Optimal allocation of s-identical multi-functional spares in a series system, *IEEE Transactions on Reliability* **48**, 118–126.
- [29] Prasad, V.R. & Kuo, W. (2000). Reliability optimization of coherent systems, *IEEE Transactions on Reliability* **49**, 323–330.
- [30] Rocco, C.M., Moreno, J.A. & Carrasquero, N. (2000). A cellular evolution approach applied to reliability optimization of complex systems, in *Proceedings Annual Reliability and Maintainability Symposium*, Los Angeles, pp. 210–215.
- [31] Sung, C.S. & Cho, Y.K. (1999). Branch-and-bound redundancy optimization for a series system with multiple-choice constraints, *IEEE Transactions on Reliability* **48**, 108–117.
- [32] Sung, C.S. & Cho, Y.K. (2000). Reliability optimization of a series system with multiple-choice and budget constraints, *European Journal of Operational Research* **127**, 158–171.
- [33] Taguchi, T., Yokota, T. & Gen, M. (1998). Reliability optimal design problem with interval coefficients using hybrid genetic algorithms, *Computers and Industrial Engineering* **35**, 373–376.
- [34] Wattanapongsakorn, N. & Levitan, S.P. (2004). Reliability optimization models for embedded systems with multiple applications, *IEEE Transactions on Reliability* **53**, 406–416.
- [35] Yalaoui, A., Châtelet, E. & Chu, C. (2005). A new dynamic programming method for reliability & redundancy allocation in a parallel-series system, *IEEE Transactions on Reliability* **54**, 254–261.
- [36] You, P.S. & Chen, T.C. (2005). An efficient heuristic for series-parallel redundant reliability problems, *Computers and Operations Research* **32**, 2117–2127.
- [37] Zhao, R. & Liu, B. (2003). Stochastic programming models for general redundancy-optimization problems, *IEEE Transactions on Reliability* **52**, 181–191.
- [38] Zhao, R. & Liu, B. (2004). Redundancy optimization problems with uncertainty of combining randomness and fuzziness, *European Journal of Operation Research* **157**, 716–735.
- [39] Coit, D.W. & Smith, A.E. (1998). Redundancy allocation to maximize a lower percentile of the system time-to-failure distribution, *IEEE Transactions on Reliability* **47**, 79–87.
- [40] Coit, D.W. & Smith, A.E. (2002). Genetic algorithm to maximize a lower-bound for system time-to-failure with uncertain component Weibull parameters, *Computers and Industrial Engineering* **41**, 423–440.
- [41] Kim, K.O. & Kuo, W. (2003). Percentile life and reliability as a performance measures in optimal system design, *IIE Transactions* **35**, 1133–1142.
- [42] Prasad, V.R., Kuo, W. & Kim, K.O. (2001). Maximization of a percentile life of a series system through component redundancy allocation, *IIE Transactions*, **33**, 1071–1079.
- [43] Levitin, G., Lisnianski, A., Ben-Haim, H. & Elmakis, D. (1998). Redundancy optimization for series-parallel multi-state systems, *IEEE Transactions on Reliability* **47**, 165–172.
- [44] Levitin, G. & Lisnianski, A. (1998). Structure optimization of power system with bridge topology, *Electric Power Systems Research* **45**, 201–208.
- [45] Levitin, G. & Lisnianski, A. (1999). Joint redundancy and maintenance optimization for multistate series-parallel systems, *Reliability Engineering and System Safety* **64**, 33–42.
- [46] Levitin, G. & Lisnianski, A. (2000). Survivability maximization for vulnerable multi-state systems with bridge topology, *Reliability Engineering and System Safety* **70**, 125–140.
- [47] Levitin, G. (2001). Redundancy optimization for multi-state system with fixed resource-requirements and unreliable sources, *IEEE Transactions on Reliability* **50**, 52–59.
- [48] Levitin, G. & Lisnianski, A. (2001). Reliability optimization for weighted voting system, *Reliability Engineering and System Safety* **71**, 131–138.
- [49] Levitin, G. & Lisnianski, A. (2001). Structure optimization of multi-state system with two failure modes, *Reliability Engineering and System Safety* **72**, 75–89.
- [50] Levitin, G. & Lisnianski, A. (2001). Optimal separation of elements in vulnerable multi-state systems, *Reliability Engineering and System Safety* **73**, 55–66.
- [51] Levitin, G. & Lisnianski, A. (2001). A new approach to solving problems of multi-state system reliability optimization, *Quality and Reliability Engineering International* **17**, 93–104.

- [52] Levitin, G. (2002). Optimal allocation of multi-state retransmitters in acyclic transmission networks, *Reliability Engineering and System Safety* **75**, 73–82.
- [53] Levitin, G. (2002). Optimal allocation of elements in a linear multi-state sliding window system, *Reliability Engineering and System Safety* **76**, 245–254.
- [54] Levitin, G. (2002). Optimal series-parallel topology of multi-state system with two failure modes, *Reliability Engineering and System Safety* **77**, 93–107.
- [55] Levitin, G. (2003). Optimal allocation of multi-state elements in linear consecutively connected systems with vulnerable nodes, *European Journal of Operational Research* **150**, 406–419.
- [56] Levitin, G. (2003). Optimal allocation of multistate elements in a linear consecutively-connected system, *IEEE Transactions on Reliability* **52**, 192–199.
- [57] Levitin, G. & Lisnianski, A. (2003). Optimizing survivability of vulnerable series-parallel multi-state systems, *Reliability Engineering and System Safety* **79**, 319–331.
- [58] Levitin, G. (2003). Optimal multilevel protection in series-parallel systems, *Reliability Engineering and System Safety* **81**, 93–102.
- [59] Levitin, G., Dai, Y., Xie, M. & Poh, K.L. (2003). Optimizing survivability of multi-state systems with multilevel protection by multi-processor genetic algorithm, *Reliability Engineering and System Safety* **82**, 93–104.
- [60] Levitin, G. (2005). Uneven allocation of elements in linear multi-state sliding window system, *European Journal of Operational Research* **163**, 418–433.
- [61] Lisnianski, A., Levitin, G. & Ben-Haim, H. (2000). Structure optimization of multi-state system with time redundancy, *Reliability Engineering and System Safety* **67**, 103–112.
- [62] Massim, Y., Zeblah, A., Meziane, R., Benguediab, M. & Ghouraf A. (2005). Optimal design and reliability evaluation of multi-state series-parallel power systems, *Nonlinear Dynamics* **40**, 309–321.
- [63] Noureelfath, M. & Dutuit, Y. (2004). A combined approach to solve the redundancy optimization problem for multi-state systems under repair policies, *Reliability Engineering and System Safety* **84**, 205–213.
- [64] Ramirez-Marquez, J.E. & Coit, D.W. (2004). A heuristic for solving the redundancy allocation problem for multi-state series-parallel system, *Reliability Engineering and System Safety* **83**, 341–349.
- [65] Busacca, P.G., Marseguerra, M. & Zio, E. (2001). Multi-objective optimization by genetic algorithms: application to safety systems, *Reliability Engineering and System Safety* **72**, 59–74.
- [66] Coit, D.W., Jin, T. & Wattanapongsakorn, N. (2004). System optimization with component reliability estimation uncertainty: a multi-criteria approach, *IEEE Transactions on Reliability* **53**, 369–380.
- [67] Eleghbede, C. & Adjallah, K. (2003). Availability allocation to repairable systems with genetic algorithms: a multi-objective formulation, *Reliability Engineering and System Safety* **82**, 319–330.
- [68] Marseguerra, M., Zio, E., Podofillini, L. & Coit, D.W. (2005). Optimal design of reliable network systems in presence of uncertainty, *IEEE Transactions on Reliability* **54**, 243–253.
- [69] Ravi, V., Reddy, P.J. & Zimmermann, H.J. (2000). Fuzzy global optimization of complex system reliability, *IEEE Transactions on Fuzzy Systems* **8**, 241–248.
- [70] Sasaki, M. & Gen, M. (2003). A method of fuzzy multi-objective nonlinear programming with GUB structure by hybrid genetic algorithm, *International Journal of Smart Engineering Design* **5**, 281–288.
- [71] Sasaki, M. & Gen, M. (2003). Fuzzy multiple objective optimal system design by hybrid genetic algorithm, *Applied Soft Computing* **2**, 189–196.
- [72] Shelokar, P.S., Jayaraman, V.K. & Kulkarni, B.D. (2002). Ant algorithm for single and multiobjective reliability optimization problems, *Quality and Reliability Engineering International* **18**, 497–514.
- [73] Suman, B. (2003). Simulated annealing-based multi-objective algorithm and their application for system reliability, *Engineering Optimization* **35**, 391–416.
- [74] Bueno, V.C. (2005). Minimal standby redundancy allocation in a k-out-of-n: F system of dependent components, *European Journal of Operation Research* **165**, 786–793.
- [75] Romera, R., Valdés, J.E. & Zequeira, R.I. (2004). Active-redundancy allocation in systems, *IEEE Transactions on Reliability* **53**, 313–318.
- [76] Levitin, G. & Lisnianski, A. (1999). Importance and sensitivity analysis of multi-state systems using the universal generating function method, *Reliability Engineering and System Safety* **65**, 271–282.
- [77] Levitin, G. (2004). A universal generating function approach for the analysis of multi-state systems with dependent elements, *Reliability Engineering and System Safety* **84**, 285–292.
- [78] Levitin, G. (2003). Reliability of multi-state systems with two failure-modes, *IEEE Transactions on Reliability* **52**, 340–348.
- [79] Levitin, G. (2003). Linear multi-state sliding-window systems, *IEEE Transactions on Reliability* **52**, 263–269.
- [80] Levitin, G. (2002). Evaluating correct classification probability for weighted voting classifiers with plurality voting, *European Journal of Operational Research* **141**, 596–607.
- [81] Levitin, G. (2003). Reliability evaluation for acyclic transmission networks of multi-state elements with delays, *IEEE Transactions on Reliability* **52**, 231–237.
- [82] Levitin, G. (2001). Reliability evaluation for acyclic consecutively connected network with multistate elements, *Reliability Engineering and System Safety* **73**, 137–143.
- [83] Liu, P.X., Zuo, M.J. & Meng, M.Q. (2003). Using neural network function approximation for optimal design of continuous-state parallel-series systems, *Computers and Operations Research* **30**, 339–352.

- [84] Jin, T. & Coit, D.W. (2001). Variance of system-reliability estimates with arbitrarily repeated components, *IEEE Transactions on Reliability* **50**, 409–413.

Further Reading

- Cui, L., Kuo, W., Loh, H.T. & Xie, M. (2004). Optimal allocation of minimal & perfect repairs under resource constraints, *IEEE Transactions on Reliability* **53**, 193–199.
- Gupta, R. & Agarwal, M. (2006). Penalty guided genetic search for redundancy optimization in multi-state series-parallel power system, *Journal of Combinatorial Optimization* **12**, 257–277.
- Ha, C. & Kuo, W. (2006). Reliability redundancy allocation: an improved realization for nonconvex nonlinear programming problems, *European Journal of Operational Research* **171**, 24–38.
- Ha, C. & Kuo, W. (2006). Multi-paths iterative heuristic for redundancy allocation: the tree heuristic, *IEEE Transactions on Reliability* **55**, 37–43.
- Konak, A., Coit, D.W. & Smith, A.E. (2006). Multi-objective optimization using genetic algorithms: a tutorial, *Reliability Engineering and System Safety* **91**, 992–1007.
- Kuo, W., Chien, K. & Kim, T. (1998). *Reliability, Yield and Stress Burn-in*, Kluwer.
- Kuo, W. & Zuo, M.J. (2003). *Optimal Reliability Modeling: Principles and Applications*, John Wiley & Sons.
- Levitin, G. (2004). Consecutive k-out-of-r-from-n system with multiple failure criteria, *IEEE Transactions on Reliability* **53**, 394–400.
- Levitin, G. (2004). Reliability optimization models for embedded systems with multiple applications, *IEEE Transactions on Reliability* **53**, 406–416.
- Lin, F., Kuo, W. & Hwang, F. (1999). Structure importance of consecutive-k-out-of-n systems, *Operations Research Letters* **25**, 101–107.
- Lisianski, A. & Levitin, G. (2003). *Multi-State System Reliability, Assessment, Optimization and Applications*, World Scientific Publishing.
- Mahapatra, G.S. (2006). Fuzzy multi-objective mathematical programming on reliability optimization model, *Applied Mathematic and Computation* **174**, 643–659.
- Marseguerra, M. & Zio, E. (2000). Optimal reliability/availability of uncertain systems via multi-objective genetic algorithms, Proceedings Annual Reliability and Maintainability Symposium, Los Angeles, pp. 222–227.
- Salazar, D., Rocco, C.M. & Galvan, B.J. (2006). Optimization of constrained multiple-objective reliability problems using evolutionary algorithms, *Reliability Engineering and System Safety* **91**, 1057–1070.
- Tian, Z. & Zuo, M.J. (2006). Redundancy allocation for multi-state systems using physical programming and genetic algorithms, *Reliability Engineering and System Safety* **91**, 1049–1056.
- Yamachi, H., Tsujimuraa, Y., Kambayashia, Y. & Yamamoto, H. (2006). Multi-objective algorithm for solving N-version program design problem, *Reliability Engineering and System Safety* **91**, 1083–1094.
- Yang, J.-E., Hwang, M.-J., Sung, T.-Y. & Jin, Y. (1999). Application of genetic algorithm for reliability allocation in nuclear power plants, *Reliability Engineering and System Safety* **65**, 229–238.
- Yu H., Yalaoui, F., Châtelet, E. & Chu, C. (2007). Optimal design of a maintainable cold-standby system. *Reliability Engineering and System Safety* **92**, 85–91.
- Zafropoulos, E.P. & Dialynas, E.N. (2004). Reliability and cost optimization of electronic devices considering the component failure rate uncertainty, *Reliability Engineering and System Safety* **84**, 271–284.

WAY KUO AND RUI WAN

Expectation Maximization Algorithm

Introduction

Dempster *et al.* [1] introduce the expectation maximization (EM) algorithm as a technique to produce maximum-likelihood estimates (MLEs) in problems where the data are incomplete. Latent variables or missing data points can be an intrinsic feature of the problem or they can be artificially imputed to ease computation.

The EM algorithm appeared in various guises before the paper by Dempster *et al.* [1], but this paper provided a unifying framework for the algorithm. Newcomb [2] is one of the earliest known works using the EM algorithm in which it is used for fitting a mixture of normal distributions. Blight [3] finds the MLE of an exponential family for censored data using an EM algorithm.

In principle, the EM process is straightforward to implement and has broad applicability. It is a two-step iterative algorithm consisting of an expectation step (E-step) followed by a maximization step (M-step). During the E-step, the expected value of the log-likelihood of the complete data (i.e., the observed and unobserved data) is computed. In the M-step, the expected complete data log-likelihood is maximized producing an MLE of the model parameters. The parameter estimates produced during the M-step are then used in a new E-step and the cycle continues until convergence.

While the EM algorithm will not decrease the observed data likelihood function (see Dempster *et al.* [1]), there is no guarantee that the resulting sequence of parameter estimates will converge to the global maximum of the likelihood. Multiple runs of the algorithm from different parameter starting values should help avoid local maxima.

McLachlan and Krishnan [4] provide an excellent review of the EM algorithm and its extensions. Other good resources include Pawitan [5, Chapter 12] and Tanner [6, Chapter 4].

A summary of the EM algorithm is provided in the next section including methods for determining convergence (Convergence Diagnostics). Many generalizations of the EM algorithm have been developed; a selection are introduced in the section

titled “Generalizations”. The section titled “Implementation of the EM Algorithm” illustrates implementation within the context of a finite mixture model. In conclusion, some properties of the EM algorithm are discussed.

EM Algorithm

Let $\mathbf{x} = (x_1, x_2, \dots, x_M)$ denote the observed data and $\mathbf{z} = (z_1, z_2, \dots, z_M)$ denote the unobserved (or latent) data. The likelihood function is $L(\phi) = p(\mathbf{x}|\phi)$ and a MLE $\hat{\phi}$ satisfies $L(\hat{\phi}) \geq L(\phi)$ for all ϕ . The complete data likelihood is of the form $L_c(\phi) = p(\mathbf{x}, \mathbf{z}|\phi)$.

The EM algorithm involves the following steps.

1. Let $\phi^{(0)}$ be an initial estimate of the unknown parameter. Let $t = 0$.
2. E-step: Compute

$$Q(\phi, \phi^{(t)}) = \int_{\mathbf{z}} p(\mathbf{z}|\mathbf{x}, \phi^{(t)}) \log p(\mathbf{x}, \mathbf{z}|\phi) d\mathbf{z} \quad (1)$$

3. M-step: Maximize $Q(\phi, \phi^{(t)})$ with respect to ϕ and let $\phi^{(t+1)}$ be the maximizing value.
4. Check for convergence. If converged, stop. Otherwise, increment t and return to Step 2.

The EM algorithm is guaranteed to give a sequence of estimates that have monotonically non-decreasing likelihood values, that is $L(\phi^{(t+1)}) \geq L(\phi^{(t)})$. The algorithm can converge to local maxima of the likelihood function. Wu [7] and McLachlan and Krishnan [4, Chapter 3] provide a detailed overview of the convergence properties and stationary points of the EM algorithm.

Convergence Diagnostics

Convergence is a feature of any iterative algorithm which must be considered. Aitken’s acceleration criterion [8–10] is often used as a convergence diagnostic for the EM algorithm.

Denote by $\{l^{(t)} = \log L(\phi^{(t)})\}$ the sequence of log-likelihood values that emerge from the EM algorithm. Assuming l^* is the limiting value of these log-likelihoods for iteration $(t + 1)$ of the EM algorithm, it follows that

$$l^{(t+1)} - l^* \approx a(l^{(t)} - l^*) \quad \text{for } 0 \leq a \leq 1 \quad (2)$$

2 Expectation Maximization Algorithm

$$\Rightarrow l^{(t+1)} - l^{(t)} \approx (1 - a)(l^* - l^{(t)}) \quad (3)$$

Thus, if $a \approx 1$, a small increase in the log-likelihood value does not necessarily imply that the EM algorithm is close to convergence. It can be shown that estimating a by the ratio of successive increments leads to the Aitken acceleration estimate of the limiting log-likelihood value

$$l^* \approx l^{(t)} - \frac{\{l^{(t+1)} - l^{(t-1)}\}}{\{l^{(t)} - l^{(t-1)}\}} \{l^{(t+1)} - l^{(t)}\} \quad (4)$$

The EM algorithm should be stopped if $|l^{(t+1)} - l^*| < \epsilon$, where ϵ is a prespecified tolerance level. Since this criterion is not an exact indicator of convergence, multiple runs of the algorithm with random starts must be employed.

Generalizations

Many extensions of the EM algorithm have also been developed. Dempster *et al.* [1] define the generalized expectation maximization (GEM) algorithm which increases the value of the likelihood at the M-step without actually completing a full maximization. The GEM algorithm weakens the requirement of maximizing $Q(\phi, \phi^{(t)})$ with respect to ϕ to a requirement of finding $\phi^{(t+1)}$ such that $Q(\phi^{(t+1)}, \phi^{(t)}) \geq Q(\phi^{(t)}, \phi^{(t)})$.

The expectation conditional maximization (ECM) algorithm [11] is a GEM algorithm for multivariate parameter problems which replaces the maximization of $Q(\phi, \phi^{(t)})$ with respect to ϕ with a series of conditional maximizations with respect to components of the ϕ vector. The multicycle ECM algorithm [11] is a variant of the ECM algorithm that maximizes only with respect to a subset of the components of the ϕ vector between E-steps; the multicycle ECM algorithm can converge faster or slower than the ECM algorithm.

The **Monte Carlo** expectation maximization (MCEM) algorithm [12, 13] approximates the expectation calculation in the E-step with a Monte Carlo estimate of

$$Q(\phi, \phi^{(t)}) \approx \frac{1}{S} \sum_{s=1}^S \log p(\mathbf{x}, \mathbf{z}^{(s)} | \phi) \quad (5)$$

where $\mathbf{z}^{(s)}$ are samples from $p(\mathbf{z} | \mathbf{x}, \phi^{(t)})$.

The stochastic expectation maximization (SEM) algorithm [14] for mixture models can be viewed as a special case of the MCEM algorithm with $S = 1$.

The alternating expectation conditional maximization (AECM) algorithm [15] is an extension of the ECM algorithm that allows for different definitions of the complete data at some of the conditional M-steps; this approach is advantageous in many latent variable modeling situations.

McLachlan and Krishnan [4] provide extensive details of further extensions of the EM algorithm, with an emphasis on algorithms that accelerate the EM convergence. Lange *et al.* [16] and Hunter and Lange [17] show that the EM algorithm fits into a larger class of optimization algorithms which use majorization techniques.

Implementation of the EM Algorithm

Finite mixture models provide an example of a framework in which the use of the EM algorithm is particularly advantageous. Mixture models arise in reliability analysis where failures occur due to multiple causes (e.g., [18, 19]). Extensive reviews of mixture models and their applications are given by Everitt and Hand [20], Titterton *et al.* [21], McLachlan and Basford [22], and McLachlan and Peel [10].

A finite mixture model uses a finite collection of K components to construct a model; in reliability analysis, these could be K different modes of failure. It is assumed that the probability of mixture component k is π_k – these values are called *mixing proportions*. In addition, an observation x_i from mixture component k has a probability density $p(x_i | \theta_k)$, where θ_k denotes unknown model parameters. Hence, the resulting model with parameter $\phi = (\underline{\pi}, \underline{\theta})$ for a single observation x_i is

$$p(x_i | \phi) = p(x_i | \underline{\theta}, \underline{\pi}) = \sum_{k=1}^K \pi_k p(x_i | \theta_k) \quad (6)$$

that is, a K -component mixture model.

Directly finding MLEs of the parameters in the mixture model is complicated because the likelihood function involves a product of terms of the form equation (6). However, the use of the EM algorithm allows for relatively straightforward model fitting for many component models, $p(x_i | \theta_k)$.

Finite mixture models can be formulated as having the component membership labels for each observation as missing data. Maximization of the likelihood function is simplified by augmenting the data to include the missing membership labels for each observation. Furthermore, the EM algorithm provides estimates not only of the model parameters but also of the unknown component memberships of the observations.

Suppose we have observed data $\mathbf{x} = (x_1, x_2, \dots, x_M)$ that have arisen from a K -component mixture model. The term *complete* data refers to the combination of both the observed data and the missing membership variables. The complete data is denoted by $(\mathbf{x}, \mathbf{z}) = \{(x_1, z_1), (x_2, z_2), \dots, (x_M, z_M)\}$, where

$$z_i = (z_{i1}, z_{i2}, \dots, z_{iK}) \quad \forall i = 1, \dots, M \quad (7)$$

with

$$z_{ik} = \begin{cases} 1 & \text{if observation } x_i \text{ belongs to component } k \\ 0 & \text{otherwise} \end{cases}$$

That is, the missing data z_i is an indicator of the component membership of observation x_i . On convergence of the EM algorithm, the estimated values of z_{ik} are the estimated conditional probabilities of observation x_i belonging to component k .

Under a finite mixture model, the complete data log-likelihood is formulated as follows. The density of the observed data is

$$p(\mathbf{x}|\underline{\theta}, \underline{\pi}) = \prod_{i=1}^M \sum_{k=1}^K \pi_k p(x_i|\theta_k) \quad (8)$$

Accounting for the missing component membership labels $\mathbf{z}$, it follows that the complete data likelihood is

$$p(\mathbf{x}, \mathbf{z}|\underline{\theta}, \underline{\pi}) = \prod_{i=1}^M \prod_{k=1}^K \{\pi_k p(x_i|\theta_k)\}^{z_{ik}} \quad (9)$$

The EM algorithm involves two steps, an E-step followed by an M-step. The E-step involves calculating the expectation of the complete data log-likelihood. Since the complete data log-likelihood is linear in the missing variables, the E-step simply replaces each z_{ik} by its expected value. The M-step then proceeds to maximize the expected complete data log-likelihood to estimate the model parameters.

Specifically, the EM algorithm proceeds as follows:

1. Initialize: Choose starting values for $\underline{\pi}^{(0)}$ and $\underline{\theta}^{(0)}$. Let $t = 0$.
2. E-step: Compute

$$\hat{z}_{ik}^{(t)} = \frac{\pi_k^{(t)} p(x_i|\theta_k^{(t)})}{\sum_{k'=1}^K \pi_{k'}^{(t)} p(x_i|\theta_{k'}^{(t)})} \quad (10)$$

for $i = 1, \dots, M$ and $k = 1, \dots, K$ where $\hat{z}_{ik}^{(t)}$ is the estimated posterior probability of observation x_i belonging to component k .

3. M-step: Maximize the complete data log-likelihood function

$$\sum_{i=1}^M \sum_{k=1}^K \hat{z}_{ik}^{(t)} \{\log \pi_k + \log p(x_i|\theta_k)\}$$

to yield new parameter estimates $\underline{\pi}^{(t+1)}$ and $\underline{\theta}^{(t+1)}$. Increment t by 1.

4. Convergence: Repeat the E-step and M-step until convergence (see the section titled ‘‘Convergence Diagnostics’’). The final parameter values are the MLEs $\hat{\underline{\pi}}$ and $\hat{\underline{\theta}}$.

Finite Mixture of Exponentials

The finite mixture of exponential distributions can be used to model data that has multiple modes of failure [18]. Suppose that there are K modes of failure which occur with probabilities $\pi_1, \pi_2, \dots, \pi_K$. In addition, suppose that lifetimes of components within each mode of failure are modeled using an exponential distribution with mean θ_k . Hence, the data are modeled using a finite mixture of exponentials, that is,

$$p(x_i) = \sum_{k=1}^K \pi_k p(x_i|\theta_k) = \sum_{k=1}^K \frac{\pi_k}{\theta_k} \exp\left(-\frac{x_i}{\theta_k}\right) \quad (11)$$

Let z_{ik} be an indicator for the unknown mode of failure of item i . The joint density of the complete data (x_i, z_i) is

$$p(x_i, z_i) = \prod_{k=1}^K \left\{ \frac{\pi_k}{\theta_k} \exp\left(-\frac{x_i}{\theta_k}\right) \right\}^{z_{ik}} \quad (12)$$

The EM algorithm proceeds as follows:

1. Initialize: Choose starting values for $\underline{\pi}^{(0)}$ and $\underline{\theta}^{(0)}$. Let $t = 0$.

4 Expectation Maximization Algorithm

2. E-step: Compute

$$\hat{z}_{ik}^{(t)} = \frac{(\pi_k^{(t)} / \theta_k^{(t)}) \exp(-x_i / \theta_k^{(t)})}{\sum_{k'=1}^K (\pi_{k'}^{(t)} / \theta_{k'}^{(t)}) \exp(-x_i / \theta_{k'}^{(t)})} \quad (13)$$

for $i = 1, \dots, M$ and $k = 1, \dots, K$.

The current value of the log-likelihood can be calculated as

$$l^{(t)} = \sum_{i=1}^M \log \left\{ \sum_{k=1}^K \frac{\pi_k^{(t)}}{\theta_k^{(t)}} \exp \left(-\frac{x_i}{\theta_k^{(t)}} \right) \right\} \quad (14)$$

which only uses terms required for computing $\hat{z}_{ik}^{(t)}$.

3. M-step: Maximize the complete data log-likelihood function

$$\sum_{i=1}^M \sum_{k=1}^K \hat{z}_{ik}^{(t)} \left\{ \log \pi_k - \log \theta_k - \frac{x_i}{\theta_k} \right\}$$

to yield new parameter estimates

$$\pi_k^{(t+1)} = \frac{\sum_{i=1}^M \hat{z}_{ik}^{(t)}}{M} \text{ and } \theta_k^{(t+1)} = \frac{\sum_{i=1}^M \hat{z}_{ik}^{(t)} x_i}{\sum_{i=1}^M \hat{z}_{ik}^{(t)}} \quad (15)$$

Increment t by 1.

4. Convergence: Repeat the E-step and M-step until convergence. The final parameter values are MLEs $\hat{\pi}$ and $\hat{\theta}$.

Properties

Convergence

The EM algorithm is a linearly convergent algorithm in the neighborhood of the MLE, that is,

$$\|\phi^{(t+1)} - \hat{\phi}\| < a \|\phi^{(t)} - \hat{\phi}\| \quad (16)$$

for some value $a < 1$. Other iterative optimization methods (e.g., Newton–Raphson) have quadratic convergence, that is,

$$\|\phi^{(t+1)} - \hat{\phi}\| < a \|\phi^{(t)} - \hat{\phi}\|^2 \quad (17)$$

for some $a < 1$. Hence, one criticism of the EM algorithm is its slow convergence rate. However, the EM algorithm is very stable and for many problems it involves simple calculations; for many problems it does not involve computing and inverting a Hessian matrix (in contrast to Newton–Raphson). Hence, the EM algorithm is still used for a wide variety of problems where faster alternatives exist.

Standard Errors

Early criticisms of the EM algorithm focused on the fact that the algorithm does not automatically produce standard errors (or an approximate **covariance** matrix of the MLEs). However, McLachlan and Krishnan [4] and McLachlan and Peel [10] detail methodologies that provide approximate standard errors of parameter estimates derived during an EM algorithm. In particular, they demonstrate ways in which the observed information matrix may be calculated within the EM framework.

The observed information matrix could be directly estimated after convergence of the EM algorithm, but in general evaluation of the second-order derivatives of the log-likelihood is cumbersome. However, the observed information matrix can be expressed in terms of derivatives of the complete data log-likelihood which are introduced within the EM algorithm.

McLachlan and Krishnan [4] deal with the particular case of independent data where standard errors can be produced without additional work beyond the necessary EM algorithm calculations. Meilijson [23] defines the *empirical* observed information matrix which can be expressed in terms of the score functions of the complete data log-likelihood. Thus the empirical observed information matrix can be constructed from the final estimates produced by the EM algorithm. Computation of second-order partial derivatives is therefore avoided. The covariance matrix of $\hat{\phi}$ is then approximated by the inverse of the empirical observed information matrix.

References

- [1] Dempster, A.P., Laird, N.M. & Rubin, D.B. (1977). Maximum likelihood from incomplete data via the EM algorithm (With discussion), *Journal of the Royal Statistical Society. Series B* **39**(1), 1–38.

-
- [2] Newcomb, S. (1886). A generalized theory of the combination of observations so as to obtain the best result, *American Journal of Mathematics* **8**(4), 343–366.
 - [3] Blight, B.J.N. (1970). Estimation from a censored sample for the exponential family, *Biometrika* **57**, 389–395.
 - [4] McLachlan, G.J. & Krishnan, T. (1997). *The EM Algorithm and Extensions*, John Wiley & Sons, New York.
 - [5] Pawitan, Y. (2001). *In All Likelihood: Statistical Modelling and Inference Using Likelihood*, Oxford University Press, Oxford.
 - [6] Tanner, M.A. (1996). *Tools for Statistical Inference*, 3rd Edition, Springer-Verlag, New York.
 - [7] Wu, C.-F.J. (1983). On the convergence properties of the EM algorithm, *The Annals of Statistics* **11**(1), 95–103.
 - [8] Böhning, D., Dietz, E., Schaub, R., Schlattmann, P. & Lindsay, B. (1994). The distribution of the likelihood ratio for mixtures of densities from the one-parameter exponential family, *Annals of the Institute of Statistical Mathematics* **46**, 373–388.
 - [9] Lindsay, B. (1995). *Mixture Models: Theory, Geometry and Applications*, Institute of Mathematical Statistics, Hayward.
 - [10] McLachlan, G.J. & Peel, D. (2000). *Finite Mixture Models*, John Wiley & Sons, New York.
 - [11] Meng, X.-L. & Rubin, D.B. (1993). Maximum likelihood estimation via the ECM algorithm: a general framework, *Biometrika* **80**(2), 267–278.
 - [12] Wei, G.C.G. & Tanner, M.A. (1990). A Monte Carlo implementation of the EM algorithm and the poor man's data augmentation algorithms, *Journal of the American Statistical Association* **85**, 699–704.
 - [13] Booth, J.G. & Hobert, J.P. (1999). Maximizing generalized linear mixed model likelihoods with an automated Monte Carlo EM algorithm, *Journal of the Royal Statistical Society. Series B* **61**, 265–285.
 - [14] Celeux, G. & Diebolt, J. (1985). The SEM algorithm: a probabilistic teacher algorithm derived from the EM algorithm for the mixture problem, *Computational Statistics Quarterly* **2**, 73–82.
 - [15] Meng, X.-L. & VanDyk, D. (1997). The EM algorithm – an old folk song sung to a new fast tune (with discussion), *Journal of the Royal Statistical Society. Series B* **59**(3), 511–567.
 - [16] Lange, K., Hunter, D.R. & Yang, I. (2000). Optimization transfer using surrogate objective functions, *Journal of Computational and Graphical Statistics* **9**(1), 1–59; With discussion, and a rejoinder by Hunter and Lange.
 - [17] Hunter, D.R. & Lange, K. (2004). A tutorial on MM algorithms, *The American Statistician* **58**(1), 30–37.
 - [18] Davis, D.J. (1952). Analysis of some failure data, *Journal of the American Statistical Association* **47**, 113–150.
 - [19] Fowlkes, E.B. (1979). Some methods for studying the mixture of two normal (lognormal) distributions, *Journal of the American Statistical Association* **74**, 561–575.
 - [20] Everitt, B.S. & Hand, D.J. (1981). *Finite Mixture Distributions*, Chapman & Hall, London.
 - [21] Titterton, D.M., Smith, A.F.M. & Makov, U.E. (1985). *Statistical Analysis of Finite Mixture Distributions*, John Wiley & Sons, Chichester.
 - [22] McLachlan, G.J. & Basford, K.E. (1988). *Mixture Models: Inference and Applications to Clustering*, Marcel Dekker, New York.
 - [23] Meilijson, I. (1989). A fast improvement to the EM algorithm on its own terms, *Journal of the Royal Statistical Society. Series B* **51**(2), 127–138.

ISOBEL C. GORMLEY AND THOMAS B.
MURPHY

Dynamic Programming Methods in Repair and Replacement

Introduction

In this article, we focus on the usefulness of dynamic programming (DP) ideas and methods in the study of repairable systems with a specific emphasis on replacement and other repair strategies. After an overview of early work in this area, which concentrated on optimal planned replacement of stochastically failing equipment, we briefly survey later development on different formulations of more pragmatic repair schemes that have been proposed over the years, as a prelude to demonstrating how different notions of imperfect repair can be encompassed in a unified stochastic DP setting.

The name *dynamic programming* was coined by Richard Bellman as a shorthand description of a body of theory and methods for solving optimization problems in multistage decision processes, illustrated in his now classic book [1]. Such problems are typically characterized by a system described *via* proxy variables (*system states*) that change (undergo deterministic or, stochastic *transitions*) as a result of decisions (*actions*) chosen from a set of available options. Corresponding to each transition is a reward or, penalty (negative reward) which depends on the chosen action associated with the transition as well as the system state before and/or after the transition. This describes a typical stage of the multistage process, which is then repeated in the next stage. The planning horizon of a DP problem is the finite or infinite number of stages for which decisions on the choice of actions have to be made. A policy, or, strategy is any well-defined rule to determine the choice of actions based possibly on the system's past history to date. The goal is to maximize the value (expected value, under stochastic regimes) of the stream of rewards over the entire planning horizon, and identify strategies which achieve, or nearly achieve the optimum value function. The "value" of a reward stream is usually defined as the total cumulative reward with or without discounting. In some problems, the average reward per unit time is considered as an alternative value function.

Such optimization problems over a finite, or infinite horizon is conceptually equivalent to finding the global maximum of an objective function with respect to a large number of variables – a task that is often fraught with analytical as well as computational difficulties, referred to collectively as the *curse of dimensionality*. Bellman's basic simplifying idea, embodied in his celebrated "principle of optimality" was to break down the problem into one of sequential choices, which together with backward induction can routinely solve the optimization problem when the planning horizon and the set of available actions are both finite. Over the years, these methods have correspondingly found extensive use in many application areas in engineering and operations research. Subsequently, in a series of papers, Blackwell [2, 3] and Strauch [4] set the original ideas of Bellman on a firm theoretical foundation, including a general (and not problem specific) proof of Bellman's principle of optimality, that covers both finite and infinite horizon problems. It is this conceptual framework on which rests much of our following discussions on DP formulations of repairable systems. It is important to remark at this point that a DP formulation to describe the temporal evolution of a system can be a useful vehicle for its modeling and analysis even when there is no choice of actions to be made; i.e., when there is no explicit optimization task at hand, as our discussions in the section titled "Illustrative Applications" will illustrate.

The basic premise of the importance of repairs, apart from generally lesser costs compared to replacement, is that it allows the reuse of equipment, even though such repaired equipment will typically have altered reliability characteristics compared to the original as a result of repair – as reflected in the distribution of postrepair survival times. Contemporary concerns about finite environmental resources further underscores the importance of such "recycling" of equipment and systems. Additionally, there may be economic incentives for considering replacement or, less than perfect repair options, since such repairs are typically less costlier than replacement. The impact of a repair strategy to maintain an equipment is reflected in the point process of failures as it evolves over time. In applications, the effectiveness of a repair strategy is determined by appropriate objective functions that are statistical functionals of the underlying failure process. Examples of such objective functions are (a) the expected number of failures/repairs over a given

time interval, (b) the total expected costs/benefits of a repair scheme, and so on. It is here in the modeling and analysis of the effectiveness of such repairs that DP ideas and methods can be fruitfully exploited.

To the best of our knowledge, as of today, there is no standard commercial or, open-source general-purpose DP software that would allow users to conveniently specify and solve arbitrary DP problems. Instead, the most common and fruitful approach to the computational implementation of DP solutions to real-world problems appears to be developing application-specific DP software. In our search for DP software on repair and replacement, we came across only one such modular software “CompEcon (toolbox for MATLAB)” which includes solvers for discrete-time/discrete, or continuous variable DP problems, and was developed to accompany a book by Miranda and Feckler [5]. The “CompEcon Toolbox” can be found at <http://www4.ncsu.edu/~pfackler/compecon/toolbox.html>. A similar search for a general purpose DP software, led us to “dProg”, a software package for specifying DP algorithms that – given a recursive definition of the problem – generates codes for solving the problem using DP. The “dProg” package, developed by Thomas Mailund, is available as a freeware at <http://www.daimi.au.dk/~mailund/dprog.html>, although it may be of relatively limited use in our context of repairable systems.

The section titled “Early Work, Minimal Repairs, and Other Imperfect Repair Schemes” describes some early work exploiting DP ideas to investigate problems in reliability, followed by a summary of relevant modeling approaches adopted by different authors to address possible notions of imperfect repairs (IRs). It is then shown that these different formulations can be synthesized in a common setting. A general DP formulation for the analysis of a repairable system is then described. Various IR schemes proposed by different authors are then seen to be special cases. The section titled “Illustrative Applications” shows the usefulness of a DP formulation to study repair and replacement problems drawing on some recent work in the literature.

Early Work, Minimal Repairs, and Other Imperfect Repair Schemes

Early Work and Repair-Replacement

The early literature on the system maintenance aspect of reliability naturally focused on replacements as

the analytically tractable option for modeling system repair, since the tools of renewal theory, a well-developed theme in probability, are then immediately applicable. By contrast, attention to modeling other relatively more pragmatic notions of repair had to wait until the 1970s and 1980s. Although not directly related to the theme of repair and replacement, a work by Bellman and Dryfus [6] appears to be the earliest reference we can find of the use of DP methods for maximizing system reliability under given constraints. Another example of such use of DP, unrelated to replacements and repair, is a work by Kettelle [7] on optimum allocation of **redundancy**.

Among the early classic illustrations of the application of DP methods for the analysis of optimal replacement policies are the works by Jorgenson and Radner [8] and Radner and Jorgenson [9]. They consider a series system of $(M + 1)$ independent components, $M \geq 1$, when components $i = 1, 2, \dots, M$ have exponential lifetimes that can be continuously monitored. The remaining other component (component $i = 0$) is assumed only to have a positive density with infinite support, and cannot be inspected except at replacement times. Using a DP formulation, the authors construct an optimal replacement policy, called the (n_i, N) policy, that maximizes the total expected discounted value of time that the system is in good operational state, less costs, and can be described as follows. For $i = 1, 2, \dots, M$, (a) replace component i alone when it fails, if the corresponding age of component 0 is between 0 and n_i , (b) replace components i and 0 together, when component i fails and component 0 has age between n_i and N at that time, and (c) replace component 0 alone, if all components $1, 2, \dots, M$ have survived until component 0 has reached age N . The same policy remains optimal for other alternative criteria such as expected total discounted time in good state, and the ratio of expected total discounted good time to expected total discounted costs.

Work on **replacement strategies** have continued over the years, with researchers addressing different variations on the assumptions regarding failure and repair time distributions, rewards and costs with or without discounting, finite or infinite planning horizon, and the objective function (suitably defined total or long-run average net reward). Chen and Feldman [10] consider optimal replacement policies with **minimal repair** (MR) and age-dependent costs. Also notable among works on the optimal

repair–replacement theme are Lam Yeh [11, 12], Stadjé and Zuckerman [13], and Boshuizen [14]. In particular, Boshuizen’s [14] analysis derives the optimal rules without imposing some reliability theoretic and monotonicity conditions, as employed in Lam Yeh [11, 12] and Stadjé and Zuckerman [13], by solving an optimal stopping problem first in a finite horizon setting, and then extending it to the infinite horizon case.

Barlow and Proschan [15] investigated sequential replacement policies for finite-time horizons in which the next scheduled replacement age depends on the time remaining. They prove that without loss of generality, it is enough to restrict attention to nonrandomized replacement policies, prove the existence of a least-cost sequential policy and demonstrate how to construct them. Although there is no direct reference to DP methods in their arguments, since a formal theory of stochastic DP [2, 3, 4, 16] was still in its relative infancy, such ideas are clearly recognizable from a careful reading of their arguments for optimality.

Minimal Repairs and Other Imperfect Repair Schemes

A repair mode is defined by its effect on the equipment as described by the distribution of its postrepair life. If X is the lifelength of a new unit and L_x denotes the lifelength of an instantaneously repaired unit that failed at age x , then setting $L_x \stackrel{d}{=} X$ defines a *replacement* by a statistically identical copy ($\stackrel{d}{=}$ denoting equality in distribution), whereas $L_x \stackrel{d}{=} X - x | X > x$, the *residual life* at the age x at failure, defines the so-called MR option. A *perfect repair* (PR), by contrast, is defined as a repair that restores an equipment to its original condition, as parameterized by its *age* when it started working prior to the most recent failure. Thus, for a currently failed equipment with failure having occurred at age x (which implies that it last started functioning at some age t satisfying $0 \leq t < x$), a PR is defined by setting $L_x \stackrel{d}{=} X - x | X > x$. Note, a renewal (or replacement) is equivalent to PR only if the currently failed unit was new when it started working.

For most equipments that exhibit some form of aging over time, replacement and MR conceptually represent two extreme positions in terms of the extent to which a failed equipment’s **degradation** is arrested and restored by repairs. The need for considering

more realistic notions of repair, which are in a broad sense intermediate between MR and PR, was mainly voiced by the engineering community, which then led to efforts to formulate alternative notions of IRs. The notions of IR, refer to those repair assumptions that are, in a broad sense, more realistic than the traditional renewal (replacement) hypothesis.

The first IR model was proposed by Brown and Proschan [17]. It underwent a significant and nontrivial extension by Block *et al.* [18]. The Block–Borges–Savits (BBS) model randomizes the choice of repair between a replacement/PR or MR with a probability that is age dependent (choosing a replacement/PR or MR with probabilities $p(x)$ ($1 - p(x)$, respectively), if failure occurred at age x). The Brown–Proschan (BP) model is covered by the BBS model as a special case when $p(x)$ is a constant. Bhattacharjee [19] provided an alternative shock model argument for the distribution of the sojourn time between consecutive PRs, possibly interspersed with a random number of MRs.

Kijima [20] proposed and investigated two models of IR, based on the idea that if an equipment failed at a certain age, after repair, its age is not necessarily restored to the age when it failed (as is the case with MR), but may be different depending on how good or bad the repair is. This effective postrepair age is referred to as the *virtual age*. In Kijima’s models, the virtual age process V_k after the k th failure is linear in V_{k-1} and the lifetime X_k after the $(k - 1)$ st failure (the k th functioning lifetime) such that the slope of V_k depends on the “degree of repair” a assuming values in $[0, 1]$ with $a = 0, 1$ implying PRs and MRs, respectively, whereas intermediate values of $a \in (0, 1)$ reflect varying degrees of repair effectiveness.

Several other virtual age type models, which are basically extensions of Kijima’s models, have been considered by Finkelstein [21], Makis and Jardine [22], and Dagpunar [23]. Virtual age models capture the idea of arresting degradation in performance by system age recovery as a result of repairs. Stadjé and Zuckerman [24, 25] investigate analytical and numerical methods to determine least expected discounted cost age-dependent repair strategies that allow for varying degrees of repair over an infinite planning horizon. Guo *et al.* [26], in a review article, provide a summary of various virtual age models and provide references to work on optimal repair and replacement policies based on such models.

4 Dynamic Programming Methods in Repair and Replacement

A Unified Setting for Imperfect Repair (IR) Models

Bhattacharjee [27] argued that a reasonable way to conceptualize the effect of a given repair mode or action a is to make the repaired unit have an effective *virtual age* Y after repair, whose distribution A_z^a depends on the repair mode a and the age z at failure. Thus, if $L_{z,a}$ denotes the postrepair life after instantaneous repair using repair mode a on failure at age z , the effect of repair is defined by the postrepair survival probability,

$$P\{L_{z,a} > t\} := E\bar{F}_Y(t) = \int_0^\infty \bar{F}_y(t) A_z^a(dy) \quad (1)$$

where $\bar{F}_y$ is the tail of the distribution of the *residual life* at age y , given by $\bar{F}_y(x) = \bar{F}(x+y)/\bar{F}(y)$. Two repair modes a, a' such that $L_{z,a} \stackrel{d}{=} L_{z,a'}$, all $z \geq 0$, are equivalent ($a \sim a'$), as their effects on equipment degradation cannot be distinguished.

The *virtual age* idea was also considered by Baxter *et al.* [28] to study the interfailure time distribution

$$B(t) = \int_0^\infty F_Y(t) P(dy) = E_P F_Y(t) \quad (2)$$

with P governing the distribution of the virtual age after repairs. Bhattacharjee's [27] formulation in equation (1) corresponds to parameterizing the virtual age distribution as $P \equiv A_z^a$, leading to the distribution of postrepair life. If $P(Y = z) = A_z^a\{z\} < 1$, we may interpret the value of effective postrepair virtual age Y of the unit which failed at age z , as the difference between the degree of effort needed for a PR *versus* the actual repair effort measured in units of operating time.

Starting with a unit (of age $x \geq 0$), which then fails at age $z(\geq x)$, it follows that the distribution of the postrepair survival times under PR, MR, BP, and BBS repair modes are given by,

$$P\{L_{z,a} > t\} = \begin{cases} \bar{F}_x(t), & a := \text{PR} \\ \bar{F}_z(t), & a := \text{MR} \\ p\bar{F}_x(t) + q\bar{F}_z(t), & a := \text{BP} \\ p(z)\bar{F}_x(t) + q(z)\bar{F}_z(t), & a := \text{BBS} \end{cases} \quad (3)$$

where F_x denotes the *residual life* distribution at age x , with $F_0 \equiv F$. When the MR option is chosen in the BP and BBS models, note that the additional age accumulated at the most recent failure is $z - x \geq$

0; setting $S := F_x$, it follows that the postrepair life distribution, conditional on the choice of MR, is

$$\begin{aligned} \bar{S}_{z-x}(t) &= \frac{\bar{S}(z-x+t)}{\bar{S}(z-x)} \equiv \frac{\bar{F}_x(z-x+t)}{\bar{F}_x(z-x)} \\ &= \frac{\bar{F}(z+t)}{\bar{F}(z)} = \bar{F}_z(t) \end{aligned} \quad (4)$$

justifying equation (3). Note also that, if the currently failed unit began working as a new equipment ($x = 0$), then a PR in equation (3) corresponds to a replacement. As a general formulation of the notion of *IR*, equation (1) is sufficiently flexible to allow different variations by suitably choosing the virtual age distribution A_z^a . Such choices can include nonparametric constraints on the postrepair life distribution, such as specifying (a) *maximum allowable improvement in virtual age*: $0 \leq A_z\{0\} < 1$ and/or, $A_z(0, \epsilon) = 0$, for some $\epsilon > 0$, (b) *a limit on maximum possible degradation*: $A_z(z+L, \infty) = 0$, $L > 0$; i.e., virtual age after repair is at most an additional L over the age at failure, or (c) *specify a best/worst case scenario for the mean virtual age*: the expected virtual age after repair is no more than (at least, respectively), the age at failure $\int_0^\infty y A_z(dy) \leq (\geq) z$.

A Dynamic Programming Framework for Repairable Systems

The DP model described below is based on Bhattacharjee [27], drawing on the foundations of a theory of stochastic DP developed by Blackwell [2, 3] and Strauch [4]. The *states* s of the system are defined by a pair $s = (t, x)$, where

t = remaining clock-time, for a given time horizon,

x = *virtual age* of the unit, which starts functioning at t

The set $S = \{(t, x) : x \geq 0\}$ is the state space and the states of the system correspond to those time epochs, when a unit starts working, possibly after a repair. The subset $S^* = \{(t, x) : x \geq 0, t \leq 0\}$ is the set of *trap states* in the sense that the process ends on entering S^* . (For an infinite time horizon, it is simpler to define the system's state s as the virtual age x of the unit at those time epochs when the unit starts working.) A repair *action* is a repair mode

that one can implement with its effect on postrepair life described by equation (1). The set of available choices of repair actions a is $A = \{a : a \in A\}$.

Let $T_x \stackrel{d}{=} X - x \mid X > x \sim F_x$ denote the residual life at age x ($T_0 \stackrel{d}{=} X \sim F$). Note also that if we start with an unit of age x (i.e., at a state $(\cdot, x)$), its time to next failure is T_x with distribution function (d.f.) F_x , and *not* F . When we choose a repair action $a \in A$ at a state $s = (t, x) \notin S^*$, our system moves to a new state s' , possibly depending on $(s \equiv (t, x), a)$, according to the transition scheme:

$$s = (t, x) \xrightarrow{a} s' = (t - T_x, V_{x,a}) \quad (5)$$

where the repaired unit's virtual age $V_{x,a}$ depends on its age x when it last began functioning and the repair mode a executed when it subsequently failed. Thus, the conditional probability distribution $q(ds' \mid (s, a))$, governing the transition in equation (5) is determined by the joint distribution of $(T_x, V_{x,a})$. The particular form of the virtual age $V_{x,a}$ in turn depends on the chosen repair action a ; e.g.,

$$V_{x,a} = \begin{cases} 0, & \text{for } a = \text{Replacement} \\ x, & \text{for } a = \text{PR} \\ x + T_x, & \text{for } a = \text{MR} \\ x + aT_x, & \text{Model I (Kijima), } 0 \leq a \leq 1 \\ a(x + T_x), & \text{Model II (Kijima), } 0 \leq a \leq 1 \end{cases} \quad (6)$$

where $a \in A = [0, 1]$ reflects the “degree of repair” in Kijima’s models [20]. Here, increasing values of a represent progressively worse quality of repair. When $a = 1$, both models of Kijima reduce to MR. By contrast, $a = 0$ represents a PR in Model I, but a replacement in Model II.

A *repair strategy* is a sequence $\pi = \{\pi_1, \pi_2, \dots\}$ that spells out a sequential plan of repair actions to choose after each failure, possibly based on the past history. Formally, the choice of repair after the n th failure is governed by π_n , a conditional distribution on the available set A of repair actions, given the past history $(s_1, a_1 \dots s_{n-1}, a_{n-1}, s_n)$ of failures and repairs. Note that, based on our formulation of the state space, a repair strategy need only specify each π_n when the most recent state $s_n \notin S^*$.

Among the simpler strategies, which are intuitively appealing, are Markovian and/or stationary ones. A strategy $\pi = \{\pi_n, n \geq 1\}$ is Markovian if each π_n is a regular conditional distribution on the

set of repair actions A , given the state at the most recent failure. If each π_n is degenerate, i.e., if $\pi_n \equiv f_n : S \mapsto A$; then $\pi = \{f_n, n \geq 1\}$ is a *nonrandomized Markov* repair strategy. If $f_n \equiv f : S \mapsto A$ is time-homogeneous, then $\pi := f^\infty \equiv \{f, f, \dots\}$ is a *nonrandomized stationary* strategy of repairs.

To illustrate some of the IR repair modes and the corresponding repair schemes considered in the literature, described in the section titled “Early Work, Minimal Repairs, and Other Imperfect Repair Schemes”, as repair strategies in a DP framework; let 0 and 1 denote a PR and an MR mode, respectively. If $A = \{1\}$, then the only available strategy is “forever minimal repairs” (FMR): $\pi = 1^\infty$. Similarly, if $\{0, 1\} \subset A$, then available choices include the age-dependent BBS repairs, as the randomized stationary strategy $\pi := \pi_1^\infty$, where

$$\begin{aligned} \pi_1(\{0\} \mid (t, x)) &= p(T_x), \\ \pi_1(\{1\} \mid (t, x)) &= 1 - p(T_x); \quad (t, x) \notin S^* \end{aligned} \quad (7)$$

noting that the age of the currently failed unit, which started working as a possibly old unit of age x , is $T_x \stackrel{d}{=} X - x \mid X > x$. The Brown–Proschan (BP) repair scheme corresponds to the special case of the BBS strategy when $p(\cdot) \equiv p$ with $0 < p \leq 1$. The case $p \equiv 1$ is the strategy of “forever perfect repairs” (FPR), which is equivalent to a strategy of replacements when the process starts with a new equipment. This equivalence is clear by noting that while replacement by a new copy of the failed equipment is technically an option that is distinct from PR; setting $x = 0$ in the stochastic transition for a PR option:

$$s = (t, x) \xrightarrow{\text{PR}} s' = (t - T_x, x) \quad (8)$$

the transition corresponding to $x = 0$ (new equipment) is $s = (0, x) \xrightarrow{\text{PR}} s' = (t - X, 0)$, and thus equivalent to a replacement.

Similarly, for various “degree of repair” models considered by Kijima [20] and others; $A = [0, 1]$ parameterizes the degree of repair efforts, as described in equation (6); any repair strategy is a sequence of distributions on the unit interval, given the past.

For a given initial state $s = (t, x)$, a repair strategy π determines a probability P_π on $H := ASAS\dots$, the set of future histories of the system. Rewards between transitions are defined by a

bounded real-valued function r on SAS such that $r(s, a, s')$ is the reward for choosing repair action a , at state s , used to repair the equipment when it fails. The choice of r is determined by our objective. For example, to compute the expected number of failures (repairs) in a time interval $(0, t)$, take $r(s, a, s') = 1$, for all a , if s and $s' \notin S^*$, and $= 0$ otherwise. A history $h = (s \equiv s_1, a_1, s_2, a_2, \dots)$, starting at s , generates a sequential stream of random rewards $r_n := r(s_n, a_n, s_{n+1})$ resulting in a total reward $R := \sum_{n=1}^{\infty} \beta^n r_n$ defined on SH , with a discount factor $\beta \in (0, 1]$ (in the undiscounted case $\beta = 1$). The *return* (total expected reward) of a repair strategy π is

$$I_\pi(s) := E_{P_\pi} R = \int_H R(s, h) P_\pi(dh), \quad s \in S \quad (9)$$

I_π is finite under discounting, or other suitable conditions on r . The *optimal return* is $I^*(s) := \sup_\pi I_\pi(s)$, $s \in S$. An optimal repair strategy π^* , if it exists, is one for which $I_{\pi^*} = I^*$, all $s \in S$.

Illustrative Applications

We give two examples of how the preceding ideas can be employed to investigate repair and replacement problems. It should be emphasized that even if our interest is confined to the analysis of performance of a specific repair strategy, such as the Brown–Proschan (BP) repair scheme, and not in pursuing possible optimality, ideas from DP can often allow us an elegant way to analyze specific repair schemes.

Computing the Sojourn Time to the Next Perfect Repair

In the Brown–Proschan (BP) repair scheme, which defines a stationary randomized repair strategy with constant probabilities p and $(1 - p)$, respectively, of choosing a PR or an MR repair option; suppose that there is a *terminal payoff* $= 1$ on entering the trap states S^* , and that all other payoffs are zero. Beginning with a unit of age x , let T_x^* be the *sojourn time to the next* PR. If $X_1, X_2, \dots$ are the postrepair lives under successive MRs, and the next PR occurs at the τ -th failure ($\tau < \infty$, w.p. 1, if $p > 0$), then the *expected total return* of the BP-repair strategy is clearly,

$$E_{BP} \{1_{\{t - X_1 - X_2 - \dots - X_\tau < 0\}}\} = P_{BP}(T_x^* > t) \quad (10)$$

the tail of the sojourn time distribution to the next PR. To compute the sojourn time tail (equation 10) by using DP ideas, consider an operator L mapping nonnegative bounded functions into itself, such that for $(t, x) \notin S^*$,

$$\begin{aligned} Lu(t, x) &:= pE1_{T_x > t} + qE\{1_{T_x > t} + 1_{T_x \leq t} u(t - T_x, x + T_x)\} \\ &= \bar{F}_x(t) + q \int_0^t u(t - y, x + y) F_x(dy) \end{aligned} \quad (11)$$

and $Lu(t, x) = 1$, if $(t, x) \in S^*$. Then $Lu(t, x)$ is our expected reward if we use the BP repair mode at the first failure after (t, x) and then stop with our terminal payoff $u(\cdot, \cdot)$ at next state. Let $\mathbf{0}$ denote the function which identically vanishes on S . Then, it can be shown that,

Theorem 1 Bhattacharjee [27]

1. L is a contraction.
2. $L^n \mathbf{0}(t, x) \rightarrow P_{BP}(T_x^* > t)$, as $n \rightarrow \infty$.
3. $u^*(t, x) := 1_{(t, x) \in S^*} + 1_{(t, x) \notin S^*} [\bar{F}_x(t)]^p$, is the unique fixed point of L .
4. $P_{BP}(T_x^* > t) = [\bar{F}_x(t)]^p$, for $(t, x) \notin S^*$.

Brown and Proschan's [17] result on the sojourn time distribution follows as a special case of claim 4, with $x = 0$.

Bounding the Renewal Function via Dynamic Programming

Blackwell [29] first pointed out the utility of exploiting DP ideas to derive sharp probability bounds. Renewal functions, which represent the expected number of failures (repairs) under a replacements-only strategy, can be explored using such ideas, to provide sharp bounds on renewal functions under various nonparametric constraints on the failure distribution generating the renewal process.

The renewal function can be viewed as the total expected payoff, in a game against a passive nature, in which the states are points $t \geq 0$ on the half-line, and stochastic transitions are defined by $t \longrightarrow \max(t - X, 0)$, where $X \geq 0$ with df F . The game stops on hitting the origin. Corresponding to this transition, we earn a "reward" $1_{\{X < t\}}$. Corresponding to successive moves dictated by an independent and

identically distributed sequence $X_1, X_2, \dots$ with df F , the total expected payoff is clearly,

$$\begin{aligned} E \left(\sum_{n=1}^{\infty} 1_{\{\sum_{j=1}^n X_j \leq t\}} \right) \\ = E \max \left\{ n : \sum_{j=1}^n X_j \leq t \right\} := M(t) \end{aligned} \quad (12)$$

the renewal function. In the spirit of the section titled “A Dynamic Programming Framework for Repairable Systems”, and the work of Blackwell [29], a method to construct sharp upper bounds to the renewal function was developed by Bhattacharjee [30], by considering an operator U mapping the set of real-valued, nonnegative nondecreasing functions Q , such that $Q(0) = 0$, $Q(\infty) = \infty$, into itself, defined by

$$UQ(t) := F(t) + \int_0^t Q(t-x)F(dx), \quad t \geq 0 \quad (13)$$

to show that,

1. $U^n \mathbf{0}(t) \rightarrow M(t)$ as $n \rightarrow \infty$,
2. U is monotone, and $UQ \leq Q$ implies $M \leq Q$ pointwise,
3. $M(t)$ is the minimal fixed point of U .

The well-known upper bound $M(t) \leq (t/\mu)$ for NBUE (new better than used in expectation) life distributions with a mean μ can be derived as an application of this result, by choosing $Q(t) := t/\mu$. When F has the stronger NBU (new better than used) property, the same method is used by Bhattacharjee [30] to show that

$$F \text{ is NBU implies } M(t) \leq -\ln\{1 - F(t)\}, \quad t \geq 0$$

i.e., the “hazard function” is an improved sharp upper bound for the renewal function under the NBU hypothesis.

One can thus explore the performance of repair strategies for appropriate payoff functions corresponding to state transitions that capture a desired objective function, such as the total expected number of failures (repairs) in a time interval $(0, t)$, starting with a unit of age $x \geq 0$. The corresponding technical issues that must be addressed in such analysis often reduce to finding the minimal solution or, tight bounds to the solution, of an integral equation. Such

an approach can also be used to explore optimal or near-optimal repair strategies. When we superimpose suitable costs (of repair) and rewards between transitions (incomes generated during equipment functioning cycles), and possibly an infinite-time horizon with or without discounting, we then have a corresponding economic optimization problem, for which DP methods are eminently suitable.

References

- [1] Bellman, R.S. (1957). *Dynamic Programming*, Princeton University Press, Princeton.
- [2] Blackwell, D. (1965). Discounted dynamic programming, *Annals of Mathematical Statistics* **36**, 226–235.
- [3] Blackwell, D. (1966). Positive dynamic programming, *Proceedings of the Fifth Berkeley Symposium on Mathematical Statistics and Probability*, University of California Press, Vol. I, pp. 415–418.
- [4] Strauch, R.E. (1966). Negative dynamic programming, *Annals of Mathematical Statistics* **37**, 871–890.
- [5] Miranda, M.J. & Feckler, P.L. (2004). *Applied Computational Economics*, MIT Press.
- [6] Bellman, R. & Dryfus, S. (1958). Dynamic programming and the reliability of multicomponent devices, *Operations Research* **6**(2), 200–206.
- [7] Kettelle, J.D. Jr. (1962). Least-cost allocation of reliability investment, *Operations Research* **10**(2), 249–265.
- [8] Jorgenson, D. & Radner, R. (1962). Optimal replacement and inspection of stochastically failing equipment, in *Studies in Applied Probability and Management Science*, K.J. Arrow, S. Karlin & H. Scarf, Stanford University Press, Stanford, Chapter 12.
- [9] Radner, R. & Jorgenson, D. (1963). Opportunistic replacement of a single part in the presence of several monitored parts, *Management Science* **10**(1), 70–84.
- [10] Chen, M.C. & Feldman, R.M. (1997). Optimal replacement policies with minimal repair and age-dependent costs, *European Journal of Operational Research* **98**, 75–84.
- [11] Lam, Y. (1988). A note on the optimal replacement problem, *Advances in Applied Probability* **20**(2), 479–482.
- [12] Lam, Y. (1991). An optimal repairable replacement model for deteriorating systems, *Journal of Applied Probability* **28**(4), 843–851.
- [13] Stadge, W. & Zuckerman, D. (1990). Optimal strategies for some repair replacement models, *Advances in Applied Probability* **22**(3), 641–656.
- [14] Boshuizen, F.A. (1997). Optimal replacement strategies for repairable systems, *Mathematical Methods of Operations Research* **45**(1), 133–144.
- [15] Barlow, R.E. & Proschan, F. (1962). Planned replacement, in *Studies in Applied Probability and Management Science*, K.J. Arrow, S. Karlin & H. Scarf, eds, Stanford University Press, Stanford, Chapter 4.

- [16] Blackwell, D. (1962). Discrete Dynamic Programming, *Annals of Mathematical Statistics* **33**(2), 719–726.
- [17] Brown, M. & Proschan, F. (1983). Imperfect repair, *Journal of Applied Probability* **20**(4), 851–859.
- [18] Block, H., Borges, W.S. & Savits, T.H. (1985). Age dependent minimal repair, *Journal of Applied Probability* **22**(2), 370–385.
- [19] Bhattacharjee, M.C. (1987). New results for the Brown-Proschan model of imperfect repair, *Journal of Statistical Planning and Inference* **16**(3), 305–316.
- [20] Kijima, M. (1989). Some results for repairable systems with general repair, *Journal of Applied Probability* **26**(1), 89–102.
- [21] Finkelstein, M.S. (1993). On some models of general repair, *Microelectronics and Reliability* **33**(5), 636–666.
- [22] Makis, V. & Jardine, A.K.S. (1993). A note on optimal replacement policy under general repair, *European Journal Operations Research* **69**, 75–82.
- [23] Dagpunar, J.S. (1998). Some properties and computational results for a general repair process, *Naval Research Logistics* **45**(4), 391–405.
- [24] Stadje, W. & Zuckerman, D. (1991). Optimal repair policies with general degree of repair in two maintenance models, *Operations Research Letters* **11**(2), 77–80.
- [25] Stadje, W. & Zuckerman, D. (1991). Optimal maintenance strategies for repairable systems with general degree of repair, *Journal of Applied Probability* **28**(2), 384–396.
- [26] Guo, R., Ascher, H. & Love, E. (2001). Towards practical and synthetical modelling of repairable systems, *Economic Quality Control* **16**(2), 147–182.
- [27] Bhattacharjee, M.C. (2000). A unified framework for modeling and analysis of repairable systems, in *Perspectives in Statistical Sciences*, A.K. Basu, J.K. Ghosh, P.K. Sen & B.K. Sinha, eds, Oxford University Press, pp. 58–74.
- [28] Baxter, L., Kijima, M. & Tortorella, M. (1996). A point process model for the reliability of a maintained system subject to general repair, *Communications in Statistics: Stochastic Models* **12**(1), 37–65.
- [29] Blackwell, D. (1964). Probability bounds via dynamic programming, *Proceedings of the Symposia in Applied Mathematics*, American Mathematical Society, Providence, Vol. XVI, pp. 277–280.
- [30] Bhattacharjee, M.C. (1999). Dynamic programming, renewal functions and perfect vs. Minimal repair comparisons, in *Frontiers in Reliability Analysis*, A.P. Basu, ed, World Scientific Press, Singapore, pp. 63–70.
- Aven, T. (2007). General minimal repair models, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons.
- Bergman, B. (2007). Stationary replacement strategies, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons.
- Ebrahimi, N. (2007). Multivariate age and multivariate renewal replacement, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons.
- Ebrahimi, N. (2007). Multivariate imperfect repair models, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons.
- Popova, E. (2007). Replacement strategies, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons.
- Rigdon, S. (2007). Imperfect repair, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons.
- Samaniego, F. & Hollander, M. (2007). Imperfect repair, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons.
- Shiu, S.H. (2007). Age-dependent minimal repair and maintenance, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons.

Related Articles

We provide a shortlist of articles, in this volume, which cover the themes of system replacement and repair from other perspectives that may be useful for a comprehensive study of these aspects of reliability theory. For a review of replacement strategies, see the articles by Bergman (**Stationary Replacement Strategies**) and Popova (**Replacement Strategies**). An overview of repairable systems is provided by Ascher (**Repairable Systems Reliability**). Minimal repair models are covered in the articles by Aven (**General Minimal Repair Models**), and Shiu (**Age-Dependent Minimal Repair and Maintenance**); while Rigdon (**Imperfect Repair, Counting Processes**) and Samaniego and Hollander (**Imperfect Repair**) discuss the broader theme of imperfect repairs. For a survey of multivariate approaches to replacement and imperfect repair, see Ebrahimi (**Multivariate Age and Multivariate Renewal Replacement**) and Ebrahimi (**Multivariate Imperfect Repair Models**).

MANISH C. BHATTACHARJEE

Further Reading

Ascher, H. (2007). Repairable systems reliability, *Encyclopedia of Statistics in Quality and Reliability*, John Wiley & Sons.

Splines and other Metamodels in Reliability Analysis

Introduction

In structural reliability, the limit state function (LSF) defines some measure of performance of a system, by means of a reliability index or probability of failure, to a vector of observed properties of the system. This function might be of very complex form, leading to the use of various approximations to both evaluate it and to fit it to data. In this contribution, the metamodel approach is looked at, in which the use of a relatively simple function to approximate the state function. Four possible examples are looked at and compared: response surfaces (RSs), splines (SPs), **neural networks** (NN), and ordinary kriging (OK). In addition, an optimization procedure for the calculation of the reliability index is presented.

Expressing the LSF in a closed form is often not possible in the reliability analysis of complex systems. In principle, any algorithm used for structural reliability evaluation may be adapted to handle implicit LSFs. However, when the gradients of the LSF are required, which is the case for first-order reliability methods (FORM) and second-order reliability methods (SORM), the performance is affected when direct or analytical differentiation are not available [1–3]. In the latter, system behavior is not included either.

The use of an RS in combination with a simulation method to increase the efficiency is widely spread [3–5]. The main idea is that the response consisting of a complex function of input variables is approximated by a simple function of the input variables. If the RS is capable of handling the system behavior, the reliability analysis can be performed on the RS, instead of using the original problem [2]. Complex system behavior might however not be captured accurately by a low-order polynomial. Ideally, no functional form is preset and a “universal estimator” is used. Therefore, other extensions are proposed, in general called *metamodels* [6, 7].

Techniques employing simulation procedures such as **Monte Carlo** and directional sampling (DS) inherently account for the system behavior. The disadvantage of these techniques is the number of samples – and thus calls to the implicit LSF – to come up with a sufficiently accurate value of the system failure probability. Importance sampling reduces the number of samples to a certain extent, but the overall performance remains far below the performance of FORM or SORM methods.

In the next section possible metamodels are reviewed. In the section titled “Reliability Analysis Framework”, the accompanying reliability analysis framework is presented. In the section titled “Example: A Series System with Four Branches”, an example application is worked out. The conclusions of this contribution are covered in the final section.

Possible Metamodels

Response Surface (RS) – Low-Order Polynomial

In its most general form, the RS consists of a constant, linear, mixed, and quadratic terms [1–5]:

$$g_{MM}(\mathbf{x}) = a_0 + \sum_{i=1}^n a_i x_i + \sum_{i=1}^n \sum_{j \geq i}^n a_{ij} x_i x_j \quad (1)$$

where $g_{MM}(\mathbf{x})$ is the metamodel of the LSF $g(\mathbf{x})$ in which $\mathbf{x}$ is a vector with size $1 \times n$; n the number of random variables defining the problem at hand; and a_0, a_i, a_{ij} the least-square regression coefficients.

Splines (SP)

The piecewise continuous polynomials are based on the thin-plate smoothing SP concept, that is able to deal with multivariate input and a single output – outcome of the LSF [7, 8]. Because of its simple form, evaluation afterward is straightforward. The LSF can be written as follows:

$$g_{MM}(\mathbf{x}) = \sum_{j=1}^{n_{LSFE}} a_j \psi(\mathbf{x} - \mathbf{x}_{c_j}) \quad (2)$$

where a_j : the smoothing SP coefficients ($1 \times n_{LSFE}$); $\mathbf{x}_{c_j}$: the coordinates at which the limit state function evaluation (LSFE) is available ($1 \times n_{LSFE}$); and $\psi(\mathbf{x})$: the thin-plate SP basis function, that equals

$$\psi(\mathbf{x}) = \|\mathbf{x}\|^2 \log(\|\mathbf{x}\|^2) \quad (3)$$

2 Splines and other Metamodels in Reliability Analysis

This thin-plate smoothing SP $g_{\text{MM}}(\mathbf{x})$ is the unique minimizer of the weighted sum

$$sE(g_{\text{MM}}(\mathbf{x})) + (1-s)R(g_{\text{MM}}(\mathbf{x})) \quad (4)$$

with $E(g_{\text{MM}}(\mathbf{x}))$ an error measure and $R(g_{\text{MM}}(\mathbf{x}))$ a roughness measure.

Neural Networks

NN reply to the criteria of a “universal estimator” and therefore offer an interesting perspective [7, 9]. The weighting functions that are the basis of the network are very flexible and adapt to virtually any kind of functional behavior. Therefore, there are no limits as to shape, dimension, singularity, or type of function.

A single neuron k is not capable of describing an arbitrarily relationship in an accurate way. The output of neuron k as a function of the input signals x_i is described as follows:

$$y_k = \varphi \left(\sum_{i=1}^{n_{\text{LSFE}}} w_{k,i} x_i + \theta_k \right) \quad (5)$$

where $x_i (i = 1, \dots, n_{\text{LSFE}})$ are the input signals; $w_{k,i}$: weights; θ_k : bias term; and φ : activating function (often sigmoidal). Multilayer networks are quite powerful. A network of two layers, where the first layer is sigmoid and the second layer is linear, can be trained to approximate any function (with a finite number of discontinuities) arbitrarily well. Different learning rules exist as a procedure for modifying the weights ($w_{k,i}$) and biases (θ_k) of a network [9]. These algorithms use the gradient of the performance function – direction of the steepest descent – to determine how to adjust the weights to minimize the average error between the network outputs and the target outputs. To further increase the convergence speed, quasi-Newton learning rules are applied, such as the Levenberg–Marquardt learning rule. The Hessian – a matrix of second-order derivatives of the performance function – is calculated numerically based on the gradient.

Ordinary Kriging

Preferably an error measure should be linked to the availability of information in the surrounding region. When more information is available in the sample region, the quality of the metamodel will probably be better, compared to other directions or points in which the availability of information is limited. OK provides

an answer to that requirement. OK is widely used to estimate a value at a point ($\mathbf{x}_0$) of a region for which a variogram is known, using data in the neighborhood ($\mathbf{x}_i; i = 1, \dots, n$) of the estimation location [7, 10]:

$$g_{\text{MM}}(\mathbf{x}_0) = \sum_{i=1}^{n_{\text{LSFE}}} w_i g_{\text{LSF}}(\mathbf{x}_i) \quad (6)$$

where $g_{\text{LSF}}(\mathbf{x}_i)$ is the known outcome of the LSF in point $\mathbf{x}_i$; $g_{\text{MM}}(\mathbf{x}_0)$ the estimated outcome in point $\mathbf{x}_0$; and w_i are the OK unit sum weights warranting the unbiasedness. Kriging is the best linear unbiased estimator. The estimation variance equals:

$$\begin{aligned} \sigma_{\varepsilon_0}^2 &= E[(g_{\text{MM}}(\mathbf{x}_0) - g_{\text{LSF}}(\mathbf{x}_0))^2] \\ &= - \sum_{i=1}^{n_{\text{LSFE}}} \sum_{j=1}^{n_{\text{LSFE}}} w_i w_j \gamma(\mathbf{x}_j - \mathbf{x}_i) + 2 \sum_{i=1}^{n_{\text{LSFE}}} w_i \gamma(\mathbf{x}_0 - \mathbf{x}_i) \end{aligned} \quad (7)$$

where $\gamma(\mathbf{h})$ is the semivariogram [10]. To minimize the estimation variance with the constraint of the weights, a Lagrange multiplier (λ) is introduced. The minimized mean-squared prediction error or kriging (or prediction) variance can be rewritten into:

$$\sigma_{\varepsilon_0}^2 = \lambda + \sum_{i=1}^{n_{\text{LSFE}}} w_i \gamma(\mathbf{x}_0 - \mathbf{x}_i) - \gamma(\mathbf{x}_0 - \mathbf{x}_0) \quad (8)$$

OK also is an exact interpolator. In the case in which $\mathbf{x}_0$ corresponds to one of the data points $\mathbf{x}_i$ in which the outcome is known, the estimated outcome will be identical to the outcome of the LSF. In all other cases a direct measure of the prediction error is obtained. The semivariogram $\gamma(\mathbf{h})$ represents the spatial correlation between the available data points. The success of the kriging procedure very much depends on the knowledge of the semivariogram. Furthermore, to be able to solve the set of equations, the semivariogram has to be positive definite.

Reliability Analysis Framework

The goal of this section is to calculate the reliability index or failure probability in an optimal way, in combination with the building of the metamodel. Therefore, this is performed, parallel with the construction of the metamodel. Therefore, in Step 1, a limited number of experiments are performed, for

which the outcome is calculated based on the real LSF. In Step 2, the initial metamodel is estimated. In Step 3, which is an iterative step, the procedure is repeated until a sufficiently accurate result is obtained.

Building a metamodel and performing the reliability analysis on the metamodel instead of the real LSF, might be a shift of effort. The number of LSFs required to build an accurate metamodel might be equally large. Therefore, an adaptive scheme is applied, that has three main steps that are described below ([11, 12], Figure 1):

- **Step 1**

The value of each random variable is increased individually until the root ($\lambda_{LSF,i}$) of the limit state

is found in the standard normal space, u -space (uncorrelated multivariate space, in which all random variables have zero mean and standard deviation equal to one). First a linear estimate is calculated, based on the outcome in the origin of the u -space (standard normal space) $(0, 0, 0, \dots, 0)$ and the point $(0, 0, \dots, \pm \beta_0, \dots, 0)$. When known, the start value β_0 is set equal to the expected reliability index β . Experience shows that $\beta_0 = 3$ performs well too. Further approximations are based on a piecewise continuous polynomial fit of second-order through the outcome of the different iteration points. Generally after three to four iterations the root is found. As a result an initial set of data is available linking the input variables with the corresponding output using the initial LSF. This initial dataset is used in Step 2.

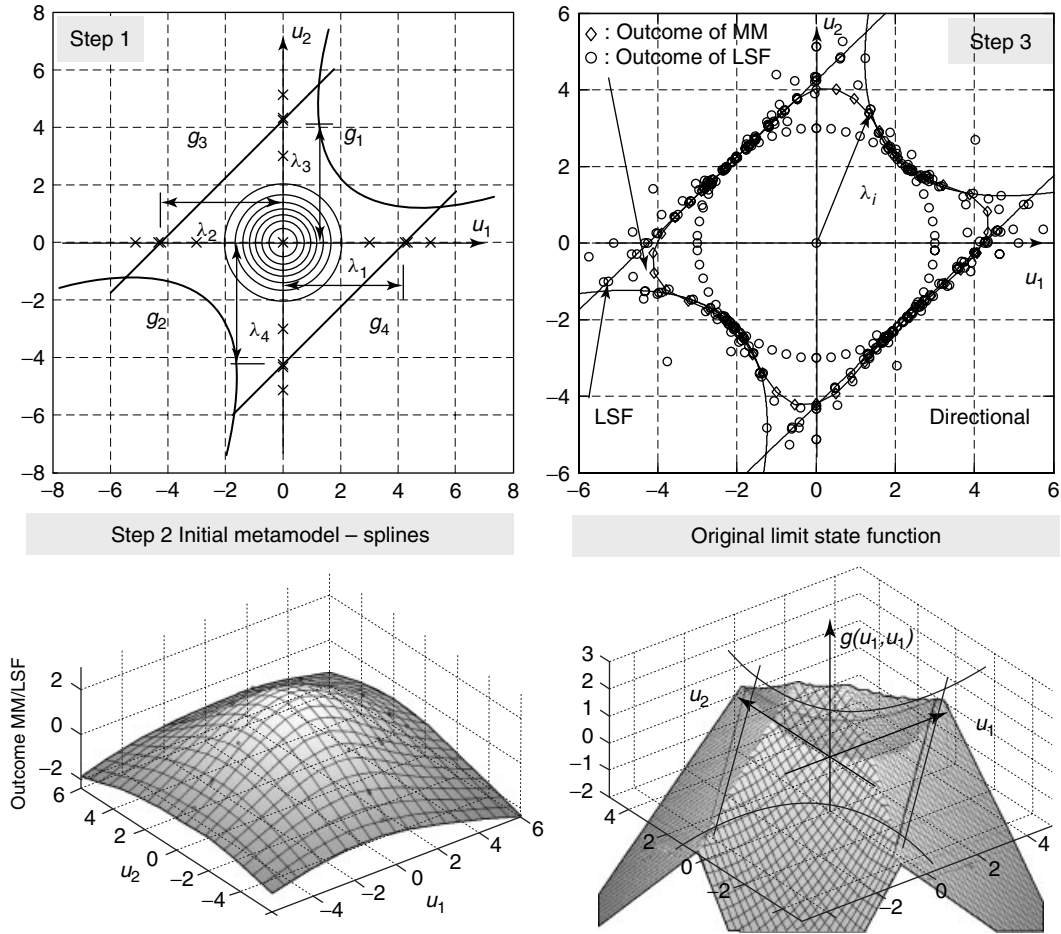


Figure 1 Different steps in the reliability framework

• Step 2

An initial metamodel (RS, SP, NN, OK) is fit through the data. Depending on the type of metamodel, different techniques are used to fit the metamodels: **least squares** in the case of RS; minimizing a weighted function with an error measure and roughness measure in the case of SP [8]; through the estimated semivariogram in the case of OK minimizing the estimation variance introducing a Lagrange multiplier; and one of the several learning rules available in case of NN, such as the quasi-Newton learning rule called *Levenberg–Marquardt*.

• Step 3

The metamodel is adapted and the reliability index is updated in an iterative procedure until the required accuracy is reached. Therefore, coarse DS or importance sampling Monte Carlo (ISMC) is performed on the metamodel. If the sample or its outcome is judged to be outside the important region, then the outcome based on the metamodel is used. If, on the other hand, the sample or its outcome is in the important region, a new LSFE is performed. The metamodel is updated as soon as new data are available.

Directional Sampling – Distinction between Important and Unimportant Regions

In the iterative procedure, Step 3, crude DS is performed on the metamodel. It is proved that the expected value $E(\hat{p}_f)$ of all contributions ($\hat{p}_i$) is an

unbiased estimate of the global failure probability (p_f) [13]:

$$p_f = E(\hat{p}_f) = \frac{1}{N} \sum_{i=1}^N \hat{p}_i \text{ with } \hat{p}_i = \chi_n^2(\lambda_{MM,i} \text{ or } \lambda_{LSF,i}) \quad (9)$$

For each sample, a first estimate of the distance $\lambda_{MM,i}$ to the origin in the standard normal space, u -space, is made based on the metamodel. An additional distance λ_{add} is used to make distinction between important and less important directions: $\lambda_{MM,i} < \lambda_{min} + \lambda_{add}$, in which λ_{min} is the minimum distance found so far, Figure 2(a). In case of a low-order RS, SP, or NN, a global error measure is used. The additional distance (λ_{add}) is based on the difference between the root of the real LSFE ($\lambda_{LSF,i}$) and the root of the metamodel ($\lambda_{MM,i}$) in the sample direction. A 99% **confidence interval** is taken as measure for λ_{add} :

$$\lambda_{add} = \max\{\text{abs}(\mu(\varepsilon_{rt}) \pm t_{0.01, n_{rt}-1} \times \sigma(\varepsilon_{rt}))\} \text{ with} \\ \varepsilon_{rt,i} = \lambda_{LSF,i} - \lambda_{MM,i} \quad (10)$$

When a good agreement is found between the roots of the real LSF and the metamodel, a low value of the additional distance (λ_{add}) is obtained; in other cases the distance will be larger.

In the case of OK, a local error measure is used. The kriging variance can be calculated for the point $\mathbf{x}_0$ in which the outcome is estimated, equation (11).

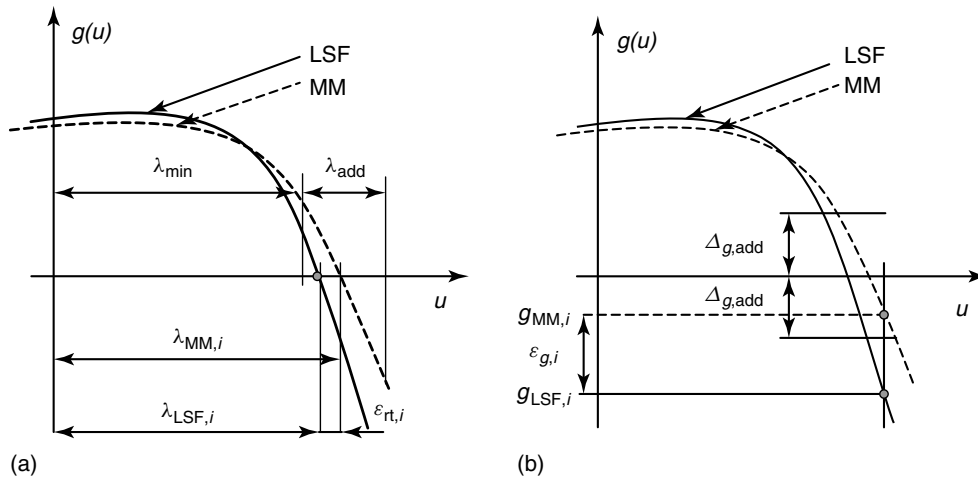


Figure 2 Additional distance (a) in the input space (λ_{add}) and (b) in the outcome space ($\Delta_{g,add}$)

Again a 99% confidence interval is taken as measure for λ_{add} :

$$\lambda_{\text{add}} = t_{0.01, n_{\text{LSFE}}-1} \times \sigma_{\varepsilon_i} \quad (11)$$

When data points are available in the neighborhood, the prediction error (σ_{ε}) and additional distance (λ_{add}) will be small. When the region is lacking information, the prediction error (σ_{ε}) and thus additional distance (λ_{add}) will be large. The higher the additional distance, the more LSFEs are performed. In case the obtained root ($\lambda_{\text{MM},i}$) has a relatively high contribution – $\lambda_{\text{MM},i} < \lambda_{\text{min}} + \lambda_{\text{add}}$ – to the estimated global failure probability (p_f), the root of the real response is used instead of the RS: $\lambda_{\text{LSF},i}$. The metamodel and the additional distance (λ_{add}) are updated with these new data. In case the contribution is less important – $\lambda_{\text{MM},i} > \lambda_{\text{min}} + \lambda_{\text{add}}$ – the contribution based on the root of the metamodel ($\lambda_{\text{MM},i}$) is used, avoiding time-consuming LSFEs.

Importance Sampling Monte Carlo – Distinction between Important and Unimportant Regions

Step 3 in the analysis can be performed using ISMC on the metamodel as well [11, 12]. Again, the expected value $E(\hat{p}_f)$ of all contributions ($\hat{p}_i$) is an unbiased estimate of the global failure probability (p_f) [13]:

$$p_f = E(\hat{p}_f) = \frac{1}{N} \sum_{i=1}^N \left\{ I[g(\mathbf{v}_i) < 0] \frac{f_{\mathbf{x}}(\hat{\mathbf{v}}_i)}{h_{\mathbf{v}}(\hat{\mathbf{v}}_i)} \right\} \quad (12)$$

The efficiency is believed to be comparable. To distinguish the important from the nonimportant region, the magnitude of the outcome values is checked. When the values are small, $g_{\text{MM},i} < \Delta_{g,\text{add}}$, the sample is believed to be near the limit between safe and unsafe, Figure 2(b). In that case the outcome is recalculated on the basis of the LSF. This only requires a single LSFE. The data are used to update the metamodel, the sampling function (variance) and the additional distance in the outcome space ($\Delta_{g,\text{add}}$).

Again, in the case of a low-order RS, SP, or NN, a global error estimate is used. The additional distance in the outcome space is related to the error between the metamodel and the LSFE. A 99% confidence interval is preset:

$$\Delta_{g,\text{add}} = \max\{\text{abs}(\mu(\varepsilon_g) \pm t_{0.01, n_{\text{LSFE}}-1} \times \sigma(\varepsilon_g))\} \text{ with} \\ \varepsilon_{g,i} = g_{\text{LSF},i} - g_{\text{MM},i} \quad (13)$$

In the case of OK, a local error measure is used, based on the kriging variance that can be calculated for the point $\mathbf{x}_i$ in which the outcome is estimated, equation (11). Again a 99% confidence interval could be chosen as measure for the additional distance in the outcome space $\Delta_{g,\text{add}}$:

$$\Delta_{g,\text{add}} = t_{0.01, n_{\text{LSFE}}-1} \times \sigma_{\varepsilon_i} \quad (14)$$

Example: A Series System with Four Branches

The LSF, often treated as a benchmark example in literature [11, 12] reads as follows (see also Figure 1):

$$g = \min \begin{cases} 3.0 + 0.1(u_1 - u_2)^2 - (u_1 + u_2)\sqrt{2} \\ 3.0 + 0.1(u_1 - u_2)^2 + (u_1 + u_2)\sqrt{2} \\ (u_1 - u_2) + 3.5\sqrt{2} \\ (u_2 - u_1) + 3.5\sqrt{2} \end{cases} \quad (15)$$

Both random variables (u_1, u_2) are standard normal distributed random variables. The different branches have comparable contribution to the system failure probability. The results of the reliability analysis are summarized in Table 1. For comparison, the random sampling is started from the same seed in each analysis. The following information is given in the subsequent rows: the used simulation procedure (DS or ISMC), the metamodel used (none, RS, SP, NN, OK), the resulting system reliability index (β) and corresponding failure probability (p_f), the required number of sample directions or points (N), n_{LSFE} and the additional distance in the standard normal space or outcome space λ_{add} (for DS) or $\Delta_{g,\text{add}}$ (for ISMC) respectively at the end of the analysis.

The different techniques all result in a good estimate of the system reliability index. Some differences in efficiency are apparent:

- (Crude) DS offers an optimal solution method because of the nearly radial shape of the LSF at hand. This results in a lower number of samples (N) as well as the number of LSFE (n_{LSFE}) compared to procedures using Monte Carlo with importance sampling.
- The low-order polynomial RS is not capable of estimating the real LSF behavior accurately. This can be seen from the relatively high values for λ_{add} (DS) or $\Delta_{g,\text{add}}$ (ISMC), respectively. Of

Table 1 Series system with four branches: results of reliability analysis^(a)

Adaptive MM	Simulation method: directional sampling (DS)					Importance sampling Monte Carlo (ISMC)				
	none	RS ^(b)	SP ^(c)	NN ^(d)	OK ^(e)	none	RS	SP	NN	OK
$\beta(\equiv 2.85)$	2.79	3.03	3.00	3.05	2.95	2.84	2.84	2.81	2.76	2.88
p_f	0.0026	0.001	0.001	0.001	0.0015	0.0022	0.0020	0.0020	0.0029	0.0020
N	2750	500	375	800	420	4750	5402	3753	5651	3500
n_{LSFE}	9192	830	107	67	107	4750	3877	724	125	146
$\sigma(\beta)$	0.16	0.15	0.15	0.15	0.15	0.09	0.14	0.14	0.19	0.14
$\lambda_{\text{add}}/\Delta_{g,\text{add}}$	–	1.77	0.21	0.03	var	–	1.83	0.05	0.04	var

^(a) Overall preset accuracy $\sigma(\beta) = 0.15$.

^(b) RS: pure quadratic polynomial **response surface**.

^(c) SP: splines.

^(d) NN: feed-forward neural network ($2 - n_n - 1$: with n_n the number of neurons in the hidden layer – $n_n = n + n_{\text{LSFE}}/4$; training function: scaled conjugate gradient).

^(e) OK: ordinary kriging.

course, SPs have to work well for this particular example, because the LSF is a kind of piecewise continuous low-order polynomial.

- With respect to OK it is clear that the overall efficiency is similar to SPs and NNs. The number of LSFEs is similar. Both additional distances (λ_{add} and $\Delta_{g,\text{add}}$) did become variable distances (var). For each estimation point, the value is a function of the available information in the neighborhood. Figure 3 illustrates the prediction error in the standard normal space, dependent on the estimation point. Areas in which more information is available will have a small prediction error. In other regions, the prediction error will be higher. In case of RS, SP, or NN, the additional distances are a global error measure of the quality of the metamodel accounting for all

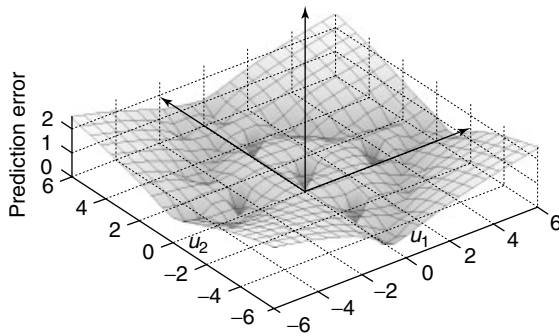
information available at that moment. A linear semivariogram is used to represent the spatial correlation. Although the number of data is insufficient to obtain an accurate fit of the real correlation, the applied semivariogram results in an acceptable efficiency.

Conclusions

This contribution describes a methodology to calculate the reliability of structural systems. The focus is on a methodology that results in an accurate value of the system reliability index, within a limited number of calls to the original LSF. It covers the problem that is often encountered in practice, where the LSF is only specified implicitly. On top of that, several failure modes contribute to the global failure probability.

Therefore, use is made of a procedure based on DS or Monte Carlo with importance sampling and an adaptive metamodel. Using an adaptive metamodel, the LSF is only evaluated if needed. Besides low-order polynomial RSs, the use of SPs, NNs, or OK is incorporated.

On the basis of a benchmark example, their mutual efficiency is discussed. From this example it is evidenced that RSs work well for problems that can be easily described by a low-order polynomial. In case the LSF is more complex, the use of SPs, NNs, and OK is demonstrated to be beneficial. The latter enables local error estimation.

**Figure 3** Prediction error – ordinary kriging (OK)

References

- [1] Gomes, H.M. & Awruch, A.M. (2004). Comparison of response surface and neural network with other methods for structural reliability analysis, *Structural Safety* **26**, 49–67.
- [2] Guan, X.L. & Melchers, R.E. (2001). Effect of response surface parameter variation on structural reliability estimates, *Structural Safety* **23**, 429–444.
- [3] Bucher, C.G., Chen, Y.M. & Schuëller, G.I. (1998). Time variant reliability analysis utilizing response surface approach, in *Proceedings of the 2nd IFIP WG 7.5 Conference on Reliability and Optimization of Structural Systems*.
- [4] Schuëller, G.I. (2002). Past, present and future of simulation based structural analysis, in *Proceedings of Euro SIBRAM Conference*, P. Marek, A. Haldar, M. Gustar & P. Tikalsky, eds, Institute of Theoretical and Applied Mechanics, Czech Academic of Sciences, Prague, pp. 17–27.
- [5] Gayton, N., Bourinet, J.M. & Lemaire, M. (2003). CQ2RS: a new statistical approach to the response surface method for reliability analysis, *Structural Safety* **25**, 99–121.
- [6] Schueremans, L. & Van Gemert, D. (2003). Sampling techniques with an adaptive meta-model to increase the efficiency in structural system reliability, in *Proceedings of the 11th IFIP WG7.5 Conference on Reliability and Optimization of Structural Systems*, Canada.
- [7] Schueremans, L. & Van Gemert, D. (2005). Benefit of splines and neural networks in simulation based structural reliability analysis, *Structural Safety* **27**(3), 246–261.
- [8] De Boor, C. (2001). A practical guide to splines, *Applied Mathematical Sciences*, Revised Edition, Springer-Verlag, New York, 1978, p. 392.
- [9] Hagan, M.T., Demuth, H.B. & Beale, M.H. (1996). *Neural Network Design*, PWS Publishing, Boston.
- [10] Wackernagel, H. (1995). *Multivariate Geostatistics: An Introduction with Applications*, Springer-Verlag, Berlin, Heidelberg.
- [11] Schueremans, L. (2001). *Probabilistic evaluation of structural unreinforced masonry*, Ph. D. dissertation, Katholieke Universiteit Leuven, Belgium.
- [12] Waarts, P.H. (2000). *Structural reliability using finite element methods: an appraisal of DARS: directional adaptive response surface sampling*, Ph. D. Thesis, Delft University of Technology, Delft.
- [13] Melchers, R.E. (1999). *Structural Reliability, Analysis and Prediction*, Ellis Horwood Series in Civil Engineering, John Wiley & Sons, Chichester.

Related Articles

Asymptotic Reliability Analysis of Very Large Structural Systems; Data Mining in Quality and Reliability; Data Mining, Evaluation Techniques in; Monte Carlo Methods, Univariate and Multivariate; Neural Networks in Statistical Process Control; Optimal Reliability Design-Modeling.

LUC SCHUEREMANS AND DIONYS VAN
GEMERT

Data Mining, Evaluation Techniques in

Introduction

We use data mining techniques to identify models that facilitate accurate descriptions and predictions of real life processes. We are also typically interested in gaining insights into the underlying process that generates the data. To give an example, banks using data mining techniques to identify fraudulent use of credit cards are not only interested in identifying fraud before authorizing the transaction, but are also interested in being able to communicate to the irate customer why the transaction was declined. Furthermore, they are interested in being able to understand fraud patterns in order to design prevention measures.

Data mining techniques, unlike hypothesis driven approaches, do not justify the model with theorems about the processes that produce the data – such as consumer behavior models or market equilibrium models. Rather, data mining techniques, such as **neural networks**, decision trees, support vector machines, and their bagged or boosted versions, search through the immense space of possible functional forms and parameter values to identify a model that represents the data well; and is useful for the business/science task at hand (*see Data Mining in Quality and Reliability*). Some of these representations are simpler, and thus faster to process, easier to communicate and understand than others.

Therefore, our evaluation criteria for data mining techniques, roughly speaking are (a) the accuracy of the model – which could be tied to cost/benefit implications, and (b) ease of understanding, communicating and processing, which are all correlated with the complexity of the model – where, other things being equal, simpler models are preferred.

For supervised data mining tasks, such as classification and regression, the simplest evaluation method of accuracy is a comparison of the actual values to those predicted by the model. A more accurate model is the one whose predicted results are closer to the actual values. There are however four complicating issues. The first issue is *overfitting*. Models that *memorize* the pattern for each instance in the data will be very accurate for the data sample that was used to build the model, i.e., *the training data*;

but they will not be as accurate when new instances are seen, especially if there is noise or irrelevant features. The sensible remedy is to evaluate the model on data that has not been “seen” by the method; i.e., the *test data*. This brings us to the second issue of typical *scarcity of relevant data*; i.e., data that duly represents all segments, e.g., fraudulent transactions of different types. The third issue is that the *consequences of an inaccuracy* may not be the same. Going back to the fraud detection example, the cost of a fraudulent transaction being honored is arguably higher than a legitimate transaction being denied; further the cost of the fraudulent transaction will depend on its amount. The fourth and final issue is that the process that is generating the data may *evolve over time*. In the case of fraud detection, as old tricks are discovered new ones are developed by the fraudsters.

In this article, we focus on evaluating data mining models after feature selection and parameter **estimation** have already been performed. While evaluation techniques are typically independent of the data mining task, the goodness measures are task specific. For example, misclassification rate is a common goodness measure for classification tasks, but it is not meaningful for regression, clustering, or market basket analysis.

This article is organized as follows. In the next section we introduce the common goodness measures by data mining task along with illustrative examples. In the following section we introduce evaluation techniques. Finally, we provide conclusions and pointers for further reading.

Goodness Measures by Data Mining Task

We consider the following data mining tasks: classification, regression, clustering, associations, and rank prediction.

Classification

For classification tasks, the *misclassification rate*, also called the *error rate*, is the basic measure of goodness of the classification task. It is defined as follows:

$$\text{Misclassification rate} = \frac{\text{Number of misclassified instances}}{\text{Number of all classified instances}} \quad (1)$$

2 Data Mining, Evaluation Techniques in

Table 1 Example dataset with the actual class and predictions by two alternative algorithms

Instance number	Actual class	Prediction by Algorithm 1	Prediction by Algorithm 2
1	A	A	A
2	A	A	B
3	B	B	B
4	B	A	B
5	B	A	B
6	A	A	B
7	B	B	B
8	A	A	A
9	B	B	A
10	A	A	A

The smaller the misclassification rate the better. For example, in Table 1, the misclassification rate is 20% for Algorithm 1 and 30% for Algorithm 2.

The classification is summarized in a *confusion matrix*, to help assess the accuracy for particular classes. The *misclassification cost* is a useful measure when the cost of misclassifying an instance depends on the class the instance belongs. The misclassification cost is calculated by weighing each misclassified instance by the respective misclassification cost. Continuing with the earlier example, Table 2 displays the confusion matrix for Algorithms 1 and 2, as well as the unit misclassification costs: classifying a *B* as *A* is three times as costly as misclassifying an *A* as *B*. Consequently, the misclassification costs are 6 and 5 for Algorithms 1 and 2, respectively. Even though the first algorithm is preferred on the basis of misclassification rate, the second one is more attractive when the misclassification costs are considered.

If the model produces probability estimates for an instance to belong to each class and these estimates are used as an input to a secondary decision mechanism to classify the particular instance, then it makes sense to evaluate how close these estimates are to the actual probabilities. *Quadratic loss* and *informational loss* functions are used for this purpose.

Quadratic loss has well-studied properties in statistics (also known as *mean squared error*), and it is defined as follows:

$$MSE = \sum_i \sum_j (p_{ij} - I_{ij})^2 \quad (2)$$

Here, p_{ij} is the predicted probability that the instance i belongs to class j and I_{ij} takes value 1 if instance i belongs to class j , and 0 otherwise. The square root of *mean squared error* (MSE) is called *root mean squared error* (RMSE).

Informational loss measure has its origins in computer science and is defined as

$$Informational\ Loss = -2 \sum_i \sum_j I_{ij} \log(p_{ij}) \quad (3)$$

Both of these loss measures are minimized for the model that estimates the true probabilities. In fact the informational loss becomes the entropy measure, and the quadratic loss becomes the variance of the underlying data when the model estimates the true probabilities. On the other hand, when the model does not exactly estimate the true probabilities, these two measures may prefer different models. The entropy measure prohibitively penalizes the model that assigns zero probability to the actual class even for a single observation, as it becomes infinity, while the quadratic error has a cap on the loss from a particular observation. Another difference is that under the multinomial case, while the informational loss measure is only affected by the probability assigned to the actual class, the quadratic loss measure is influenced by how the rest of the probabilities are assigned. The quadratic loss measure prefers those models that do not have another strong candidate, i.e., equally distribute the probabilities amongst the nonselected classes.

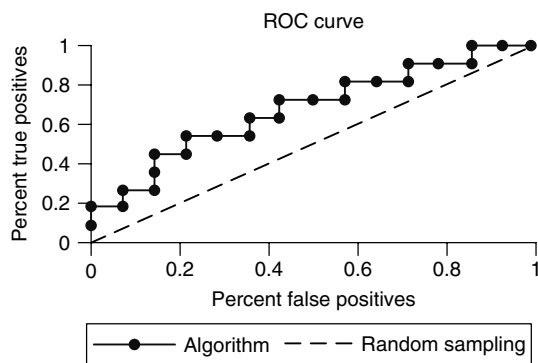
The *receiver operating characteristic (ROC) curve* is a convenient way of depicting the trade-off between only using the model where it can predict the outcome with a high confidence and thus have a low false positive rate; and using the model to find

Table 2 The confusion matrices for the two algorithms and the misclassification costs

Algorithm 1 Confusion Matrix			Algorithm 2 Confusion Matrix			Misclassification Costs		
	Predicted A	Predicted B		Predicted A	Predicted B		Predicted A	Predicted A
Actual A	5	0	Actual B	3	2	Actual A	0	1
Actual B	2	3	Actual B	1	4	Actual B	3	0

Table 3 Illustrative example for ROC curve and lift chart

Actual class	Predicted probability of class A	Cumulative % false positives	Cumulative % true positives	Cumulative % of all cases	Lift
A	0.89	–	0.09	0.04	2.27
A	0.87	–	0.18	0.08	2.27
B	0.82	0.07	0.18	0.12	1.52
A	0.81	0.07	0.27	0.16	1.70
B	0.75	0.14	0.27	0.20	1.36
A	0.69	0.14	0.36	0.24	1.52
A	0.65	0.14	0.45	0.28	1.62
B	0.63	0.21	0.45	0.32	1.42
A	0.59	0.21	0.55	0.36	1.52
B	0.45	0.29	0.55	0.40	1.36
B	0.41	0.36	0.55	0.44	1.24
A	0.40	0.36	0.64	0.48	1.33
B	0.39	0.43	0.64	0.52	1.22
A	0.37	0.43	0.73	0.56	1.30
B	0.35	0.50	0.73	0.60	1.21
B	0.35	0.57	0.73	0.64	1.14
A	0.34	0.57	0.82	0.68	1.20
B	0.30	0.64	0.82	0.72	1.14
B	0.29	0.71	0.82	0.76	1.08
A	0.21	0.71	0.91	0.80	1.14
B	0.15	0.79	0.91	0.84	1.08
B	0.08	0.86	0.91	0.88	1.03
A	0.05	0.86	1.00	0.92	1.09
B	0.05	0.93	1.00	0.96	1.04
B	0.03	1.00	1.00	1.00	1.00

**Figure 1** ROC curve for the example

as many cases of the target class as possible, and but having a higher false positive rate. The former situation corresponds to the bank calling only the most suspicious cases (i.e., those assigned a high probability) a fraud and thus avoids annoyance problems but ending up allowing quite a few fraudulent transactions go undetected. In the latter,

the bank disrupts transactions at the slightest hint of fraud and manages to catch a high proportion of fraudulent transactions, but incurs a high annoyance cost. The illustrative example in Table 3 has 25 instances that are sorted in decreasing order of predicted probability to belong to class A. Figure 1 depicts the associated ROC curve, assuming that the target class is A. The y axis is the true positive percentage (cumulative true positives/total positives) and the x axis is the false positive percentage (cumulative false positives/total negatives). The diagonal corresponds to the performance of the random targeting, i.e., the naïve performance. In general, the larger the area under the curve the better model accuracy.

A close relative of the ROC curve that is used predominantly in targeting applications in marketing is the *lift chart*. The idea is that acting on, e.g., 10% of the cases, we should be catching more than 10% of the positive cases, for the algorithm to be beneficial. Lift is defined as the ratio of the probability of the target class in sample (that we

4 Data Mining, Evaluation Techniques in

act on) to the probability of the target class in the population [1].

$$\text{Lift} = \frac{\text{Cumulative percentage of true positives acted on}}{\text{Cumulative percentage of all cases acted on}} \quad (4)$$

The higher the lift, the better the model performance. We produce the lift chart for the previous example in Figure 2, where the last two columns in Table 3 form the axes. In general, the lift starts high and decreases with larger targeted sample.

Both the lift chart and ROC curve can be used (a) to select a model and (b) to select a cutoff probability to act on an instance; however, they stop short of pointing the optimal model and target sample. The missing piece is the cost of misclassification and the benefit from correct classification, which can be incorporated to these calculations to estimate the expected profit.

Regression

For *regression tasks*, the basic idea of goodness measures is that the actual and the predicted values should be as close to each other as possible. Therefore, low *RMSE* and *mean absolute deviation (MAD)* and *mean absolute percentage error (MAPE)* values are desired. While they all prefer more accurate models, *RMSE* is more sensitive to large errors, and *MAPE* considers the error relative to the size of the prediction. The R^2 which is based on the squared error,

expresses the percentage of the variability explained by the model in a scale of 0–100%.

$$RMSE = \sqrt{\sum_i (e_i)^2 n^{-1}}; MAD = \sum_i |e_i| n^{-1};$$

$$MAPE = \sum_i |e_i| |x_i|^{-1} n^{-1} \quad (5)$$

$$R^2 = 1 - \frac{\sum_i (e_i)^2}{\sum_i (x_i - \bar{x})^2} \quad (6)$$

Here, e_i stands for the error, i.e., the difference between the actual and predicted value, for instance i , x_i is the actual observation value, for instance i , $\bar{x}$ is the average of the actual observation values, and n is the number of observations.

Other measures, such as an information criterion (AIC) proposed by Akaike [2], and Bayesian information criterion (BIC) – also called *Schwarz information criterion (SIC)* [3], measure model accuracy while penalizing models that are more complicated.

$$AIC = -2 \log L + 2k \quad (7)$$

$$BIC = -2 \log L + k \log(n) \quad (8)$$

Here, k is the number of free parameters in the model, and L is the likelihood of the observed data given the model, while n is the number of observations.

Minimum description length (MDL) principle developed by Rissanen [4] seeks to combine measures of model accuracy and model simplicity, and suggests that the best model is the one that provides the shortest description of the data, where the description length resembles the number of digits in a binary string used to code the transmissions [5].

Rank prediction

While *rank prediction tasks* are frequently handled as a scoring task, the implied assumption of equal (or when logs are taken uniformly decreasing) distance between the subsequent instances may be misleading. Other more suitable measures include nonparametric correlation coefficients commonly used in statistics such as the *Spearman's ρ* and *Kendall's τ* [6]. Spearman's ρ is basically the Pearson correlation coefficient applied on the rankings. Kendall's τ represents the difference between the probability

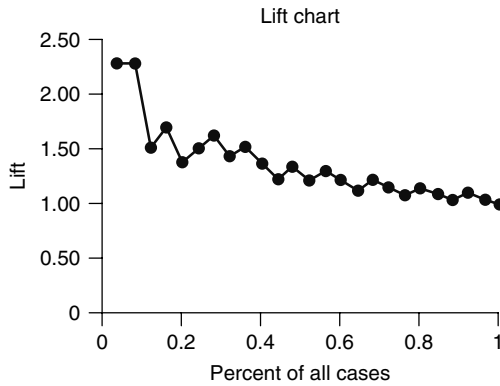


Figure 2 Lift chart for the example

that the observations have the same order and the probability that they are ordered differently, based on the concordant and discordant pairs of predicted and actual values.

Clustering

For *clustering tasks*, the goodness measures used are (a) *within cluster distance* (or variability) and (b) *between cluster distance* (or variability). The R^2 measure defined as the ratio of within cluster variability to between cluster variability combines the two measures. As expected, this measure improves as the number of clusters increases. *MDL* has been suggested as a natural fit for clustering evaluation.

Market Basket Analysis

For *market basket analysis*, or *association tasks*, the evaluation issue is not accuracy, since the result from all algorithms is the same. The relevant goodness measure is processing speed [7].

Evaluation Techniques

To summarize our discussion in the introduction section, evaluation techniques should be able to answer the following questions: how accurate is the model? Does it have repeatable performance over different data sets and over time as the underlying process evolves? Is it just memorizing the data? Is it accurate in situations that are of high importance?

Estimating Performance of Algorithms

The *resubstitution estimate* refers to the evaluation of the model on the training data set. For example, in the earlier example of Table 1, the resubstitution misclassification rate of Algorithm 1 was 20%. The resubstitution estimate, of course, overestimates the true accuracy of the model, since the model's parameters have been optimized to yield minimum misclassification rate on this very sample. In other words, we are not able to evaluate whether and to what extent the model is memorizing the dataset down to the irrelevant facts, i.e., noise, and will misclassify when it is presented a different random sample from the same underlying data generation process.

An unbiased accuracy estimate can be obtained by evaluating the accuracy of the model on a *holdout dataset*, i.e., a representative dataset that was not used for training the model, which includes feature and technique selection and parameter estimation steps. (Especially in **time series** analyzes, models trained on a particular time period are tested on holdout data that is from a more recent time period to evaluate whether the model is still valid.) As the overall dataset is partitioned into training and test data, stratified sampling can be used to ensure representativeness. For classification tasks stratification can be performed in terms of the proportion of classes, while for regression tasks proportion of high/medium/low actual values can be taken as basis.

Representativeness, of course, goes beyond just the response variable, the relationship between the x 's are also of high importance for training the model and for evaluating its accuracy. However, the single holdout dataset approach does not provide a variability estimate for the accuracy. It would be nice to be able to repeatedly train the algorithm on a representative set and test on a different dataset, and see how robust the accuracy is. The problem is that relevant and clean data representative of the underlying process is a scarce resource.

The *v-fold cross-validation* addresses most of these concerns. The idea is to partition the available dataset to v (typically 10) pieces randomly, repeatedly leaving out one piece of the data for testing, while training the algorithm on the remaining data. Accuracy is then calculated for each test set. The average becomes the point estimate, and the standard deviation becomes the variability estimate of the algorithm's accuracy on random samples coming from the same underlying data generation process. For added assurance on representativeness of the samples resulting from the folds, this technique can be used with stratification.

A common extension is the *10-by-10-fold cross-validation* in which case the 10-fold cross-validation is repeated 10 times to get a better estimate of the standard deviation.

A special case of the v -fold cross-validation is the *leave-one-out cross-validation*, also called the *jack-knife estimate* [8]. Here, for a dataset with n observations, we perform n -fold cross-validation. Notice that the average and standard deviation of the accuracy will not vary with the sequence of the points left out. While this technique eliminates the sampling issues,

the computational effort is substantially higher, especially for large datasets.

We can use statistical theory to provide a $(1-\alpha)\%$ **confidence interval** for the true accuracy as follows:

$$\bar{a} \pm t_{v-1, \alpha/2} \text{std}(a) v^{-0.5}$$

where $\bar{a}$ is the average accuracy, and $\text{std}(a)$ is the standard deviation of accuracy calculated from the samples. $P(t > t_{v-1, \alpha/2}) = \alpha/2$, where t has the t distribution with $v - 1$ **degrees of freedom**. For example, constructing the 95% confidence interval for 10-fold cross-validation, v is 10, $t_{9, 0.025}$ is 2.26 as can be read from tables available in any elementary statistics book, or retrieved with the *tin* function in Excel. We can make these intervals tighter if we are happy to accept a lower confidence level such as 90%, whereas higher confidence levels imply a larger interval.

A related sampling scheme to the v -fold cross-validation is the **bootstrap** [8], where the training data is sampled from the original data with replacement, and those observations that are not picked for training form the test dataset. This process is repeated many times and results are averaged. While the bootstrap provides a good method of evaluation for very small datasets, it is not commonly used for large datasets.

Comparing Performance of Algorithms

Frequently, we are concerned with comparing the accuracy of two algorithms operating on the same data. Considering the uncertainty in the accuracy estimates, comparing the point estimates; i.e., the average accuracy of each algorithm is misleading. Even if two algorithms have the same true accuracy, the probability that their accuracy on a sample will be the same is very low, and it is zero for regression tasks. What should be compared are the confidence intervals of the accuracy, if these are overlapping, then there is not sufficient evidence to conclude that these two algorithms perform differently in terms of accuracy on the data coming from the same underlying data generation process.

There are three ways in which we can be more decisive about the relative accuracy of algorithms. The first involves decreasing confidence levels, say from 95 to 90%. The second involves increasing the number of partitions for the v -fold cross-validation, thus reducing the t -statistic and the $v^{-0.5}$ term. Notice

that the effect of adding more folds decreases as the number of folds increases, for example, the decrease in the multiplier of the $\text{std}(a)$; i.e., $t_{v-1, \alpha/2} v^{-0.5}$, from $v = 10$ to $v = 20$ is only 0.27 while the decrease from $v = 10$ to $v = 5$ is 0.63, as can be seen from Table 4.

The third way to be more decisive about the relative accuracy of algorithms is to construct the test statistic of the paired difference: the difference of the accuracy measures for the two algorithms is calculated at each step of the v -fold cross-validation, thus ensuring that the differences observed in the algorithm's accuracy are not stemming from being exposed to different data samples. The average of the differences provides a point estimate for the difference in accuracy between the two algorithms. The standard deviation of the v differences is used to test the hypothesis that the difference in accuracy between the two algorithms is zero. The test statistic $\bar{d} v^{0.5} / \text{std}(d)$ is compared against the t distribution with $v - 1$ degrees of freedom. Typically, a p value less than 5% indicates reasonable evidence that the accuracy of the two algorithms is not the same. Taking the alternative hypothesis, as any of the two could be more accurate, the t -statistic with 9 degrees of freedom and $\alpha = 0.025$ can be used as a reference value (2.262). Hence, if the absolute value of the test statistic exceeds this reference value, we can conclude that the true accuracy of the algorithms is different.

While we have focused on accuracy evaluation in this section, the above evaluation techniques can be used with other goodness measures as well, such as processing speed.

Conclusion and Further Reading

The knowledge discovery in databases (KDD) and data mining movement initially emphasized finding

Table 4 The multiplier of $\text{std}(a)$ in the confidence interval of accuracy

		v in v -fold cross-validation		
		5	10	20
Confidence level	90%		0.611	
	95%	1.388	0.754	0.480
	99%		1.083	

actionable/interesting knowledge in the data. “Interestingness” was defined as a somewhat subjective concept containing validity, novelty, usefulness, and simplicity [9]. For further references on developments on interestingness measures for data mining we refer the reader to the survey by Geng and Hamilton [7].

In conclusion, the two main criteria that are used in the evaluation of data mining algorithms are accuracy and simplicity. While the MDL principle suggested by Rissanen [4] is generating a lot of interest [5], and while there are some measures related to the MDL such as AIC and BIC, due to practical considerations, accuracy and simplicity are typically evaluated separately. Tenfold cross-validation is the standard method [10].

References

- [1] Berry, M.J.A. & Linoff, G.S. (2004). *Data Mining Techniques: For Marketing, Sales, and Customer Support*, 2nd Edition, John Wiley & Sons, pp. 81–84.
- [2] Akaike, H. (1981). Likelihood of a model and information criteria, *Journal of Econometrics* **16**, 3–14.
- [3] Schwarz, G. (1978). Estimating the dimension of a model, *Annals of Statistics* **6**(2), 461–464.
- [4] Rissanen, J. (1983). A universal prior for integers and estimation by minimum description length, *The Annals of Statistics* **11**(2), 416–431.
- [5] Hansen, M.H. & Yu, B. (2001). Model selection and the principle of minimum description length, *Journal of the American Statistical Association* **96**(454), 746–774.
- [6] Rosset, S., Perlich, C. & Zadrozny, B. (2005). Ranking-based evaluation of regression models, in *Fifth IEEE International Conference on Data Mining*, Houston, Texas, pp. 1–8.
- [7] Geng, L. & Hamilton, H.J. (2006). Interestingness measures for data mining: a survey, *ACM Computing Surveys* **38**(3), 1–32.
- [8] Efron, B. & Gong, G. (1983). A leisurely look at the bootstrap, the jackknife, and cross-validation, *The American Statistician* **37**(1), 36–48.
- [9] Fayyad, U., Piatetsky-Shapiro, G. & Smyth, P. (1996). Knowledge discovery and data mining: towards a unifying framework, *Proceedings of the Second International Conference on Knowledge Discovery and Data Mining (KDD-96)*, AAAI Press, Portland, August 2–4.
- [10] Witten, I. & Frank, E. (2000). *Data Mining, Practical Machine Learning Tools and Techniques with Java Implementations*, Morgan Kaufmann Publishers, pp. 119–156.

Related Articles

Data Mining in Quality and Reliability; Processing of Survey Data; Statistics in Data and Information Quality

ÖZDEN F. GÜR ALI

Applications of Extreme Statistics in Science and Engineering

Generalities

Applications of **risk analysis** in science and engineering abound. For some theoretical background and the resulting statistical issue, we refer to **Risk Analysis, Extremes in**. Before we illustrate this use of extreme value theory we should mention that real life applications have been responsible for major developments of the theory. It is actually remarkable that historically the basic Fisher–Tippett result was proved long before the appearance of Cramér’s treaty of the central limit problem in Cramér [1]. In spite of some earlier attempts like in Gumbel [2], the genuine development of statistics of extremes only took off in the last two decades of the twentieth century. The appearance of theoretical texts like de Haan [3] and Leadbetter *et al.* [4] have instigated this development.

Nevertheless, many problems related to extremal thinking already appeared in specific applications. Large claims have been modeled by **Pareto**-type distributions already in Jung [5] in 1964. In the 1970s, wind speed data with their extremes have been investigated and modeled using extreme value procedures, see Simiu *et al.* [6].

Moreover, applications have been responsible for a change in statistical thinking in that the usual central limit paradigms had to be replaced by intrinsic extreme value thinking. The fact that investigators were forced to draw statistical conclusions from a meager set of rare events is one of the reasons for this approach. Moreover, extreme values are often only rough guesses like in the case of casualties from major natural catastrophes. Perhaps in some cases one might be able to expand the meager data set by proper inclusion of calibrated data that refer to other timescales and/or geographical areas.

In the examples below we will avoid to emphasize again and again the important role played by the extremes. The latter should be clear from the specific context. In an attempt not to overload the list of references we only refer to the main historical contributions and/or to recent publications where more information and references can be found. For

an early publication with a variety of applications of extreme value theory, see Galambos *et al.* [7]. For a recent textbook treatment of *operational risk*, see Panjer [8].

Operations Research

Network Models

Owing to the extreme heterogeneity in machines, subnetworks, and administrative organizations over which networks are laid, customers must be aware that certain constituent components can cause problems with the performance and/or reliability of the network.

Van Uitert and Borst [9] investigate the quality of service in integrated-services networks. The overall workload may be influenced if not controlled by one particular flow that has long-tailed traffic characteristics. As studied in Resnick [10] heavy-tailed **time series** result from the increasing availability of large data sets from communication networks in which heavy-tailed distributions are ubiquitous and long-range dependence may be present. For an application of bivariate extreme value theory to Internet traffic data, see Heffernan and Resnick [11].

Reliability

Structural systems are often subject to severe forms of risk. Load combinations for bridges, corrosion, metal fatigue are but a few key words that come to mind when dealing with such risks. For a compilation of contributions in this area, see Maes and Huyse [12]. Structural systems are subject to deterioration over time due to a variety of reasons: improper use, exposure to aggressive environment, wear, extreme accidents, and so on. Usually, these reasons can be subdivided into two not necessarily disjoint categories; the first being aging that affects the lifetime of the structure on a continuous basis, the second being shocks that act on a discrete timescale. Stochastic modeling of such events suggests the use of a combination of stochastic differential equations coping with the continuous scale and a superimposed discrete-time process dealing with the shocks. See Crowder *et al.* [13] for a textbook treatment and Ciampoli [14] for an example of a reinforced concrete shear wall in a nuclear power plant, subject to deterioration.

2 Applications of Extreme Statistics

Many forms of structure fail due to combinations of extremal processes. In Coles and Tawn [15], multivariate extremes are introduced as elements in the design process.

- Heffernan and Tawn [16] studied data on *wave impacts on ships*. The data were obtained from physical model studies for different models of bulk carriers and a broad range of sea and ship conditions. The generalized Pareto distribution is used to describe the distribution of wave impact pressures exceeding some threshold and a **Poisson process** for the number of such impacts.
- **Software reliability** has been treated in Tsokos and Nadarajah [17] where the reliability of a new software package is determined prior to being marketed.
- In **genetic algorithms** one only seems to use random variables with normal-like behavior. Nevertheless the performance of these algorithms often proves to be controlled by extremal rather than central values, see Adler *et al.* [18].
- *Corrosion* of pipelines is the subject of Isogai *et al.* [19] where the pit depth of metal oxide films is modeled while the stress is modeled by a compound process.
- Probably one of the most important applications of extreme value thinking concerns metal *fatigue strength* where the largest inclusions are at the origin of this phenomenon. The mathematical treatment of this problem is made complicated by the fact that inference on three-dimensional inclusions has to be based on two-dimensional cross sections, see Murikami and Beretta [20] and Beretta and Anderson [21].

Geophysics

The whole area of geophysics offers vast possibilities to apply extreme value theory. For a general text, see Smith [22].

Hydrology

- *River floods and droughts* are the extremes when considering river levels. An overshoot of a certain safety threshold causes floods, but droughts are caused when the level gets below a lower safety margin. There already exists a wide

variety in treated problems and corresponding methodologies. Let us first look at the river levels at a single point. As is illustrated in the section on genetic algorithms, hydrologists can draw statistical conclusions on the basis either of one or more annual maxima or of floods above a given threshold, see Li *et al.* [23] for this approach. Engeland *et al.* [24], apply different extremal models and estimation procedures to fit flood extremes and droughts durations.

Multivariate procedures have also found applications in hydrology. Shiau [25] combines flood volume and flood peak in a bivariate extreme value model. Nagaraja *et al.* [26] look at extreme river levels of the same river but at different sites. Extreme floods are modeled by fractal theory in Turcotte [27]. See Hutson [28] for the construction of a **confidence interval** for the **quantiles** in a hydrological data set.

- *Rainfall* data add a spatial component to extreme value theory. Cowpertwait [29] models rain storms by a Poisson process while the storms themselves generate a cluster of heavy and light cells. For a time series approach, see Coles and Tawn [30]. For a Bayesian approach, see Coles and Tawn [31]. Joe *et al.* [32] treat bivariate data on sulfate and nitrate levels taken from a major study of acid rain.

Oceanography

It goes without saying that extreme sea levels are among the most dangerous natural events. The number of human casualties at recent coastal floods can hardly be estimated. As explained in Dixon and Tawn [33], the sea level is the composition of meteorological surge and astronomical tidal processes. It exhibits temporal nonstationarity due to a combination of long-term trends in the mean level, the deterministic tidal component, surge seasonality, and interactions between tide and surge. Ignoring nonstationarity usually leads to an underestimation of extreme sea levels.

Ocean waves are described by crest height, trough height, and their interrelationship. Return periods of sea storms in which either the maximum crest or trough height exceeds a fixed threshold are discussed in Arena [34]. Extreme-event analysis is applied to the time series of sea data in specific locations by Tirozzi *et al.* [35]. For a trivariate model for extremal sea-level data, see Tawn [36]. For one based on

the Gumbel distribution, see Escalante-Sandoval and Raynal-Villasenor [37]. For a bivariate extremal treatment of wind velocities and wave heights, see Liu *et al.* [38].

de Haan and Rootzén [39] have investigated height requirements for *sea dikes*. Bruun and Tawn [40] study the combination of extreme water levels and wave heights. For a generalization to Gaussian fields and seas, see Podgórski *et al.* [41].

Meteorology, Climatology

Risks in climatology offer a wealth of applications of extreme value theory. The assumptions of independence and/or of stationarity in the data are mostly violated. Fortunately as shown *inter alia* in Leadbetter *et al.* [4] and subsequent research, the extreme value distributions remain useful candidates to describe yearly maxima. Estimations are often improved by relying on more of the largest data values. For an early attempt in this direction, see Buishand [42].

- *Meteorological data* generally have no alarming effect as long as they are situated in a fixed band around the average. Once the boundaries of this safe region are crossed, extreme value theory comes into action. Let us mention a couple of examples, first of all, extreme *temperatures*. Davison and Ramesh [43], offer a semiparametric approach to sample extremes. Beirlant *et al.* [44] treat a case study of temperature data using a Markov model. Hall and Tajvidi [45] use a non-parametric approach to estimate temporal trends in a weakly dependent time series of maximal temperatures.
- *Wind* encompasses a number of meteorologically different subjects like storms, tornadoes, hurricanes, typhoons, and so on. Kestens and Teugels [46] survey some of the modeling problems. See also Simiu *et al.* [6] and Beirlant *et al.* [44] for large-scale investigations of wind speed data. Since data on hurricane wind speed are often unreliable, Casson and Coles [47] have developed a two-stage approach starting from a spatial pressure model and the resulting associated wind speeds. In Teugels and Vandewalle [48], it is shown that casualty data from hurricanes are well modeled by a Pareto distribution with an extreme value index (EVI) $\gamma > 1$, indicating that the data do not support the finiteness of the variance, even of the mean.

- Other delicate risk-oriented topics within climatology are *global warming* and *ozone depletion*. For attempts to model the latter using extreme value theory, see Smith [49] and Dasgupta and Bhaumik [50]. Both subjects are commonly discussed within the realm of *environmetrics*, a scientific area where the role of extremes can hardly be overestimated. Pollution levels for rivers do not give rise to specific actions as long as they are considered to be about average. Extreme value theory comes into play as soon as the pollution level hits a value that is too high. Many similar examples can be found when one is interested in the values of environmental indicators.

Seismology

Earthquakes form another major area of application of extreme value analysis. Pisarenko and Sornette [51] analyzed shallow earthquakes over the period 1977–2000. Also in Beirlant *et al.* [44] the same data have been subjected to further extremal analysis.

Economics

Also in **finance** and insurance, heavy-tailed models have been popular. See for example the survey by Zinchenko [52] and the textbook by Embrechts *et al.* [53].

Finance

One of the key Black–Scholes assumptions is that price movements in financial markets (assets, credit, etc.) follow the lognormal distribution. In practice, however this is not always the case. Empirical distributions exhibit far more frequent extreme outcomes of *heavy tails*. This observation has forced financial mathematics to introduce concepts like heavy-tailed distributions, long-range dependency, and extreme value behavior. For a recent comprehensive treatment, see the textbook by Malevergne and Sornette [54].

- A first area of interest is the construction and estimation of *risk measures*, especially for heavy-tailed distributions. The most popular such measure is the *value at risk* but there are many alternatives, see, for example, Suárez and Carrillo [55] and Cai and Li [56].

- Another area is the modeling of one-dimensional *asset returns* with heavy tails. Like in most applications, extreme value statistics can only offer partial answers since such large returns are rare. Underestimation usually results from ignoring the potential heaviness of the tail of the distribution; but then again overestimation is a consequence of taking this tail as too heavy. A first form of dependence can be found within the same financial environment when successive returns are no longer assumed to be independent. See Deo [57] and Schmidbauer and Rösch [58] for models that incorporate this dependence for stock returns.
- Another is the detection and modeling of *dependency* among heavy-tailed assets. It turns out that large risks are not solely due to the heavy tails of the asset returns but also to their collective behavior. This observation lies at the introduction of measures of dependence, often expressible in terms of copulas. Dependence can be measured by the *coefficient of tail dependence* that quantifies the probability for an asset to lose a large amount knowing that another asset (or the market) has dropped significantly, see also Cizek *et al.* [59]. In Schmidbauer and Rösch [58] an attempt is made to compare this dependence between the assets in periods of low and intense market activity. For a case study, see Poon *et al.* [60]. Another illustration of such dependence is offered by the interaction between portfolio quality and credit losses, see Lucas *et al.* [61]. For an application of Markov chain modeling, see Bortot and Coles [62].

Insurance

Historically actuarial science has been one of the first areas in which extremes have played a vital role. For a general introduction, see the article on extremes by Beirlant [63]. In *nonlife insurance*, certain portfolios have a tendency to occasionally include a large claim that jeopardizes the solvency of a portfolio. Apart from major accidents as earthquakes, hurricanes, airplane crashes, and so on, there are a vast number of occasions where large claims occur. Many illustrations from general insurance can be given.

- The distribution of the *total claim amount* in such a portfolio is very sensitive to the inclusion of

heavy-tailed distributions. Mikosch and Nagaev [64] give a treatment that puts the theories of extremes and large deviations in insurance in perspective.

- Morales [65] constructs a simulation scheme to estimate *ultimate ruin probabilities* for the classical surplus model in the presence of heavy- and medium-tailed severity distributions.
- See Beirlant *et al.* [66, 44] for treatments of *car accident* data.
- Rootzén and Tajvidi [67] look for statistical connections between large insurance claims from *windstorm losses* and wind speeds.
- *Fire portfolios* sometimes encounter large claims. Industrial fires, especially, cause a lot of side effects in loss of property, temporary unemployment, and lost contracts.
- Still other areas of nonlife insurance can give raise to large claims. We mention accidents with oil tankers and platforms, environmental spills, explosions at nuclear power stations, and so on. A special mention should be given to lawsuits resulting from medical malpractice that often lead to excessive claims from the insured.

Let us mention that *reinsurance* is a branch of actuarial practice that studies risk sharing by passing on part of the risk to one or more other insurance companies. Also the calculation of *insurance premiums* forces both the insurer and the reinsurer to be aware of potential extremes in the data. For such an example, see Jang and Krvavych [68].

References

- [1] Cramér, H. (1946). *Mathematical Methods of Statistics, Princeton Landmarks in Mathematics*, Princeton University Press.
- [2] Gumbel, E.J. (1958). *Statistics of Extremes*, Columbia University Press, New York.
- [3] de Haan, L. (1970). *On Regular Variation and its Applications to the Weak Convergence of Sample Extremes*, *Mathematical Centre Tract Vol. 32*, Mathematisch Centrum, Amsterdam.
- [4] Leadbetter, M.R., Lindgren, G. & Rootzén, H. (1982). *Extremal and Related Properties of Stationary Processes*, Springer-Verlag, New York.
- [5] Jung, J. (1964). On the use of extreme values to estimate the premium for an excess of loss reinsurance, *The ASTIN Bulletin* 3, 178–184.

- [6] Simiu, E., Changery, M.J. & Filliben, J.J. (1979). *Extreme Wind Speeds at 129 Stations in the Contiguous United States*, NBS Building Science Series 118, National Bureau of Standards, Washington, DC.
- [7] Galambos, J., Lechner, J. & Simiu, E. (1994). *Extreme Value Theory and Applications*, Kluwer Academic Publishers, Dordrecht.
- [8] Panjer, H. (2006). *Operational Risk. Modeling Analytics*, John Wiley & Sons, Hoboken.
- [9] van Uitert, M. & Borst, S. (2002). A reduced-load equivalence for generalized processor sharing networks with long-tailed input flows, *Queueing Systems* **41**, 123–163.
- [10] Resnick, S.I. (1997). Heavy tail modeling and teletraffic data, *Annals of Statistics* **25**, 1805–1869.
- [11] Heffernan, J. & Resnick, S. (2005). Hidden regular variation and the rank transform, *Advances in Applied Probability* **37**, 393–414.
- [12] Maes, M. & Huyse, L. (2004). *Reliability and Optimization of Structural Systems*, Balkema Publishers, Leiden.
- [13] Crowder, M.J., Kimber, A.C., Smith, R.L. & Sweeting, T.J. (1991). *Statistical Analysis of Reliability Data*, Chapman & Hall, London.
- [14] Ciampoli, M. (1999). A probabilistic methodology to assess the reliability of deteriorating structural elements, *Computer Methods in Applied Mechanics and Engineering* **168**, 207–220.
- [15] Coles, S.G. & Tawn, J.A. (1994). Statistical methods for multivariate extremes: an application to structural design, *Journal of the Royal Statistical Society. Series C* **43**, 1–48.
- [16] Heffernan, J. & Tawn, J.A. (2001). Extreme value analysis of a large designed experiment: a case study in bulk carrier safety, *Extremes* **4**, 359–398.
- [17] Tsokos, C. & Nadarajah, S. (2003). Extreme value models for software reliability, *Stochastic Analysis and Applications* **21**, 719–735.
- [18] Adler, R., Feldman, R. & Taqqu, M. (eds) (1998). *A Practical Guide to Heavy Tailed Data*, Birkhäuser, Boston.
- [19] Isogai, T., Katano, Y. & Miyata, K. (2004). Models and inference for corrosion pit depth data, *Extremes* **7**, 253–270.
- [20] Murikami, Y. & Beretta, S. (1999). Small defects and inhomogeneities in fatigue strength: experiments, models and statistical implications, *Extremes* **2**, 123–147.
- [21] Beretta, S. & Anderson, C.W. (2002). Extreme value statistics in metal fatigue, *Atti della XLI Riunione Scientifica della Società Italiana di Statistica* 251–260.
- [22] Smith, R.L. (2001). *Spatial statistics in environmental science*, in *Nonlinear and Nonstationary Signal Processing*, W.J. Fitzgerald, R.L. Smith, A.T. Walden & P. Younged, eds, Cambridge University Press, pp. 152–183.
- [23] Li, Y., Saxena, K.L.L. & Cong, S. (1999). Estimations of the extreme flow distributions by stochastic models, *Extremes* **1**, 423–448.
- [24] Engeland, K., Hisdal, H. & Frigessi, A. (2004). Practical extreme value modelling of hydrological floods and draughts: a case study, *Extremes* **7**, 5–30.
- [25] Shiau, J.T. (2003). Return period of bivariate distributed extreme hydrological events, *Stochastic Environmental Research and Risk Assessment* **17**, 42–57.
- [26] Nagaraja, H.N., Choudhary, P.K. & Matalas, N. (2002). Number of records in a bivariate sample with application to Missouri river flood data, *Methodology and Computing in Applied Probability* **4**, 377–391.
- [27] Turcotte, D.L. (1994). Fractal theory and the estimation of extreme floods, *Journal of Research of the National Institute of Standards and Technology* **99**, 377–389.
- [28] Hutson, A.D. (1999). Calculating nonparametric confidence intervals for quantiles using fractional order statistics, *Journal of Applied Statistics* **26**, 343–353.
- [29] Cowpertwait, P.S. (1994). A generalized point process model for rainfall, *Proceedings of the Royal Society of London. Series A* **447**, 23–37.
- [30] Coles, S.G. & Tawn, J.A. (1996). Modelling extremes of the areal rainfall process, *Journal of the Royal Statistical Society. Series B* **58**, 329–347.
- [31] Coles, S.G. & Tawn, J.A. (1996). Modelling extremes: a Bayesian approach, *Applied Statistics* **45**, 463–478.
- [32] Joe, B., Smith, R.L. & Weissman, I. (1992). Bivariate threshold methods for extremes, *Journal of the Royal Statistical Society. Series B* **54**, 171–183.
- [33] Dixon, M.J. & Tawn, J.A. (1999). The effect of non-stationarity on extreme sea-level estimation, *Journal of the Royal Statistical Society. Series C* **48**, 135–151.
- [34] Arena, F. (2002). Ocean wave statistics for engineering applications, *Rendiconti del Circolo Matematico di Palermo. Serie II. Supplemento* **70**, 21–52.
- [35] Tirozzi, B., Puca, S., Pittalis, S., Bruschi, A., Morusci, S., Ferraro, E. & Corsini, S. (2006). *Neural Networks and Sea Time Series. Reconstruction and Extreme-Event Analysis*, Birkhäuser, Boston.
- [36] Tawn, J.A. (1990). Modelling multivariate extreme value distributions, *Biometrika* **77**, 245–253.
- [37] Escalante-Sandoval, C.A. & Raynal-Villasenor, J.A. (1998). *Journal of Statistical Computation and Simulation* **61**, 313–340.
- [38] Liu, D., Wen, S. & Wang, L. (2002). Poisson-Gumbel mixed compound distribution and its application, *Chinese Science Bulletin* **47**, 1901–1906.
- [39] de Haan, L. & Rootzén, H. (1993). On the estimation of high quantiles, *Journal of Statistical Planning and Inference* **35**, 1–13.
- [40] Bruun, J.T. & Tawn, J.A. (1998). Comparison of approaches for estimating the probability of coastal flooding, *Journal of the Royal Statistical Society. Series C* **47**, 405–423.
- [41] Podgórski, K., Rychlik, I., Rydén, J. & Sjö, E. (2000). How big are the big waves in a Gaussian sea? *International Journal of Offshore and Polar Engineering* **10**, 161–169.
- [42] Buishand, T.A. (1989). Statistics of extremes in climatology, *Statistica Neerlandica* **43**, 1–30.

- [43] Davison, A.C. & Ramesh, N.I. (2000). Local likelihood smoothing of sample extremes, *Journal of the Royal Statistical Society. Series B* **62**, 191–208.
- [44] Beirlant, J., Goegebeur, Y., Segers, J. & Teugels, J.L. (2004). *Statistics of Extremes*, John Wiley & Sons, Chichester.
- [45] Hall, P. & Tajvidi, N. (2000). Nonparametric analysis of temporal trend when fitting parametric models to extreme-value data, *Statistical Science* **15**, 153–167.
- [46] Kestens, E. & Teugels, J.L. (2002). Challenges in modelling stochasticity of wind, *Environmetrics* **13**, 821–830.
- [47] Casson, E. & Coles, S. (2000). Simulation and extremal analysis of hurricane events, *Journal of the Royal Statistical Society. Series C* **49**, 227–245.
- [48] Teugels, J.L. & Vandewalle, B. (2006). Statistical analysis of catastrophic events, in *Coping with Uncertainty*, K. Marti, ed, Springer.
- [49] Smith, R.L. (1989). Extreme value analysis of environmental time series: an application to trend detection in ground-level ozone, *Statistical Science* **4**, 367–393.
- [50] Dasgupta, R. & Bhaumik, D.K. (1995). Upper and lower tolerance limits of atmospheric ozone level and extreme value distribution, *Sankhya. Series B* **57**, 182–199.
- [51] Pisarenko, V.F. & Sornette, D. (2003). Characterisation of the frequency of extreme events by the generalized Pareto distribution, *Pure and Applied Geophysics* **160**, 2343–2364.
- [52] Zinchenko, N.M. (2001). Heavy-tailed models in finance and insurance: a survey, *Theory of Stochastic Processes* **7**, 346–362.
- [53] Embrechts, P.A.L., Klüppelberg, C. & Mikosch, T. (1997). *Modelling Extremal Events for Finance and Insurance*, Springer-Verlag, Berlin.
- [54] Malevergne, Y. & Sornette, D. (2006). *Extreme Financial Risks. From Dependence to Risk Management*, Springer-Verlag, Berlin.
- [55] Suárez, A. & Carrillo, S. (2003). Computational tools for the analysis of market risk, *Computational Economics* **21**, 153–172.
- [56] Cai, J. & Li, H. (2005). Conditional tail expectations for multivariate phase-type distributions, *Journal of Applied Probability* **42**, 810–825.
- [57] Deo, R.S. (2000). On estimation and testing goodness of fit for m -dependent stable sequences, *Journal of Econometrics* **99**, 349–372.
- [58] Schmidbauer, H. & Rösch, A. (2004). Joint threshold exceedances of stock returns in bull and bear periods, *The Central European Journal of Operations Research* **12**, 197–209.
- [59] Cizek, P., Härdle, W. & Weron, R. (2005). *Statistical Tools for Finance and Insurance*, Springer-Verlag, Berlin.
- [60] Poon, S.-H., Rockinger, M. & Tawn, J. (2003). Modelling extreme-value dependence in international stock markets, *Statistica Sinica* **13**, 929–953.
- [61] Lucas, A., Klaassen, P., Spreij, P. & Straetmans, S. (2003). Tail behaviour of credit loss distributions for general latent factor models, *Applied Mathematical Finance* **10**, 337–357.
- [62] Bortot, P. & Coles, S. (2003). Extremes of Markov chains with tail switching potential, *Journal of the Royal Statistical Society. Series B* **65**, 851–867.
- [63] Beirlant, J. (2004). Extremes, in *Encyclopedia of Actuarial Science*, J.L. Teugels & B. Sundt, eds, John Wiley & Sons, Chichester, pp. 654–661.
- [64] Mikosch, T. & Nagaev, A.V. (1998). Large deviations of heavy-tailed sums with applications in insurance, *Extremes* **1**, 81–110.
- [65] Morales, M. (2004). On an approximation for the surplus process using extreme value theory: applications in ruin theory and reinsurance pricing, *North American Actuarial Journal* **8**, 44–66.
- [66] Beirlant, J., Teugels, J.L. & Vynckier, P. (1993). Extremes in non-life insurance, in *Extremes Value Theory and Applications*, J. Galambos, J. Lechner & E. Simiu, eds, Kluwer Academic Publishers, Dordrecht, pp. 489–510.
- [67] Rootzén, H. & Tajvidi, N. (1997). Extreme value statistics and windstorm losses: a case study, *Scandinavian Actuarial Journal* **70**–94.
- [68] Jang, J.-W. & Kravavych, Y. (2004). Arbitrage-free premium calculation for extreme losses using the shot noise process and the Esscher transform, *Insurance, Mathematics and Economics* **35**, 97–111.

Related Article

Risk Analysis, Extremes in.

JOZEF L. TEUGELS

Polya Trees and Their Use in Reliability and Survival Analysis

A Polya Tree Prior Centered at the Weibull Family

Polya trees generalize and add flexibility to standard parametric models such normal, Weibull, lognormal, Pareto, log-logistic, gamma, and others used for the statistical analysis of time to event data. A simple Weibull model for survival times $\mathbf{T} = (T_1, \dots, T_n)'$ assumes a distribution function G with a particular form, i.e.

$$T_1, \dots, T_n | G \stackrel{\text{i.i.d.}}{\sim} G, G(t|\alpha, \lambda) = 1 - \exp(-t^\alpha/\lambda) \quad \text{for } t \geq 0 \quad (1)$$

This type of model is referred to as *parametric* because the class of possible density shapes is indexed by a small number of parameters, here $\theta = (\alpha, \lambda)$. Parametric survival models are straightforward to fit and draw inferences from, but the trade-off for this ease of implementation is lack of flexibility to capture more complex features (e.g., multimodality, **skewness**, and clustering) that may be present in time to event data. This motivates expanding the class of distributions to include not only the Weibull family (or more generally any parametric family) but also deviations from it supported by the data. A Bayesian nonparametric Polya tree prior for G achieves this flexibility. Other approaches include the Dirichlet process [1], Dirichlet process mixtures, beta, and gamma processes (*see Dirichlet Process, Simulation of; Survival Analysis, Nonparametric; Lévy Processes, Simulation of*).

A Polya tree adds $2^{J+1} - 2$ additional parameters, $\mathcal{X}_J$, to equation (1) that collectively “adjust” the overall shape of G to add more probability mass where data are clumped and remove mass where data are sparse relative to the Weibull distribution $G(t|\alpha, \lambda) = 1 - \exp(-t^\alpha/\lambda)$. The parameters $\theta = (\alpha, \lambda)$ serve to fix the overall shape, location, and spread of the distribution G ; the parameters $\mathcal{X}_J$ refine its shape to match the observed data more closely. In fact, the parameters $\mathcal{X}_J$ adjust G on intervals derived from **quantiles** of the Weibull(α, λ)

distribution, obtained from the quantile function $G^{-1}(p|\theta) = (-\lambda \log(1 - p))^{1/\alpha}$. These intervals partition the sample space and are constructed as follows.

Initially, $(0, \infty)$ is split into two pieces, $B_\theta(0) = [0, G^{-1}(0.5|\theta))$ and $B_\theta(1) = [G^{-1}(0.5|\theta), \infty)$. These two sets partition $[0, \infty)$ and have equal probability, 0.5, under the parametric Weibull(α, λ) model. The set $B_\theta(0)$ is further split into two subsets $B_\theta(00) = [0, G^{-1}(0.25|\theta))$ and $B_\theta(01) = [G^{-1}(0.25|\theta), G^{-1}(0.5|\theta))$, each having equal probability, 0.25, under the Weibull(α, λ) model. Similarly $B_\theta(1)$ is split into two subsets $B_\theta(10) = [G^{-1}(0.5|\theta), G^{-1}(0.75|\theta))$ and $B_\theta(11) = [G^{-1}(0.75|\theta), \infty)$, each with Weibull probability 0.25. This defines the first and second partitions of the Polya tree: $B_\theta(0), B_\theta(1) \in \Pi_{\theta 1}$, and $B_\theta(00), B_\theta(01), B_\theta(10), B_\theta(11) \in \Pi_{\theta 2}$. For $\alpha = 2$ and $\lambda = 10$, $B_{(2,10)}(0) = [0, 2.63)$ and $B_{(2,10)}(1) = [2.63, \infty)$ at the first level and $B_{(2,10)}(00) = [0, 1.70)$, $B_{(2,10)}(01) = [1.70, 2.63)$, $B_{(2,10)}(10) = [2.63, 3.72)$, and $B_{(2,10)}(11) = [3.72, \infty)$ at the second.

Nested partitions are sequentially defined in this manner by splitting each set in $\Pi_{\theta j}$, having probability 2^{-j} under Weibull(α, λ), into two sets in $\Pi_{\theta, j+1}$, each having probability $2^{-(j+1)}$. Let $\epsilon_1 \epsilon_2 \dots \epsilon_j$ be a binary expansion that indexes a set $B_\theta(\epsilon_1 \dots \epsilon_j) \in \Pi_{\theta j}$. Then $B_\theta(\epsilon_1 \dots \epsilon_j) = B_\theta(\epsilon_1 \dots \epsilon_j 0) \cup B_\theta(\epsilon_1 \dots \epsilon_j 1)$. Here, $B_\theta(\epsilon_1 \dots \epsilon_j) = [G^{-1}(N/2^j), G^{-1}((N+1)/2^j))$, where N is the decimal representation of $\epsilon_1 \dots \epsilon_j$ and ranges over the integers from 0 to $2^j - 1$.

The probability of an observation T_i falling into $B_\theta(0)$ is denoted $X(0) = G(B_\theta(0))$. Then the probability of T_i falling into the companion set $B_\theta(1)$ is 1 minus this, $X(1) = G(B_\theta(1)) = 1 - G(B_\theta(0))$. Given that $T_i \in B_\theta(\epsilon_1 \dots \epsilon_j)$, the *conditional probability* of T_i falling into $B_\theta(\epsilon_1 \dots \epsilon_j 0)$ versus its companion set $B_\theta(\epsilon_1 \dots \epsilon_j 1)$ is given by $X(\epsilon_1 \dots \epsilon_j 0)$ versus $X(\epsilon_1 \dots \epsilon_j 1)$. That is, for $\epsilon_1 \epsilon_2 \dots \epsilon_j \in \{0, 1\}^j$, $X(\epsilon_1 \dots \epsilon_j 0) = P(T_i \in B_\theta(\epsilon_1 \dots \epsilon_j 0) | T_i \in B_\theta(\epsilon_1 \dots \epsilon_j))$ and $X(\epsilon_1 \dots \epsilon_j 1) = 1 - X(\epsilon_1 \dots \epsilon_j 0) = P(T_i \in B_\theta(\epsilon_1 \dots \epsilon_j 1) | T_i \in B_\theta(\epsilon_1 \dots \epsilon_j))$.

Let $E_J = \bigcup_{j=1}^J \{0, 1\}^j$ denote all binary numbers from 0 to $2^J - 1$ and let $\mathcal{X}_J = \{X(\epsilon) : \epsilon \in E_J\}$ be the set of all conditional probabilities up to level J that define G . A fully specified Polya tree has an infinite number of levels ($J \rightarrow \infty$). A finite Polya tree has J levels. If J is chosen to be large, one can further assume G to be flat on intervals $B_\theta(\epsilon)$ of finite length at level J [2]. Alternatively, one can assume

2 Polya Trees

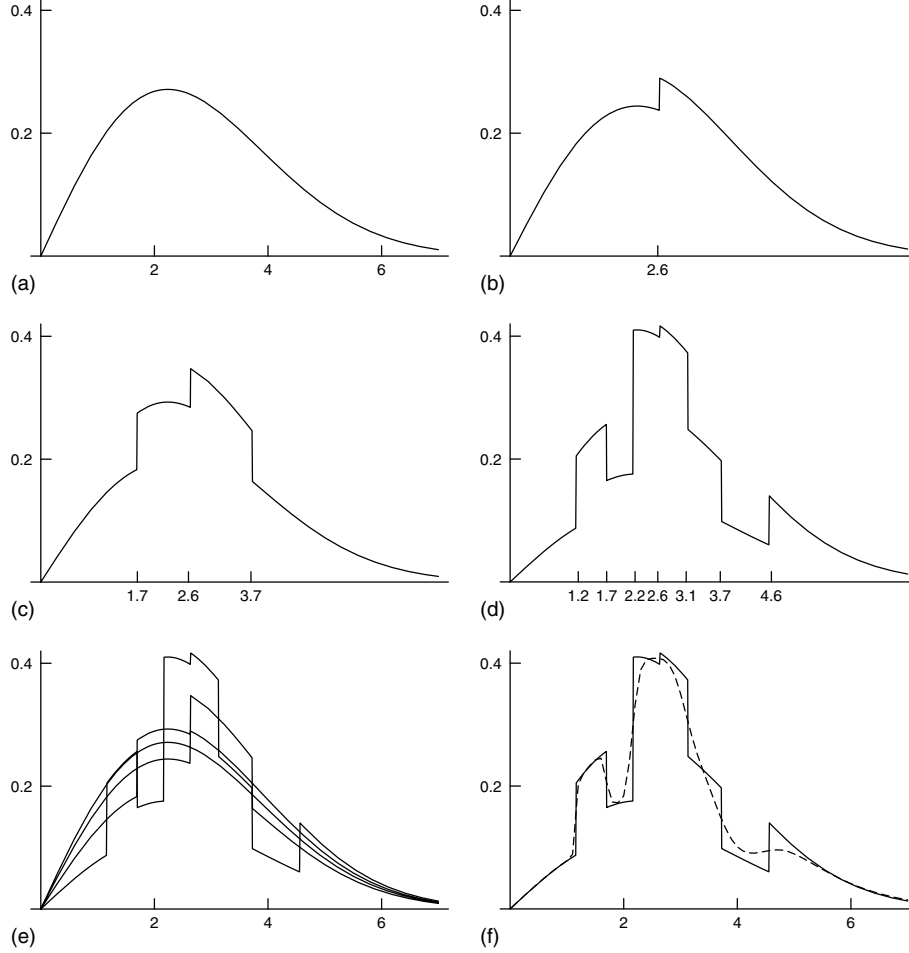


Figure 1 (a) Weibull(2, 10) ($J = 0$) density, (b–d) $J = 1, 2, 3$ respectively, (e) all four levels on same plot, (f) PT and MPT densities

that G follows the Weibull(α, λ) on sets at level J [3] – the approach taken here. When $X(\epsilon) = 0.5$ for all $\epsilon \in E_J$, G is simply the Weibull(α, λ) density we started with.

Let us see how particular choices of conditional probabilities affect the density. Figure 1(a) shows a Weibull(2, 10) density. This is also the density of G governed by a Polya tree but with all conditional probabilities $X(\epsilon_1 \cdots \epsilon_j) = 0.5$. In Figure 1(b), fixing $X(0) = 0.45$ serves to adjust the Weibull density at the first level of the tree on the sets $B_{(2,10)}(0)$ and $B_{(2,10)}(1)$. Figure 1(c) shows the density further refined at the second level from fixing the conditional probabilities $X(00) = 0.4$ and

$X(10) = 0.6$. Figure 1(d) shows a finite Polya tree with three levels obtained by further fixing $X(000) = 0.3$, $X(010) = 0.3$, $X(100) = 0.6$, and $X(110) = 0.3$. Figure 1(e) shows the original density along with the three refinements. Figure 1(f) is obtained by keeping the fixed values in $\mathcal{X}_3$ above, but by considering a mixture of Polya trees (MPT) from taking $\alpha \sim N(2, 0.1^2)$ independent of $\lambda \sim N(10, 0.5^2)$. In the mixture, $E(\alpha) = 2$ and $E(\lambda) = 10$ so the prior on G is centered at the same Weibull(2, 10) density, but by allowing the mixing parameters (α, λ) to be random the partitioning effects (e.g., jagged edges in the density at partition interval endpoints) apparent in the simple Polya tree are smoothed over. This is

because the partition sets Π_{θ_j} are random under the mixture.

Instead of treating the conditional probabilities as fixed, unknown parameters, the Bayesian considers them random. The probability of an observation T_i falling into $B_{\theta}(0)$ is modeled as $X(0) = G(B_{\theta}(0)) \sim \beta(\alpha(0), \alpha(1))$. That is,

$$\begin{bmatrix} X(0) \\ X(1) \end{bmatrix} = \begin{bmatrix} G(B_{\theta}(0)) \\ G(B_{\theta}(1)) \end{bmatrix} \sim \text{Dirichlet} \left(\begin{bmatrix} \alpha(0) \\ \alpha(1) \end{bmatrix} \right) \quad (2)$$

Recall that, in general, for $\epsilon_1 \epsilon_2 \dots \epsilon_j \in \{0, 1\}^j$, $X(\epsilon_1 \dots \epsilon_j 0) = P(T_i \in B_{\theta}(\epsilon_1 \dots \epsilon_j 0) | T_i \in B_{\theta}(\epsilon_1 \dots \epsilon_j))$ and $X(\epsilon_1 \dots \epsilon_j 1) = 1 - X(\epsilon_1 \dots \epsilon_j 0) = P(T_i \in B_{\theta}(\epsilon_1 \dots \epsilon_j 1) | T_i \in B_{\theta}(\epsilon_1 \dots \epsilon_j))$. Rewritten in compact notation, for $\epsilon \in \{0, 1\}^j$, $X(\epsilon 0) = P(T_i \in B_{\theta}(\epsilon 0) | T_i \in B_{\theta}(\epsilon))$ and $X(\epsilon 1) = 1 - X(\epsilon 0) = P(T_i \in B_{\theta}(\epsilon 1) | T_i \in B_{\theta}(\epsilon))$. These pairs of conditional probabilities are also Dirichlet:

$$\begin{bmatrix} X(\epsilon 0) \\ X(\epsilon 1) \end{bmatrix} \sim \text{Dirichlet} \left(\begin{bmatrix} \alpha(\epsilon 0) \\ \alpha(\epsilon 1) \end{bmatrix} \right) \quad (3)$$

So the prior on G is in fact a tree of pairs of conditional probabilities of falling into one set *versus* an adjoining set. These pairs of probabilities are assumed to be independent. Recall that fixing $X(\epsilon) = 0.5$ for all $\epsilon \in E_J$ ensures that G is Weibull(α, λ). By setting $\alpha(\epsilon 0) = \alpha(\epsilon 1)$ for each $\epsilon = \epsilon_1 \dots \epsilon_j \in E_J$, random G is instead *centered* at the Weibull(α, λ) distribution on all partition sets:

$$E\{G(B_{\theta}(\epsilon))\} = \int_{B_{\theta}(\epsilon)} dG(t | \alpha, \lambda) \quad (4)$$

from the fact that $E(X(\epsilon)) = 0.5$ for all $\epsilon \in E_J$. The parameters $\alpha(\epsilon_1 \dots \epsilon_j)$ are often set to $\alpha(\epsilon_1 \dots \epsilon_j) = c_j^2$ for $c > 0$ in applications. This ensures that the conditional probabilities get more highly concentrated around 0.5 as the level J increases, giving larger sets at lower levels more flexibility to adapt to large trends or clumps in the observed data.

The tree structure and independence among probability pairs in $\mathcal{X}_J$ gives the measure of any set

$$G\{B_{\theta}(\epsilon_1 \dots \epsilon_j)\} = \prod_{i=1}^j X(\epsilon_1 \dots \epsilon_i) \quad (5)$$

A key result for Polya trees is that the conditional probability pairs in $\mathcal{X}_J$ given data $\mathbf{T} = (T_1, \dots, T_n)$

are again independent Dirichlet vectors, and are thus easily sampled for a given (α, λ) . This conjugacy result has played a part in several papers that involve Polya tree priors (e.g., [2, 4, 5]) but in many models conjugacy is lost and more sophisticated methods for obtaining posterior inference, such as Markov chain Monte Carlo (MCMC), are employed (*see Markov Chain Monte Carlo, Introduction*).

The definition of a Polya tree can be more complicated than what is outlined so far, but this gives a particular development that has been used with great success in a number of papers. More general Polya trees involve definition of elements $\mathcal{A} = \{\alpha(\epsilon)\}$ one at a time and consideration of different partitions $\{\Pi_1, \dots, \Pi_J\}$. An MPT prior on G assumes $\theta = (\alpha, \lambda)$ is random in addition to the elements in $\mathcal{X}_J$. See [6] and [7] for more details and definitions. Figure 1 illustrates that the MPT prior smooths over partitioning effects apparent in a simple Polya tree; see [6, 8].

Polya Trees and Survival Modeling

The literature on survival modeling with Polya trees can be sorted into two categories: the no-covariates case focusing on the estimation of survival, hazard, and density functions; and semiparametric models that extend these models to include **covariate** information.

Ferguson [9] codifies and fleshes out various results published in the 1960s on Polya trees. This paper also contrasts Polya trees to the special case of the Dirichlet process [1]. Lavine [6, 7] carefully develops and catalogs many additional aspects of Polya tree and MPT theory and methodology; these papers were foundational, and provided core material from which many further developments grew. Lavine [6] analyzes reliability data on the lifetimes of spherical pressure vessels from four different MPT models, each with different sets $\mathcal{A} = \{\alpha(\epsilon)\}$, centered at the exponential family $G(t | \alpha) = 1 - \exp(-\alpha t)$. The data considered in [6] were observed exactly; censored survival data is considered by Muliere and Walker [10] through the construction of Polya trees with some partition points coinciding with censored observations. The resulting estimator reduces to that of Kaplan and Meier [11] as a limiting case. When the partition is fixed ahead of time, i.e. not picked to coincide with censored values, Neath [12] shows that

the posterior distribution for G is an MPT, derives the mixing distribution, and provides a reanalysis of data considered in [10, 11].

These approaches all consider survival data without covariates. Often in reliability settings such as accelerated testing, covariates for the i th experimental unit $\mathbf{x}_i$ are fixed at predetermined values and the resultant lifetime T_i , or possibly a censoring time, is recorded. In survival analysis settings, patient information such as age, treatment, gender, etc., are collected into $\mathbf{x}_i$ and a time to event T_i observed. In either case, the covariates $\mathbf{x}_i$ are linked to T_i through a probability model. In what follows, let $S_i(t) = P(T_i > t)$ be the survival function of subject i with covariates $\mathbf{x}_i$ given all model parameters.

Walker and Mallick [2] develop an accelerated failure time model

$$\log T_i = \mu + \mathbf{x}_i' \boldsymbol{\beta} + \sigma \epsilon_i \quad (6)$$

where the ϵ_i are independent and identically distributed (i.i.d.) from a distribution G assigned a Polya tree prior; this model is also considered in [13]. In applications, $\sigma = 1$ is fixed and the Polya tree is centered at a median-zero **normal distribution** with a large variance. Hanson [8] considers proportional hazards

$$S_i(t) = S_0(t)^{\exp(\mathbf{x}_i' \boldsymbol{\beta})} \quad (7)$$

proportional odds

$$\frac{1 - S_i(t)}{S_i(t)} = \exp(\mathbf{x}_i' \boldsymbol{\beta}) \frac{1 - S_0(t)}{S_0(t)} \quad (8)$$

and a reparameterized version of equation (6),

$$S_i(t) = S_0\{\exp(\mathbf{x}_i' \boldsymbol{\beta})t\} \quad (9)$$

where equations (7–9) assume the same MPT prior centered at the Weibull(α, λ) family for baseline survival S_0 . Mallick and Walker [5] develop a semiparametric transformation model

$$h(T_i) = \mu + \mathbf{x}_i' \boldsymbol{\beta} + \epsilon_i \quad (10)$$

where $h(\cdot)$ is a monotone function modeled as a transformation of a small mixture of beta cumulative distribution function (cdf's). As in equation (6), the ϵ_i are i.i.d. from a distribution G assigned a Polya tree prior. This model reduces to the proportional odds model (8) when G is fixed to be the standard

logistic distribution, reduces to the proportional hazards model (7) when G is extreme value, and reduces to equation (6) when $h(\cdot) = \log(\cdot)$ but with a fixed scale $\sigma = 1$. All of these papers analyzed data on $n = 121$ patients diagnosed with small cell lung cancer, considered in the next section. The accelerated failure time model was found to fit these data better than the proportional odds and proportional hazards models [8, 5].

Hanson and Johnson [3] allow for random σ in equation (6), rather than a simple Polya tree, smoothing out partitioning effects. A frailty version of this model has been used to analyze data on fetal lifetime in cows with parametric frailties accounting for different herd effects [14]. Damien *et al.* [15] considered the analysis of reliability data through equation (9). The proportional odds model (8) was extended by Hanson and Yang [16] to include parametric frailties. Survival data on the lifetimes of $n = 97$ men with advanced, inoperable lung cancer were analyzed by Mallick and Walker [5] and Hanson and Yang [16]. The proportional odds model was picked to be superior in both papers.

Walker and Mallick [4] consider a Polya tree prior on the random effects distribution in generalized linear mixed models, and, in particular, use a Polya tree to model frailties in the Cox model [17]. The use of Polya trees to model random effects in hierarchical models was further developed in [8]. These models are in contrast to the models considered in [5, 14, 16], which considered parametric frailty models coupled with semiparametric survival models. Although common, parametric frailty assumptions are typically made for computational convenience; the MPT model provides a powerful method to relax these strict parametric assumptions in survival and reliability modeling.

In related work, Karabatsos and Walker [18] consider estimation of distributions from several populations that are subject to a stochastic order constraint, and apply their methodology to data comprising the reaction times of schizophrenics and nonschizophrenics.

An Example: Survival Data on Lung Cancer Patients

Data on the treatment of limited-stage small cell lung cancer in $n = 121$ patients were analyzed by Walker

and Mallick [2], Mallick and Walker [5], Hanson [8], and Walker *et al.* [13] among others. In the study, it was of interest to determine which sequencing of the drugs cisplatin and etoposide increased the lifetime from time of diagnosis, measured in days, of those with limited-stage small cell lung cancer. Treatment *A* applied cisplatin followed by etoposide, while treatment *B* applied etoposide followed by cisplatin; $x_{i1} = 0, 1$ for treatments *A, B* respectively. The patients' ages in years at entry into the study, x_{i2} , were also included so $\mathbf{x}_i = (x_{i1}, x_{i2})'$. Treatment *A* was administered to 62 patients, while treatment *B* was administered to 59 patients; 23 patients were administratively right censored.

Here we consider model (9) with an MPT baseline S_0 centered at the Weibull(α, λ) family of distributions. The flat prior $p(\alpha, \lambda, \beta_1, \beta_2) \propto 1$ was assumed. The remaining parameters were fixed at $c = 1$ and $J = 4$ and inference obtained *via* MCMC [8]. The estimated age effect $\hat{\beta}_2$ and 95% credible interval is 0.004 (−0.005, 0.034); for treatment the estimated effect $\hat{\beta}_1$ and 95% credible interval is 0.36 (0.15, 0.71). Treatment *A* is estimated to prolong mean and median life by a factor of $e^{0.36} \approx 1.4$ times, or about 40%, relative to *B*.

The MPT approach models the baseline survival distribution S_0 in tandem with the regression coefficients β , so, arbitrary posterior functionals of (S_0, β) are easily obtained from the MCMC output. For example, one can obtain the hazard ratio for treatment *B versus A* at any time point t , the predictive density at t , arbitrary survival quantiles, etc. Figure 2(b) is a plot of the estimated posterior median hazard ratio for treatment *B versus A* at the mean entry age of 61,

$$hr(t|\text{data}) = \frac{h_{(1,61)}(t)}{h_{(0,61)}(t)}|\text{data} \quad (11)$$

versus time with pointwise 95% credible intervals. Here $[1 - S_x(t)]$ by just $S_x(t)$, the hazard function at t . We see that those under regimen *B* are typically about 1.5 times more likely to die at any instant *versus* those under *A*, but there is considerable variability of the hazard ratio over time. For example, we see peaks under 500 days where those under *B* are estimated to be almost three times as likely to instantaneously expire. In Figure 2(a) we have the two predictive densities at baseline age 61 for treatments *A* and *B*. Note that unlike the Weibull model, the MPT generalization allows for multimodality in these predictive distributions which may have biological interpretations in some applications.

Discussion

This paper introduces the Polya tree prior as a generalization of a parametric family of distributions, and provides a review of literature on the use of Polya trees in reliability and survival analysis. Details of model development and fitting are left to the references. Good references on Polya tree and MPT modeling include [6–9, 13].

One aspect of Polya tree distributions that can be exploited in the modeling of survival data is that the tails of an MPT G retain the parametric flavor of the centering family, essential in some cure-rate models and reliability settings where extrapolation is common. Also, MPT survival models can accommodate arbitrarily censored and truncated data. For example, interval-censored and right-truncated data are easily handled in the proportional hazards model. This is not typically the case with other nonparametric Bayesian survival models.

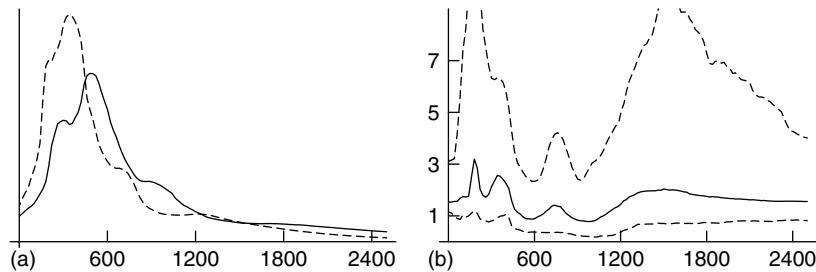


Figure 2 (a) Survival density estimates for treatments *A* (solid) and *B* (dashed) for 61-year-old patients. (b) Hazard ratio (*B vs A*) and pointwise 95% credible intervals

References

- [1] Ferguson, T. (1973). A Bayesian analysis of some non-parametric problems, *Annals of Statistics* **1**, 209–230.
- [2] Walker, S. & Mallick, B. (1999). Semiparametric accelerated life time model, *Biometrics* **55**, 477–483.
- [3] Hanson, T. & Johnson, W. (2002). Modeling regression error with a mixture of PolyA trees, *Journal of the American Statistical Association* **97**, 1020–1033.
- [4] Walker, S. & Mallick, B. (1997). Hierarchical generalized linear models and frailty models with Bayesian nonparametric mixing, *Journal of the Royal Statistical Society. Series B* **59**, 845–860.
- [5] Mallick, B. & Walker, S. (2003). A Bayesian semiparametric transformation model incorporating frailties, *Journal of Statistical Planning and Inference* **112**, 159–174.
- [6] Lavine, M. (1992). Some aspects of PolyA tree distributions for statistical modeling, *Annals of Statistics* **20**, 1222–1235.
- [7] Lavine, M. (1994). More aspects of PolyA tree distributions for statistical modeling, *Annals of Statistics* **22**, 1161–1176.
- [8] Hanson, T. (2006). Inference for mixtures of finite PolyA tree models, *Journal of the American Statistical Association* **101**, 1548–1565.
- [9] Ferguson, T. (1974). Prior distributions on spaces of probability measures, *Annals of Statistics* **2**, 615–629.
- [10] Muliere, P. & Walker, S. (1997). A Bayesian nonparametric approach to survival analysis using PolyA trees, *Scandinavian Journal of Statistics* **24**, 331–340.
- [11] Kaplan, E. & Meier, P. (1958). Nonparametric estimation from incomplete observations, *Journal of the American Statistical Association* **53**, 457–481.
- [12] Neath, A. (2003). PolyA tree distributions for statistical modeling of censored data, *Journal of Applied Mathematics and Decision Sciences* **7**, 175–186.
- [13] Walker, S., Damien, P., Laud, P. & Smith, A. (1999). Bayesian nonparametric inference for random distributions and related functions, *Journal of the Royal Statistical Society. Series B* **61**, 485–527.
- [14] Hanson, T., Bedrick, E., Johnson, W. & Thurmond, M. (2003). A mixture model for bovine abortion and fetal survival, *Statistics in Medicine* **22**, 1725–1739.
- [15] Damien, P., Galenko, A., Popova, E. & Hanson, T. (2007). Bayesian semiparametric analysis for a single item maintenance optimization, *European Journal of Operational Research* **182**, 794–805.
- [16] Hanson, T. and Yang, M. (2007). Bayesian semiparametric proportional odds models, *Biometrics* **63**, 88–95.
- [17] Cox, D. (1972). Regression models and life-tables (with discussion), *Journal of the Royal Statistical Society. Series B* **34**, 187–220.
- [18] Karabatsos, G. & Walker, S. (2007). Bayesian nonparametric inference of stochastically ordered distributions, with PolyA trees and Bernstein polynomials, *Statistics and Probability Letters* **77**, 907–913.

Related Article

Accelerated Life Models; Bayesian Reliability Analysis.

TIMOTHY E. HANSON

Calibration

Introduction

Calibration is the task of establishing common scales/units of measurement for commerce, construction, weather, and scientific study. Some of the earliest properties for which calibration was found useful include weight, mass, volume, purity, distance, area, size, speed, pressure, and temperature. The origin of the word calibration is itself more recent and has its roots in the words caliber and caliper, which initially related to the measurement and manufacturing of firearms that were capable of using bullets of similar size. As measurement devices became more sophisticated, the quantification of contents such as percentages of each component material, bulk properties, trace properties, particle size distributions, and morphology became feasible *via* calibration [1, 2].

Calibration gradually became synonymous with the gradations on scales, thermometers, barometers, and other measurement scales. The success of the industrial revolution rested largely on the ability to produce and use products with standardized properties. As measurement processes continued to become more sophisticated, calibration was then (and is now) used to generally describe the process by which a mathematical relationship is used to systematically relate known measurements (standards) to different, related, but directly measurable properties. This calibration, a mathematical relationship that is statistically estimated, allows the property of interest to be either directly measured or estimated from the indirect measurement(s).

Calibration may be direct or indirect in nature. Calibration is most direct when the measurement scale is in some sense a copy of the standard (a standard defines a known measurement quantity). This occurs with measurement devices such as a ruler, which measures distance directly on its gradations of distance. A mercury thermometer appears to allow the direct measurement of temperature, but instead relies on the manufacturer's calibration relating temperature to the expansion of mercury in a custom tube. The manufacturer has already provided a mechanism for directly reading temperature. A scale that uses springs, appears to directly measure weight, but instead is making use of the relationship between spring compression and the mass placed on the scale

to report a weight. On such devices, calibration has already been performed by the scale manufacturer and these devices may allow a limited degree of user adjustment to the manufacturer's calibration to remove scale bias. If nothing is placed on the scale, a weight of zero is expected. If the scale provides a nonzero result with nothing on it, then an adjustment knob allows the scale to be reset to zero.

For indirect calibration, an instrument measures a response to another variable and this variable can be related to the property of interest through the process of calibration. Most modern analytical instruments such as gas chromatographs (GC), ion-coupled plasma mass spectrometers (ICP-MS), liquid chromatographs (LC), and atmospheric pressure ion mass spectrometers (APIMS) rely on an indirect calibration process. All such instruments require standards to relate what each of these instruments actually measure to the quantified composition of a measured material. For example, with a GC the instrument collects response data over a time domain. Response data from a particular range of time domains can be associated with the response to the concentration of interest. Often the peak area of the data under the curve of the instrument response to the concentration of interest is used as the instrument's response for the purpose of calibration. A GC does not measure concentration directly; rather it measures a signal that can be related to concentration if the relationship between instrument response and concentration can be quantified. This is typically accomplished *via* a mathematical model whose parameter(s) are statistically estimated. Establishing such a quantification of a relationship between instrument response and the measurement of interest is calibration. Most modern measurement instruments rely on indirect calibration (the measuring instrument is not itself a standard). American National Standards Institute (ANSI) defines calibration as the process of determining the deviation of indicated values of an instrument from those of a standard with known uncertainty traceable to the National Institute of Standards and Technology (NIST).

Elements and Styles of Calibration

The elements of calibration include the following:

- an analytical instrument or methodology that serves as a scale;

2 Calibration

- a standard or a series of standards whose properties are known;
- a measurement procedure for using the standard(s) to perform calibration;
- typically a statistical **estimation**/modeling procedure to establish the instrument's scale by relating to how the instrument responds to the measurement of the standard(s); and
- a measurement procedure for the analysis of sample(s).

The process of calibration involves at least several of the following series of statistical concepts:

- accuracy (bias)
- precision
- experimental design
- modeling
- confidence and tolerance intervals
- change detection

Calibration is often performed in two very different modes: research and production. Research-mode calibration is often performed by analytical professionals who are developing and verifying the performance of a measurement method/calibration model. It involves greater effort in that more data than are ultimately required for calibration are collected and evaluated to determine either an optimal or simply a reasonable calibration model/measurement strategy. Production calibration is often performed by laboratory technicians or production personnel using established measurement protocols. Production calibration tends to be minimal in its approach; the calibration model used and the number of measurements utilized are often the minimum required for the task.

Nature of Calibration Models

Many calibration models are univariate in nature, where only one instrument response is being related to the property of interest. Univariate calibration models tend to be monotonic functional forms and tend to have high R^2 values compared to many other general regression modeling situations. The simplest quantitative calibration model is a line that passes through the origin. This model form has only one parameter, the slope. It can be estimated from as little as one measurement from a single standard (single-point calibration, see Figure 1). Although calibration

professionals often chafe at the use of a single standard for the purpose of calibration, even when it is prepared and/or measured multiple times to establish calibration, single-point calibration is often practiced in production environments. Single-point calibration typically relies on a linear model that is assumed to pass through the origin. Use of the origin assumes that when nothing is present or nothing is happening in the system, the measurement instrument will provide a response of zero. Good calibration practice is to verify that one's calibration system either behaves in the assumed manner or if it does not, that the errors introduced are acceptably small. In such a system, the calibration model is defined by the line passing through the two points; the origin and the average response to the single level of calibration standard studied. Although estimation of a single-point calibration model can be done *via* basic geometry, further discussion takes the perspective of calibration models being estimated through some form of regression analysis.

A single-point calibration model is of the form $Y = b \times X$ where Y is the average response, X is the "known" value of the single standard used to establish calibration, and b is the slope of the line. In the traditional regression framework, the "known" is the dependent variable and the variable subject to random measurement error (the response) is illustrated as the independent variable. The issue of how error enters into the calibration standard (the "known") will be discussed later. Expressing the model in this manner is most satisfactory for the statistician since such a model formulation yields the best chance of having an appropriate interpretation of model significance, confidence limits on parameters, and **confidence interval** for the model or an individual point. Such statistical quantification does require the single level of standard to have been measured more than once. This formulation is not convenient for the calibration practitioner. The intent is to use the calibration model to predict estimates of the "known" scale from measured values of response. This requires the inverse of the model, $X = Y/b$. A common calibration shortcut (especially with harder-to-invert calibration model forms) is to reverse the roles of the independent and dependent variables for the purpose of model estimation. While this reversal has *no* impact on the estimated calibration model for single-point calibration, such a variable reversal does impact multipoint calibration.

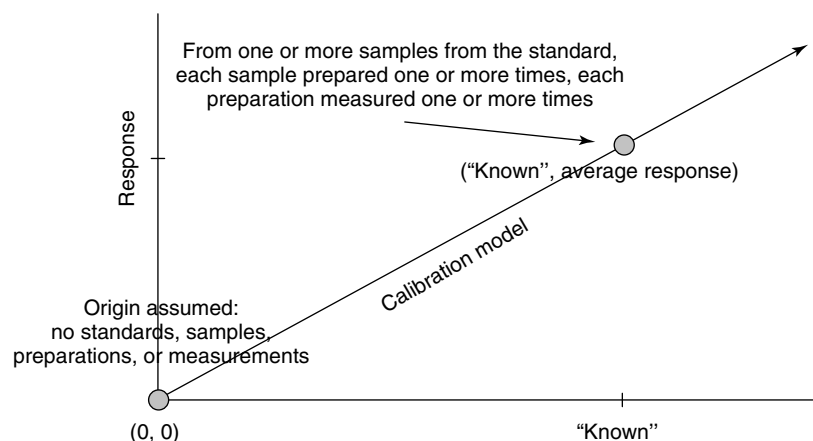


Figure 1 Single-point calibration model

In multipoint calibration, a relationship is estimated from two or more different levels of standards. The calibration relationship may be linear (Figure 2) or nonlinear in nature. When estimating such a model using regression analysis, as long as an equal number of measurements are taken on each combination of each unique level of standard utilized, the *identical* estimated calibration model will result, regardless of whether the model was estimated using the average response at each point or all of the individual measurements as separate points. If the number of measurements is unbalanced this will not be the

case. Even in the balanced data collection case where the same calibration model is estimated regardless of the data level used, the statistical interpretation of the results will differ significantly. The best statistical interpretation will result from performing the regression analysis that does best in terms of meeting the assumptions implicit to the form of regression analysis utilized.

In multipoint calibration models, reversal of the independent and dependent variables will almost always result in a *different* calibration model even when a model as simple as a line is fit. This

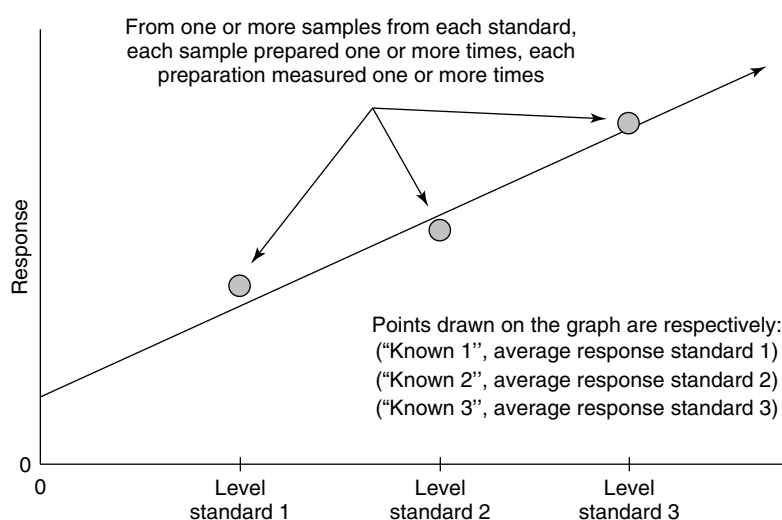


Figure 2 Linear multipoint calibration model

4 Calibration

phenomenon originates from the assumption of the typical regression algorithm that the independent variable(s) are known and the measurement of the dependent variable is subject to random error as well as the impact of the independent variable(s). Models with high explanatory power will be impacted the least by such variable reversal. The reliability of the statistical interpretation of a calibration model is contingent on having met the key assumptions utilized in model fitting. Fitting should be done where assumptions (form of the calibration model and nature of the error model) are best met using a statistically appropriate fitting criterion. Another reason calibration practitioners will often select the measured quantity to take the role of the independent variable in regression analysis (besides simpler inversion) is that any predictions as well as statistical confidence intervals will then be directly expressed in terms of what the calibration model is ultimately targeted at predicting. Such variable reversal has potential ability to corrupt much of the statistical interpretation of a calibration model, *if* doing so violates key assumptions of the statistical analysis methodology used in model estimation.

If a method's measured response is the appropriate dependent variable and the variables are not

reversed, then the statistical confidence intervals are with respect to the dependent variable. To obtain intervals for the independent variable one needs to generate tolerance intervals for the predictions (Figure 3). Specialized regression methodologies that allow for random errors in both the response variable and the levels of the standards used in calibration exist. Such errors-in-variables (EIV) model algorithms are relatively uncommon in laboratory software and are often difficult to use properly with small data collections; even more so with ones that include nonindependently generated working standards. In the classical regression framework in which the errors are implicitly assumed to be only in the dependent variable in the least-squares fitting process, variable reversal for a linear calibration model ultimately has little impact on both the prediction and the confidence intervals for an individual prediction, other than introducing some relatively negligible patterned bias in both. This occurs since linear calibration models generally have high explanatory power.

Instrument manufacturers strive to make their instruments provide as simple a calibration model as possible. Much energy is expended by instrument manufacturers to devise instruments with a linear response over as broad a range as possible. Linear

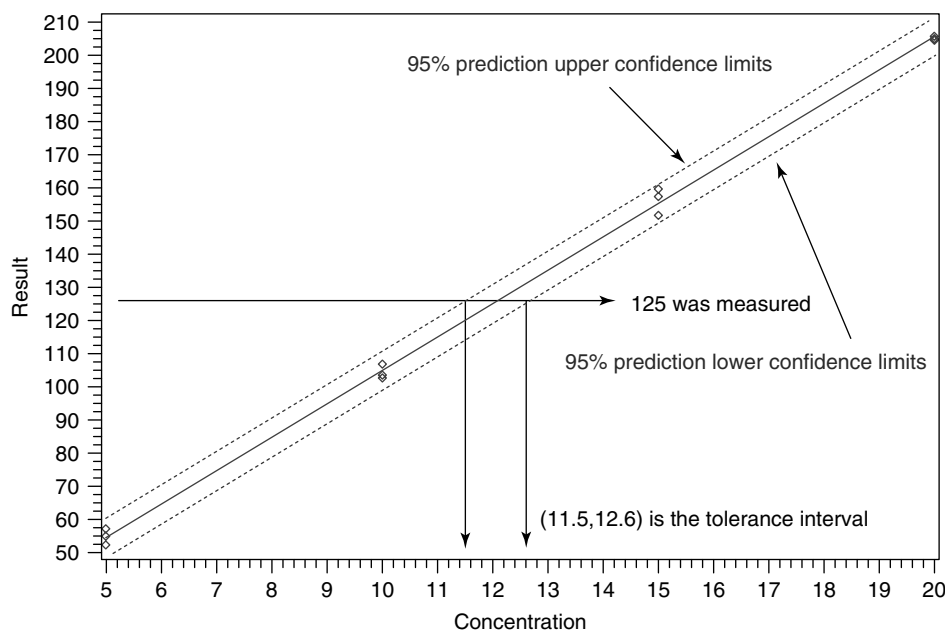


Figure 3 Tolerance interval example

models allow for easy calibration, model inversion, and relatively straightforward model estimation. The assumption of linearity in a calibration model should routinely be checked before using a linear calibration model.

Multivariate Calibration

Although measurement methods and instruments are often designed to be as univariate as feasible, this is not always possible to achieve. A more difficult calibration situation occurs when instrument responses are multivariate rather than a univariate. In such a situation, each instrument response is a joint function of at least two things that the instrument can either measure or be impacted by and can be simultaneously present in samples and standards. Such joint effects are often called *matrix effects*. The most basic tool for dealing with matrix effects is the method of *standard addition*. This is a calibration study that uses a sample with an unknown level in conjunction with one or more samples that have been spiked with known additional amounts of the component of interest. In standard addition, the calibration is typically specific to the sample being measured. If only one spiked sample is utilized, the underlying calibration model is nothing more than a line forced through the origin (single-point calibration equivalent). Attempting more complex model forms through standard addition can be a relatively risky exercise in estimation. If standard addition is not sufficiently reliable to resolve matrix effects, a multivariate calibration model is required [3].

Multivariate models require extra effort in every way. They are more difficult to invert, more difficult to specify, more difficult to name a reasonable error model for, more difficult to obtain appropriate fitting software for, more difficult to get statistical prediction intervals for the quantities of primary interest, require more calibration standards that may also be more difficult to manufacture, and typically require use of DOE (design of experiments) methodologies to obtain efficient calibration sets. Most violations of assumptions in a multivariate modeling situation add patterned bias to the calibration model's predictions. A particular aspect of the multivariate calibration model is that being inside the multivariate region bounded by the lower and upper limits for each of the related variables contained in the standard does

not a guarantee that a prediction did not originate from a "hidden extrapolation" of the model. As with all prediction methodologies, model extrapolations of calibration models typically yield less-reliable forecasts than model interpolations.

Fitting Criteria

Fitting criteria, the version of **least squares** that is used to estimate a calibration model is often critical. For some methods, the traditional "ordinary" least-squares (OLS) approach used in elementary statistics texts is sufficient. For some instruments, a weighted least-squares (WLS) approach will provide more reliable calibration models and estimates. For some (often multivariate) calibration problems, partial least squares (PLS) can be useful. For many instruments used in the measurement of trace contaminant levels, WLS analysis is more appropriate than OLS analysis [4, 5].

Checking Calibration

Calibration is typically checked in a variety of ways. One way is to independently make or obtain standards that are then not used in the process of calibration. These independent standards are then measured as samples and these measurements *versus* the known values serve as a check of the reliability of the calibration. A *recovery study* is another means by which calibration can be checked. Recovery studies involve adding a known amount of a material that also can be present in a given sample and measuring both the sample and the "spiked" sample and observing by differencing how well the measurement method recovers the value of the known spike. When multiple measurements and samples/standards are involved in such studies summary statistics such as the mean and standard deviation are used sometimes in conjunction with *t* tests, *F* tests, and ANOVA (analysis of variance) methodologies.

When calibration is an issue between multiple laboratories, a common statistical tool is a round-robin study. In a typical **round robin**, each laboratory measures the "same" samples some number of times each while possibly also doing this for multiple operators, multiple dilutions, multiple extractions, multiple preconcentrations or any other key processing steps,

and/or multiple instruments from within each laboratory. Round-robin studies are typically focused on how well different laboratories can measure the same thing. DOE methodologies are often used to build such studies and ANOVA techniques are frequently used to analyze such data.

Working Standards and Traceability

Laboratories frequently need to rework an out-of-range primary standard into the working standards that span the range of interest of their applications for calibration. For example, NIST may only provide a 1-ppm standard for the constituent in question. The practitioner may require a 1-, 2-, 5-, 10-, and 20-ppb standard for their calibration process. Frequently, such working standard manufacturing processes may involve another dilution or even a sequence of multiple dilutions. There is a great deal of statistical science involved with how to optimally perform a given dilution in terms of the best number of dilution steps required *versus* the accuracy and precision characteristics of the dilution process. A common labor-saving shortcut used by practitioners for some multistep dilution sequences is to do a serial dilution (each dilution is based on the prior dilution and not by starting the dilution sequence with the primary standard again). This strongly impacts the resulting error structure in that any precision error in the first dilution step shows up as an accuracy error in any subsequent dilution step. Dilution is not the only way in which working standards derived from primary standards are prepared; in some instances preconcentration may be required.

Calibration standards also may or may not have traceability. Traceability is the process that leads back to the primary standard and this standard's value and estimated reliability. If a working standard is traceable to a primary standard, then it may be possible to quantify the likely reliability of such a working standard through appropriate application of techniques such as simulation or propagation of errors.

Limit of Detection Concepts

Another key calibration-related concept is that of limit of detection (LOD). LOD is an important figure of merit for a measurement process when the measurement process may be utilized in a range

of values close enough to zero to provide relatively unreliable measurements. There are many other alternative acronyms for LOD or LOD-related concepts including LLOD (lower limit of detection), LDL (lower detectable limit), LOGD (limit of guaranteed detection), MDL (method detection limit), and LOQ (limit of quantification). Many LOD concepts estimate the smallest value that is statistically distinguishable from zero. Numerous organizations and technical publications have defined many differing techniques for LOD determination and for some techniques there are published variants to address perceived weaknesses [6, 7, 8, 9].

For the consumers of the data produced from calibration processes, the most useful version of a detection limit is an MDL although versions of detection limits that focus only on a detector's reliability (example: detector signal to noise definitions) or other parts of a measurement method can be useful for a measurement professional to understand as well. Good MDL definitions include the variability from the calibration process as part of the definition. This is because the reliability of any measured result that has inherently utilized calibration also has been impacted by the uncertainty introduced by the calibration process itself. Some examples of good MDL definitions in this respect include the ISO and the SEMI MDL definitions [10].

Limitations of Statistical Toolsets

Statistical tools, in current common usage in calibration practice, do not directly incorporate many of the characteristics inherent to common calibration processes. Commonly used statistical methods do not distinguish between standards that have been manufactured independently from those that have been manufactured from another standard in a serial fashion. Commonly used regression or other fitting tools do not fully incorporate the error hierarchy that takes place for a given measurement. For example, a standard may have been sampled once, prepared once, and this sample/preparation may have been measured four times. This is a very different error structure/statistical scenario from a standard having been sampled four times, with each of these four samples having been independently prepared and measured once each. The former scenario requires much less effort to obtain the four measurements than the

latter. Such important distinctions are often routinely ignored in subsequent modeling of calibration data. Much further work needs to be done in this area to develop and refine statistical tools that appropriately integrate the substructure of typical calibration processes and get these tools, with appropriate training, into the hands of analytical professionals.

Other Issues

Some of the other areas in calibration in which statistical tools or comparisons are particularly useful include tasks such as instrument tuning, selection of the best instrument response or responses to use in calibration, calibration model identification algorithms, signal averaging, drift correction, and the evaluation of the state of calibration, or determining the optimal frequency or strategy of when to recalibrate. Analytical professionals have also adopted many specific calibration referencing strategies such as the use of internal standards or external standards as part of a calibration process to have the capability of statistically removing certain systematic errors that may impact a given measurement or a group of measurements. Other bias-removal techniques are based on the use of DOE strategies [11].

References

- [1] Dietrich, C.F. (1991). *Uncertainty, Calibration and Probability: The Statistics of Scientific and Industrial Measurement*, 2nd Edition, Adam Hilger, pp. 1–5.
- [2] Sharaf, M.A., Illman, D.L. & Kowalski, B.R. (1986). *Chemometrics, Chapter 4: Calibration and Chemical Analysis*, John Wiley & Sons.
- [3] Martens, H. & Naes, T. (1989). *Multivariate Calibration, Chapter 1: Introduction to Multivariate Calibration*, John Wiley & Sons.
- [4] Meier, P.C. & Zund, R.E. (1993). *Statistical Methods in Analytical Chemistry*, John Wiley & Sons.
- [5] Ketkar, S.N. & Bzik, T.J. (2000). Calibration of analytical instrument – impact of nonconstant variance in calibration data, *Journal of Analytical Chemistry* **72**(19), 4762–4765.
- [6] Hogan, J.D. (ed) (1997). *Specialty Gas Analysis: A Practical Guidebook, Chapter 8: Limit of Detection*, John Wiley & Sons, pp. 163–190.
- [7] Bzik, T.J. (2000). Method detection limit estimation theory and practice, in *Proceedings of the CC session of the IEST ATM* Providence.
- [8] ISO 11843-1, Capability of Detection - Part 1: Terms and definitions, ISO, Genève, (1997).
- [9] ISO 11843-2, Capability of Detection - Part 2: Methodology in the linear calibration case, ISO, Genève, (2000).
- [10] SEMI International, SEMI C10 - Guide for Determination of Method Detection Limits, SEMI International Standards, Process Chemicals (or Gases), (2003).
- [11] Bzik, T.J., Henderson, P.B. & Hobbs, J.P. (1998). Increasing the precision and accuracy of top-loading balances: application of experimental design, *Journal of Analytical Chemistry* **70**(1), 58–63.

THOMAS J. BZIK

Kano Analysis

Overview

In the early 1980s, Dr. Noriaki Kano of Tokyo Rika University introduced several concepts aimed at understanding quality better as the customer defines it. They were intended for application in **quality function deployment** (QFD), but can also be used as stand-alone concepts to help companies improve their quality, and respond to market opportunities and needs.

One of these concepts is called *attractive quality* and *must-be quality* [1]. It describes the continuum over which a customer values a particular feature or attribute of a product or service.

Attractive quality is a feature that the customer does not necessarily expect but is delighted when it is present. When properly implemented, attractive-quality features typically add to a company's market share, depending on the added cost to the customer. In 2007, an automobile with a built-in GPS (global positions system) was introduced and this feature would fall on the attractive-quality side of the continuum.

Must-be quality is a feature that the customer expects to be present and will not pay extra for it. If a must-be-quality feature is not present or is poorly implemented, the company can incur serious customer dissatisfaction and possibly lose market share. Today, antilock brakes on cars would fall on the must-be-quality side of the continuum.

Dr. Kano points out that the position of any quality feature is not static on this scale over time. What starts as an attractive quality, a delight to customers when it is present, often moves to be a must-be quality over time as more companies provide this feature routinely. A good example of this is air bags on automobiles. Originally an attractive-quality feature, air bags first appeared on just the driver's side; soon all manufacturers had them. Then they were added to the passenger's side, and now all new cars have them there, too – a reflection of their must-be status. More recently, side air bags have been added, and so forth. Similarly, cameras on cell phones have moved rapidly over time from attractive quality to nearly must-be quality.

Kano analysis is also commonly called the *Kano model*. This incorporates additional analyses of quality from a customer's perspective and has been actively studied and applied [2].

References

- [1] Kano, N. (1984). Attractive quality and must-be quality, *The Journal of the Japanese Society of Quality Control* **1**, 39–48.
- [2] Walden, D. ed, (1993). Kano's methods for understanding customer-defined quality, special issue of *The Center of Quality Management Journal* **2**(4), 37.

DEBRA SHENK

Change Management, Stakeholder Assessment in

Introduction

The stakeholder assessment process is used by project leaders in change management to understand each stakeholder's level of influence, authority, and control in the organization as well as his level of involvement and commitment to support a particular project. The process is used by those leading projects that change processes, systems, equipment, software, and facilities. The goal of change management is to ensure that the organization is fully prepared to function once the project is completed and turned over to the organization to operate. Organizational changes may include new or revised roles and reporting relationships, staffing levels, facilities, procedures, education and training, and workplace culture [1, 2].

Stakeholder Definition

For this article on change management, "stakeholder" is defined as an individual or group of individuals with an interest (stake), something to be gained or something to be lost, in the outcome of the change initiative and whose interests must be recognized for the project to be successful. Freeman and McVea propose that by adopting a view of stakeholders as real people with names and faces rather than roles, project leaders can understand how collections of individuals can work together in organizations to create value that benefits all those who are affected [3].

The types of stakeholders can be both internal and external to the organization. They include customers, suppliers, end users, project team members, leaders, and functional groups having a governing responsibility such as, environment, health and safety, legal, procurement, facilities, logistics, and information technology.

Stakeholder Assessment Objective

The objective of stakeholder assessment is to ensure success of an organizational change initiative by identifying, gaining, and keeping the appropriate

levels of support and involvement of all individuals and groups having a stake in the project's outcome.

Stakeholder Assessment Process

There are two components to the assessment. First, the project leader collects data assessing each stakeholder's level of involvement and commitment to the project. Stakeholder interviews reveal whether they will most likely oppose, allow, assist, perform, or even sponsor the project as described in Table 1. Second, an assessment is made whether they will exert a high, medium, or low level of influence, authority, or control over the project. The project leader will plot the resulting assessment of each stakeholder on the matrix in Figure 1.

Stakeholder Involvement and Commitment

The benefit of understanding a stakeholder's level of support for a particular project is to (a) identify those in the organization who can support the project by providing resources and/or performing work and (b) understand early on where the project may face serious challenges from those who do not support the project and may work against the project team's efforts. With this understanding, project leaders will be able to formulate a strategy and take action to minimize the effect of those who oppose, and leverage those who support to expedite the project.

High					
Influence Authority Control					
Low					
	Oppose	Allow	Assist	Perform	Sponsor
	Involvement & Commitment				

Figure 1 Influence and Involvement Matrix

2 Change Management, Stakeholder Assessment in

Table 1 Stakeholder Involvement and Commitment Descriptions

Oppose	This stakeholder clearly opposes the project and will attempt to block an action taken by the project team.
Allow	This stakeholder may not agree with the project, but, will allow the project to proceed. He may not directly support it.
Assist	This stakeholder agrees with the project and while not taking a leadership role, will help in some way to ensure the success of the project.
Perform	This stakeholder usually has a strong interest in the project's success and will take responsibility for a significant portion of the project.
Sponsor	This stakeholder is required on every project and will provide the resources, leadership, and visibility at the top of the organization necessary for the success of the project.

Using Table 1, project leaders will select the category that best represents each stakeholder.

Stakeholder Influence, Authority, and Control

The second assessment is to understand a stakeholder's ability to affect the behavior of others in the organization through the control of information, expertise, resources, or authority. This is a key understanding in that the stakeholder's influence can impact the project in a favorable or unfavorable manner.

Stakeholder Assessment Summary

The project leader will record the results of the stakeholder assessment in a format similar to the one shown in Table 2. Each entry identifies what is at stake including what each expects to gain (or lose) once the project is completed.

Influence and Involvement Matrix

Matrix Objective

The objective of the matrix is to provide a graphic representation of the stakeholder assessment results

and from this to develop strategies to increase support for the project and remove the barriers to the project's success.

Project Sponsor

If the project sponsor has not been identified yet, the first step is to identify one. This step is very important and a project's likelihood for success is extremely low without an effective project sponsor.

The project sponsor is accountable for the project's business benefits realization and has macro-level scope control. The project sponsor will also ensure the project is appropriately resourced to achieve the intended results and will undertake regular project reviews and make necessary adjustments. The project sponsor may also hold the executive role as well. Staying connected to the executive or top governing level in the organization is necessary to ensure the project has the appropriate level of visibility and support in the organization to be successful. If the project sponsor does not hold the executive role, he will need to periodically inform the executive of the progress the project is making against the intended objectives.

Table 2 Stakeholder Assessment Summary

[illegible]

To identify a project sponsor, project leaders look for stakeholders with a high level of influence and involvement, someone who believes in the value of the project and is willing to use influence to ensure that the project is completed successfully. The matrix can be used in discussions with potential sponsors to illustrate the other stakeholders' positions on the project.

Influence and Involvement Gaps

Project leaders next look for gaps on the matrix where there is insufficient support. These influence and involvement gaps can occur where a high level of influence and low level of involvement exists. In this case, the stakeholder has power in the organization but does not currently support the project. The goal is to develop strategies to increase support for the project or remove the barriers. This can be accomplished by moving "high-influence" and "low-involvement" stakeholders to "high-influence" and "high-involvement" positions. If these stakeholders understand how the project will benefit them, they will use their influence to make it happen. However, if they continue to oppose the project, the project leader's goal is to lessen their influence by moving

them to a "low-influence" and "low-involvement" position on the matrix. This may be accomplished through intervention by the project sponsor or senior leaders of the organization.

Additionally, insufficient support occurs where a low level of influence and high level of involvement exists. These stakeholders are very supportive of the project yet lack the influencing ability to push the project along. Increasing their influence may be accomplished by involving them in governing the project, perhaps through steering committees and enhanced project roles.

Stakeholder Action Plan

With each stakeholder's current position located on the matrix, and gaps in influence and commitment identified, the project leader will develop a strategy and action plan to resolve issues and reposition stakeholders to the desired locations. Project leaders will record the action plan in a table similar to Table 3.

Table 4 provides a summary of the stakeholder assessment process and is used by project leaders as a "to-do" list to communicate with others and to keep the assessment process moving forward.

Table 3 Stakeholder Action Plan

Stakeholder	Current Position	Desired Position	Action Plan to Resolve Issues and Reposition Stakeholders

Table 4 Stakeholder Assessment Process Summary

1. Identify individuals or groups that will be most affected by the outcome of the project and list them as stakeholders.
2. Determine their current level of involvement and commitment.
3. Determine their current level of influence, authority, and control.
4. Locate each stakeholder's current position on the influence and involvement matrix.
5. Identify the interests and outstanding issues associated with each stakeholder.
6. Identify the desired location for the stakeholder and draw an arrow from their current location to the desired location.
7. Create a strategy and action plan to resolve the issues and reposition stakeholders by changing their influence and/or involvement.

Summary

Projects, by their nature, create change within organizations. Some changes are subtle and others monumental. Understanding how those who have a stake in the outcome are involved enables project leaders to plan and implement strategies for increasing support and reducing barriers to project success.

Bridges clarifies the difference between change and transition. “*Change* is situational: the new site, the new boss, the new team roles, the new policy. *Transition* is the psychological process people go through to come to terms with the new situation. Change is external, transition is internal.” Using stakeholder assessment is a way in which project leaders may favorably influence organizational changes by involving key individuals, with names and faces, in the change process [4].

References

- [1] Sexty, R.W. (2004). *The Business Environment: Identifying Stakeholders and Issues* (Website: <http://www.uccs.mun.ca/~rsexty/business8107/StakeholderConcept.htm>).
- [2] Freeman, R.R. (1984). *Strategic Management: A Stakeholder Approach*, Pitman, Boston.
- [3] McVea, J.F. & Freeman, R.R. (2005). A names and faces approach to stakeholder management: how focusing on stakeholders as individuals can bring ethics and entrepreneurial strategy together, *Journal of Management Inquiry* **14**(1), 57–69, (Web site:<http://jmi.sagepub.com/cgi/content/refs/14/1/57>).
- [4] Bridges, W. (1991). *Managing Transitions, Making the Most of Change*, Addison Wesley.

THOMAS F. TRABERT

Policy Deployment for Performance Improvement

Introduction

Organizations that win in the long term “plan their work and work their plan”. Realization of strategy – the long-term vision of an organization is achieved by a disciplined approach to setting direction and then executing that direction through the effective use of the organization’s resources. In Japan this method is called *policy deployment* – which has also been called the *secret weapon* in the Japanese management system. Policy deployment is the strategic direction-setting methodology used to identify business goals, as well as to formulate and deploy major change management projects throughout an organization. It describes how strategy cascades from vision to execution in the workplace through a collaborative engagement process that also includes implementation details like performance, self-assessment, and management review. It describes a systematic relationship between strategy development and the organization’s daily imperative to measure and manage its operations using a system that aligns the actions of its people to produce collaboration among the various business functions and processes to produce requirements for customers.

Historical Development of Policy Deployment

What were the circumstances under which policy deployment originated? Interest in strategy, market focus, and long-term, balanced planning were generated by the visits of Dr. Peter F. Drucker to Japan in the early 1950s [1]. As a result of his teaching, “policy and planning” was added to the Deming Prize checklist in 1958. Bridgestone Tire Corporation first used *hoshin kanri*, the Japanese term for policy deployment, in 1965. In 1976, Dr. Yoji Akao and Dr. Shigeru Mizuno were involved in the implementation of *hoshin kanri* in Yokagawa Hewlett–Packard (YHP) as part of its pursuit of the Deming Prize. By 1982, YHP had used *hoshin* to manage a strategic change that moved it from the

least profitable Hewlett–Packard (HP) division to the most profitable. In 1985, this *hoshin* methodology was introduced to the rest of HP as a lesson learned from the YHP Deming Prize journey. From HP this methodology was transferred to other leading companies including Proctor & Gamble, Ford, Xerox, and Florida Power & Light, involving several advisors and councilors of the Union of Japanese Scientists and Engineers (JUSE). The work of the GOAL/QPC research committee also extended the managerial technology of policy deployment and was a key ingredient in introducing policy deployment across North America and through multinational companies, to the world [2].

Foundations of Policy Deployment

Mizuno defined *hoshin kanri* as the process for “deploying and sharing the direction, goals, and approaches of corporate management from top management to employees, and for each unit of the organization to conduct work according to the plan”. *Hoshin kanri* is a comprehensive, closed-loop management planning, objectives deployment, and operational review process that coordinates activities to achieve desired strategic objectives. The word *hoshin* refers to the long-range strategic direction that anticipates competitive developments while the word *kanri* refers to a control system for managing the process [2].

This management system does not encourage random business improvement, but focuses the organization on projects that move it toward its strategic direction. It builds strength from its relationship with the daily management system that is focused on *kaizen* – continuous improvement. *Hoshin* seeks breakthrough improvement in business processes by allocating strategic business resources (both financial and human resources) to projects that balance short-term business performance to sustain improvement toward its long-term objectives. In a policy deployment management system this two-pronged approach integrates operational excellence in the daily management system with architectural design of its long-term future. This planning process contains two objectives: *hoshin* – the long-range planning objectives for strategic change that allow an organization to achieve its vision and *nichijo kanri* – the daily, routine management control system (or daily management system) that translates the strategic objectives into the

2 Policy Deployment for Performance Improvement

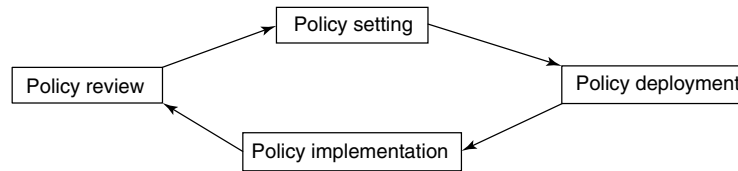


Figure 1 The system of hoshin kanri of managing policy

work that must be accomplished for an organization to fulfill its mission. The blending of these two elements into a consensus management process to achieve a shared purpose is the key to success for the policy deployment process. In a *hoshin* planning system, strategy is observed through the persistence of its vision – how it is deployed across cycles of learning in project improvement projects that move the performance of the organization’s daily management system toward its direction of desired progress.

The fundamental premise in policy deployment is that the best way to obtain desired results for an organization is for all employees to understand the long-range direction and participate in designing the practical steps to achieve the results. This form of participative management evolved and was influenced by the Japanese refinement of Drucker’s management by objectives (MBO) through the birth and growth of the quality circle movement. To manage their workplace effectively, workers must have measures of their processes and monitor these measures to assure that they are contributing to continuous improvement as well as closing the gap toward the strategic targets. Policy deployment became the tool that Japanese business leaders used to engage their workers in a strategic dialog and align their work with the consensus strategic direction of their firm. When HP first implemented *hoshin* planning, many of its business leaders explained how it worked by calling it “turbo-MBO”.

Policy deployment links breakthrough projects that deliver the long-term strategic direction to achieve sustainable business strength while, at the same time, delivering an operating plan to achieve short-term performance. The methods of policy deployment anticipate long-term requirements by focusing on annual plans and actions that must be met each year to accumulate into long-term strength. Policy deployment processes begin when senior management identifies the key issues or statements of vulnerability, where improvement will have its greatest impact on business performance. This perspective

is an essential starting point for policy deployment and allows management to focus a strategic dialog to solicit ideas from frontline workers regarding the opportunities for improvement that exist in their workplace. As Dr. Noriaki Kano of the Tokyo Science University points out, without such direction “the ship would be rudderless”. The communication of the focus area or theme for improvement provides a cohesive direction to assure alignment of the entire organization and to build consensus between the management team and the workers on business priorities.

Hoshin helps to create the type of organization that William McKnight, former CEO of 3M, expressed as his desire: “an *organization* that would continually self-mutate from within impelled forward by employees exercising their individual initiative” [3]. In this type of organization, creativity is managed through a combination of self-initiated improvement projects, which engage teams to combine individual capabilities for achieving strategic projects that make a difference to the whole organization. How does this change management process work at the frontline where these strategic *hoshin* projects interface with the organization’s routine work processes?

Perhaps the reason *hoshin kanri* took hold within HP is that this methodology demonstrated its ability to translate qualitative, directional, or strategic goals of an organization into quantitative, achievable actions that focus on fundamental business priorities achieving significant competitive breakthroughs – in short, the leaders at HP recognized *hoshin kanri* as MBO done right!^a The extension of this methodology beyond HP to other leading firms came about because HP was recognized as possessing the best practice for linking its strategic direction with its operational management systems.

Daily Management System

Policy deployment uses a systems approach to manage organization-wide improvement of key business

processes. It combines the efforts of focused teams on breakthrough projects with the efforts of intact work groups who continuously improve the performance of their work processes. All strategic change occurs in projects that accomplish those activities that are necessary to achieve the stretch business objectives that assure sustained success for the organization. Policy deployment systematically plans ways to link strategic direction with those business fundamentals that are required to run the business routine successfully. Policy deployment allows management to commission change projects for implementation and to review the implementation of a system of projects and thereby to manage change. It seeks opportunities disguised as problems – and elevates those high-priority changes required for the improvement of the daily management system and work processes into business change objectives that are accomplished as *hoshin* projects.

Routine operation of the daily management system requires a foundation in management by fact, or the combination of business measurement with statistical analysis and graphical reports that illustrates the current state of performance, historical trends, and is able to extrapolate trends through statistical inference. A key ingredient is the business fundamentals measurement system that includes the set of basic process results measures that are monitored at control points within the organization where the flow of its throughput can be managed based on the requirements that are driven (using a pull system) by the customer requirements. This measurement system should include both predictive and diagnostic capabilities.

HP embedded its daily management system into a work process measurement system that they initially called *business fundamentals tables*. Other companies refer to the set of measures that translate strategic goals into operational measures of work (in units such as quality, cost, and time) as either a customer dashboard or a balanced scorecard. These systems are used to monitor the daily operations of a business and to report to the management on the progress in the process for developing and delivering value to customers. This measurement process must operate in close-to-real time to permit process owners to take appropriate corrective action that will limit the “escapes” of defects to external customers by catching and correcting errors before they are released from the organization, and finding and fixing

mistakes as they occur at the source. Such measures of core work processes are called *business fundamentals* because they must operate under control for the business to achieve its fundamental performance objectives.

These measures must also be captured at the point where control may be exercised by process operators to adjust the real-time operation of the process and assure meeting the customer’s performance requirements on a continuing basis. As the great Dutch architect Miles van der Rohe once observed “God is in the details” and it is in these details that business must effectively operate. A daily management system defines the details of an organization’s operations. Thus, the measurement and the point at which it is both monitored and controlled are parts of the daily management system and at this point they must be related to their contribution to deliver organizational performance objectives. In the language of Six Sigma, a “Business *Y*” (such as “profitable growth”) that must be achieved is the strategic goal, while a “Process *X*” (such as “creditworthy customers”) delivers this result by process-level performance through the **transfer function** $Y = f(X)$ and the “*X*” is therefore a fundamental business measure of the organization’s daily management system.

Collins and Porras describe how leading companies stimulate improvement by *evolutionary progress*. “Evolutionary” describes progress that resembles the organic growth or the way that species evolve and adapt to their natural environments. Evolutionary progress differs from the big hairy audacious goals (BHAG) of strategic progress in two ways. First, whereas BHAG progress involves clear and unambiguous goals (“We are going to climb *that* mountain”), evolutionary progress involves ambiguity (“By trying lots of different approaches, we’re bound to stumble onto something that works; we just don’t know ahead of time what it will be”). Second, whereas a BHAG involves bold discontinuous leaps, evolutionary progress begins with small *incremental* steps or mutations, often in the form of quickly seizing unexpected opportunities that eventually grow into major – and often unanticipated – strategic shifts. Evolutionary progress represents a means to take advantage of unplanned opportunities for improvement that are observed at the point of application – the daily management system. The accumulation of many evolutionary improvements results in what looks like part of a brilliant overall

4 Policy Deployment for Performance Improvement

strategic plan [4]. Both types of change are needed to stimulate the organic growth of an enterprise. If an organization can make improvements in the “right X’s” then it will improve its performance on the critical Business Y.

Choosing Strategic Direction

Hoshin kanri begins with a process for choosing strategic change. In most firms, this process is called strategic planning. Proposed changes are usually identified to either increase the competitive performance of a process or to create the competitive “attractiveness” of a product to its targeted market. Strategic choice in both dimensions is essential to have a globally competitive organization. As pointed out by Dr. Hiroshi Osada, many Japanese companies have not paid enough attention to the critical aspects of strategy formulation as they have to the deployment of their strategy using *hoshin kanri*. This leads to an error of effectively deploying a poorly chosen strategy. When management confuses the mechanistic aspects of policy deployment with its own crucial obligation to establish strategic direction, then they create a grievous error that is truly an abrogation of leadership. An organization may effectively deploy management’s strategic choice, however, if the choice of strategy is not carefully directed it will not lead to improvement [5].

Osada notes that in traditional Japanese management systems, ideas flow “bottom-up” from the workplace to the management. However, in policy deployment, there is also a top-down approach to planning change. As Osada comments, policy deployment “is a simple tool for effectively deploying a given policy, and has therefore been broadly adopted by Japanese industry, it does not aid in policy formulation. Even when employing management by policy (MBP), therefore, the question of whether or not a given policy is appropriate will remain. It is thus possible for an inappropriate policy to be effectively deployed – to a counterproductive effect” [6]. Strategic direction must be determined by discovering the alternatives for achieving the organization’s vision and choosing the direction that will accomplish it. This direction is modified through the power of the incremental change to act as the “rudder” that steers the ship by making “finely tuned” changes to the general direction of the strategy.

What are the essential ingredients in choosing strategic direction? This process of management integrates strategic planning, change management, and project management with the performance management methods that focus on delivering results. Some specific work activities in designing and implementing a policy deployment system include the following:

- identifying critical business assumptions and areas of vulnerability,
- identifying specific opportunities for improvement,
- establishing business objectives to address the most imperative issues,
- setting performance improvement goals for the organization,
- developing change management strategies for addressing business objectives,
- defining goals project charters for implementing each change strategy,
- creating operational definitions of performance measures for key business processes, and
- defining business fundamental measures for all subprocesses to the working level.

Once a strategy has been set, the next challenge is to align the strategic direction with the work that is being performed in the daily management system.

Aligning Operations with Strategies

A critical challenge for organizations is to align their strategic direction with their daily work systems so that they work in concert to achieve the desired state. Alignment must include linking cultural practices, strategies, tactics, organization systems, structure, pay and incentive systems, building layout, accounting systems, job design, and measurement systems – *everything*. In short, alignment means that all elements of the company work together much like an orchestra leader integrates the various instruments to conduct a coordinated symphony. Organizations that apply the most mature aspects of policy deployment do not put in place any random mechanisms or processes, but they make careful, reasoned strategic choices that reinforce each other and achieve synergy. These organizations will “obliterate misalignments”. If you evaluate your company’s systems, you can probably identify some

specific items that have misaligned with its vision and impede its progress. These “inappropriate” practices have been maintained over time and have not been abandoned when they no longer align with the organizational purpose. “Does the incentive system reward behaviors inconsistent with your core values? Does the organization’s structure get in the way of progress? Do goals and strategies drive the company away from its basic purpose? Do corporate policies inhibit change and improvement? Does the office and building layout stifle progress? Attaining alignment is not just a process of adding new things; it is also a never-ending process of identifying and doggedly correcting misalignments that push a company away from its core ideology or impede progress” [7].

The System for Policy Management

Policy deployment combines both the “target and the means to achieve the target” into a consensus-generating, management decision-making process. Improvement targets are described using four elements: a performance measure to be improved, direction and rate of improvement desired, targeted improvement magnitude, and timeframe for achievement of the target. A means to achieve the target describes a set of specific actions that will be taken to deliver the desired results. These means may differ across the organization, based upon the initial, local management self-assessment, or “current state analysis” that is conducted to assess the business area’s starting point for change and determine the magnitude and nature of the performance gap to be closed by the change management or *hoshin* project to deliver the desired state.

Peter Drucker once commented “for full effectiveness all the work needs to be integrated into a unified *program for performance*” [8]. The program for performance is designed by the top management team to provide a specific, effective course of action to achieve its desired results. To achieve these results, all the dimensions of the business must be consistent with each other. This is the job of the policy deployment system.

This system for managing policy consists of *kanri* or control mechanisms that deploy business policy through four essential steps in order to execute management’s program for the business direction using a systematic sequence of steps that achieve *hoshin*

project objectives within the constraints of assigned resources. These four steps define policy setting (or establishment of *hoshin* projects), deployment (or propagation of these projects throughout the organization), implementation (or integration of the results of change into the daily management system), and review (or assessment of the results achieved from the process) (see Figure 1). These four steps will be described in the next four sections of this chapter.

Policy Setting

Policy setting is a top-down “catchball” whereby management conducts “strategic dialog” with employees to collect ideas and opinions about chronic major problems and their aspirations regarding the business future, and then processes this information in conjunction with environmental data analysis and scenario analysis to formulate the annual business change objectives (which some organization call their *hoshin* projects): strategic change projects (identified by both targets to be achieved and means for achievement). In this phase, organizations recognize the most critical projects that must be accomplished to eliminate vulnerabilities or capture the benefits from potential change initiatives or newly emerging improvement opportunities. For organizations to succeed they must undergo a rigorous analysis of both their fundamental work processes to identify business imperatives (things that must change) and their current strategic direction to determine potential business vulnerabilities from competitive, economic, or technological changes.

This system structures application of continuous improvement into strategic and operational dimensions. David Packard incessantly used the term *continuous improvement* beginning in the early 1940s – it is not a new term – but as Collins and Porras describe its adaptation in leading companies that have adapted to change, they observe that it is an essential structural ingredient in those companies that have been *Built to Last*: “Visionary companies apply the concept of self-improvement in a much broader sense than just process improvement. It means long-term investments for the future, it means investment in the development of employees; it means adoption of new ideas and technologies. In short, it means doing *everything* possible to make the company stronger tomorrow than it is today” [9].

Most organizations operate on three levels of managerial thinking: enterprise thinking assures their

long-term viability; strategic thinking focuses on products, markets, and customers; and operational thinking focuses on the daily work that delivers the organization's results. Strategies align to these three areas of focus: "Management strategies can be classified into three types – corporate strategy, business strategy, as well as functional and cross-functional strategy – depending on the level of the corporate organization to which they apply. The corporate strategy, which delineates the fundamental direction of the whole company, is certainly very important for realizing a management vision; but it would be no exaggeration to say that the success or failure of the corporate strategy is determined by the particular business strategies, since it is through these business strategies that the aims of the corporate strategy are actually implemented" [10]. The portfolio of specific strategies that any organization pursues must be managed to deliver the risk – benefit performance desired by the organization in order to achieve its desired results – whether for breakthrough results or just for incremental improvement of a specific business area.

How is this approach to planning conducted? The corporate planning process should deliver increasing business brand value to balance financial risk and reward. This planning process consists of three elements: strategic planning, business planning, and functional planning that must all fit together in an integrated planning system. The strategic planning process is conducted at the enterprise level of the business thinking to identify which business opportunities to exploit and how to sustain the ability of the organization to meet or exceed its annual performance objectives. The business planning process is conducted at the business level of thinking and its objectives are to drive market share to accelerate financial payback, build customer loyalty, and decrease market risk. At the operational level of thinking, the functional planning process improves all process performance to reduce cost, cycle time, and defects while enhancing responsiveness to customers and delivering customer satisfaction.

A business strategy should deliver "visionary" performance: strategy is the persistence of a vision over the long term – and it requires both vision setting and vision deployment to assure alignment in strategic direction. Policy deployment provides direction to guide these plans and assures that the organization moves in a coherent direction. The more robust a plan for required change, the more

effective the organization's ability is to accomplish this plan. Robustness is a function of the management team's ability to see beyond its operating horizon and understand what may occur in its planning horizon that requires its focus and attention today.

What is a planning horizon? It is the distance that an organization "sees" into the future to study and understand the potential impacts of events on its policies and prepare it for evolving situations that may impact its current or future performance. Organizations tend to have four distinct planning horizons.

- **Business foresight**

Managing for the long term to assure that the organization is not surprised by changes in the assumptions that it has made in the design of its business model and product line strategy (focusing on a 3–10-year business outlook).

- **Strategic direction**

Managing for the intermediate-term changes in technology and competitive dimensions to assure that vulnerabilities in the business model are not exploited and to bridge the chasm that may exist between product line introductions (focusing on the next 3–5 years of business operation, depending on the degree of change that is anticipated in the business environment).

- **Business plans**

Managing the short-term fluctuations of the market – a planning horizon that delivers against short-term fluctuations in demand or supply (focus on quarterly and annual operating plans).

- **Business controls**

Managing the current state of a business – a planning horizon that delivers today's performance and assures rapid responses for corrective actions required to sustain advertised service levels (focus on the daily–weekly–monthly–quarterly operating plans to deliver the targeted results in the annual plan).

How is strategic policy formulated? Strategic direction is best established using cross-functional dialog to capture all the inputs of the organization and to build a common direction, based on the consensus of how to boost organizational strengths and overcome organizational weaknesses in the face of critical business threats to capture the most important market opportunities. Most organizations have just two kinds of strategic decisions: those that may

be executed within the areas of their direct oversight of top management (e.g., personnel decisions, budgeting, merger, capital budgeting, etc.) and those that require cross-organizational collaboration for implementation. These cross-functional projects require special attention and project management to realize the objectives of the change initiative. Such change strategies that require mutual consent and collaboration are ideal for a policy deployment system. In addition to planned continuous improvement that is a result of problem solving, continuous improvement may also result from process management, whenever a process is consciously enhanced over time.

Osada encourages strategic engagement of all employees through the following ways:

- recognizing product life stage (product life cycle analysis)
- analyzing business and product position objectively (positioning analysis)
- analyzing competitiveness (competitive analysis)
- perceiving strengths and weaknesses of products (SWOT analysis)
- forecasting future competitiveness using **time series** data (time series analysis)
- maintaining transparency through visualization (visual method); involving all employees [11].

Osada encourages using seven strategic tools (the S-7 tools) in policy setting:

1. environment analysis
2. product analysis
3. market analysis
4. product–market analysis
5. product portfolio analysis (product portfolio management, PPM)
6. strategic elements analysis
7. resource allocation analysis [12].

But, all these tools and methods are employed as staff-directed preconditions for strategic planning in the planning approaches of many Western organizations where they link the three planning systems (strategic, business, and functional) with the business and environmental assessment analyses that precede strategic decision-making. While this linkage may be a bit weaker in Japanese firms, in Western firms, such issues do not appear to be a critical shortfall. However, without complete integration of these planning processes it is difficult to obtain the degree of

effectiveness in deployment of shared resources that permits breakthrough achievement to occur. What is breakthrough achievement in management? Breakthroughs represent at least an order-of-magnitude change in performance that is accomplished over a relatively short period of time. Breakthrough is achieved by developing a capability to choose the right objectives for planned change and then aggressively executing these objectives. This requires two factors: identification of what to change and the timing of when to change it. The job of top management is to decide which lever of change must be pulled in order to accomplish the desired result.

Other success factors that are significant in achieving breakthrough plans include the right action to achieve the desired state of change. Right does not mean comprehensive or exhaustive, but it implies a budgeting of energy that focuses an organization on catalytic actions that stimulate organizational response in the desired direction – applying a limited capital budget and the best people to accomplish those important objectives that they have been personally developed to concentrate on. A second success factor in the management of breakthrough projects is the capacity of an organization to convert objectives into results. Excellence comes from execution of plans, not just from planning. To execute, an organization must be mobilized to consolidate their energy and coordinate their actions to achieve shared objectives for the common good of the organization as an organism – a living entity that requires appropriate nourishment and execution of all its bodily functions. A third critical success factor for breakthrough management is the capacity of an organization to integrate specific improvements into standard operating practices that are consolidated across the entire organization for maximum leverage effect. This success factor is based on the existence of a business control management system that holds the gains from improvement projects and is able to assure that performance degradation does not occur – people do not slip back into their “old way of doing things”, but that they embrace the new methods as their routine way of working. This success factor is addressed in the policy deployment and policy implementation steps.

It should be noted that breakthrough change projects can only be accomplished if the daily work processes are operating under reasonable control. If a business is not operating under control, then

8 Policy Deployment for Performance Improvement

the “breakthrough” activity should focus on bringing itself under control before making a significant investment in strategic change. When an organization’s daily management system is operating under control, then more time is available for strategic change because management is not “fighting fires” or “expediting execution” of routine work. This requires that the management start this journey by evaluating its readiness to make strategic changes in its operating system.

To maximize the effectiveness of policy deployment, the strategic direction setting and implementation processes should identify the highest priority business process improvement projects and the requirements for accomplishing them. For instance, only some of these projects will require the degree of diagnostic sophistication that is available from a Six Sigma Black Belt while others will require intensive capital investment or software development to accomplish their objectives. Only when management chooses change projects that improve the infrastructure of its business processes – work processes whose performance contributes to the common cause variation of the business performance – can the most significant gains be realized from a comprehensive business improvement effort. How does management achieve this focus? The short answer is that management must put in place the methods to “recognize” their priority business improvement needs by linking their choice of change management projects to the strategic direction of their business so that the portfolio of change management projects drives the policy changes that are necessary to achieve the long-term performance objectives of the organization. Some of the questions that must be addressed during policy setting include the following:

- What is our business and what results do we expect to achieve? How will we know that we have achieved these results? Is this the best we could do?
- What are our assumptions about society, the economy, market and customers, technology, and knowledge? Are they still valid?
- Has anything happened that would change the dynamics of our industry or markets?
- What would change mean for our business position?

- Are there any opportunities that we should anticipate and capitalize upon to our long-term advantage?
- Where should we choose to excel? Can we take action on this opportunity?

Policy Deployment

While policy setting sets organizational policy, the policy deployment step deploys policy change to the organization by changing the way that effective work is accomplished in the daily management system.

How are the *hoshin* or strategic change objectives deployed? Policy deployment is the heart of this policy management system and has received much attention because of the “catchball” approach that aligns the objectives of the organization and then balances work by resource leveling and prioritization of improvement activities. An implementation plan for a change project is a living document – it acts like a compass to guide an organization while allowing employees to take ownership by participating in choices that define the reasoning behind the project as well as the steps in the project’s execution. Change projects can identify two types of improvement effort – either by a breakthrough project or by a continuous improvement project – to change the way work process activities function. Breakthrough activities are strategic change projects that make a significant shift in the organization’s capability to perform routine operational work processes or deliver products (either goods or services) to the marketplace. Work process continuous improvement (*kaizen*) activities are part of a daily management system that defines how work is accomplished. The *kaizen* change activities are the responsibility of all work process owners. This planning approach focuses on guidelines for major improvement projects (while small incremental or continuous improvements are made through the regular course of continuous improvement of the daily routine work). The distinction between these activities is twofold: first, more of the organization’s resources are focused on breakthrough projects, and second, accomplishment of a breakthrough project usually occurs over a multiple-year period (or as a series of coordinated improvement projects). One important management consideration in choosing breakthrough projects is that the combination of all the annual breakthrough projects (also called the *portfolio of change projects*) will define the steps that an organization chooses for

accomplishing strategic change in the range of its midterm planning horizon (1–3 years).

To achieve “saturation” of policy (which consists of both targets and the means for their achievement) in the organization – or complete deployment of change projects within the whole organization that is affected by the defined policy change – and assure collaboration of all the affected work groups, the objectives cascade of an action plan for a particular improvement project must involve not only functional deployment of policy but also engage its cross-functional aspects. It is across the functional seams of an organization where most significant difficulties are encountered and these boundaries represent focus areas for management to assure continuous collaboration in the execution of change projects and consensus among the various functional organizations that engage all the decision-making managers in the areas where the change will have a direct effect. To understand the difficulty that the boundary condition dynamics have, consider what happens as change is managed when organizations shift work activities from internal to external units (e.g., from internal manufacturing to an external contract manufacture). At such boundary conditions, conflicting objectives and political issues of the organizations often can interfere with performance improvement work and it is the job of the management team to eliminate any such barriers to the success of their project team. This also happens within organizations at their functional boundaries.

Catchball is the process that is used to build a consensus through dialog about the targets and means to achieve the change. This process is data driven and uses tools that permit management by facts. Catchball links annual change projects to midrange and long-term plans – deployment prior to annual fiscal year commencement, incorporated into target setting and annual employee objectives cascade; coordinated both vertically within functions and horizontally across processes and negotiated across the processes to allocate resources (competence, funding, and equipment) to achieve the shared and agreed objectives. The catchball process includes four activities:

- building alignment through linked cascade of means

- setting business performance targets and objectives
- cascading business objectives to the workplace (*gemba*)
- achieving alignment of improvement and effective resource allocation.

Policy deployment is a structured, systematic, and standardized process and it has an ability to empower organizations to achieve strategic change. Policy deployment includes a few key elements that assure that an organization is properly and fully engaged in change projects:

- *Policy* – a general rule or operating principle that describes a management-approved process to approach a business condition or situation based on how it chooses to control its work and manage risk. Once the right policy has been determined, then the organization can handle similar situations with a pragmatic response by adapting its policies to the concrete situation that it faces. Truly unique business situations that run counter to the critical business assumptions require the full attention of the senior management team to evaluate how these situations challenge the boundary conditions of the business model and threaten its policies of operations with change that is imposed from externalities. Policies consist of targets and means.
- *Target* – the measurable results that are to be achieved within a specific timeframe for performance. Targets have checkpoints.
- *Checkpoints* – a measurement point that is used to evaluate an intermediate state in the policy deployment process to demonstrate that progress is being made. The data collected at a checkpoint can be reported to management in interim project status reports. The checkpoint of one process is the control point of the next process – the checkpoints and control points work together to formulate a “waterfall” that cascades across the implementation plan flow and is part of the business measurement system.
- *Check items* – check items and process or project variables that are evaluated to enable organizations to understand the causes that contribute to the outcome of a particular policy.
- *Means* – the sequence of actions that an organization will take to implement a policy or choice of the management team that is an outcome of the

strategic direction-setting process. Means have control points.

- *Control points* – a point in a sequence of work activities where corrective action may be taken or countermeasures are put in place to resolve a concern or issue that has been identified at a checkpoint.
- *Control items* – control items verify whether results agree with the established goals – does the work demonstrate progress in accomplishments that will enable the final achievement of targets?
- *Deployment* – the process of engaging the entire organization in an appropriate participation in the strategic direction both vertically (within functional areas) and horizontally (across the functional areas) by creating shared ownership of the implementation actions. The entire system of deployment is “connected” from the long-term vision to the daily management activities. The plans are progressively more detailed as they are refined in deployment from the top levels of the organization to the frontline employees and teams. The plan is deployed through an organization by negotiating the means between management layers (levels) as well as across the functional departments. Targets are not negotiated.
- *Catchball* – the joint analysis process that encourages a strategic dialog between levels of organizational deployment is called a *catchball*. The means at one level become the desired outcomes of the subsequent level. This cascading of targets and means establishes the linkage and alignment of objectives across organizational levels. Mutual discussion between the parties – a two-way communication that is both top-down in general direction and bottom-up in adaptation to the workplace using the existing hierarchical management structure and matrix process structure to engage all parts of the organization in the dialog. This dialog is a negotiation process (see *nemawashi* and *sureawashi* below) that arrives at a collective wisdom to develop and refine the implementation plan.
- *Nemawashi* – negotiation – prior consultation to achieve consensus; careful preparation of the roots of a plant for transplantation; and seeking to achieve *wa* or harmony, consensus, and absence of conflict.
- *Sureawashi* – the sharpening of a sword requires four ingredients: a blade, a template for the angle to be produced, oil to ease the process, and a sharpening stone to remove the unwanted material. This analogy is used for policy deployment. The use of data makes the objectives cascade a fact-based process, not just a subjective negotiation process. Mutual consultation occurs between the organizational levels in order to test the feasibility of planned process improvements and to refine any conflicts between the objectives of the organizational layers. This process opens communication channels and establishes both agreement and alignment in the way people work. This dialog is necessary to obtain buy-in and define achievable plans that middle managers are committed to implement.
- *Shibui* – a state of uncluttered, beautifully efficient austerity – the perfect balance or harmony (*wa*) between not enough and too much – used to describe the desired state or vision of a business system.

Peter Drucker quotes the Roman law in order to focus management on the things that are most important: “*De minimis non curat praetor*” (The magistrate does not consider trifles) [13]. This warning to management against what has been called “micro-management” is a reason for senior executives to focus on the vital few issues that are critical in the business that they manage. If they do not take the time to manage these important things, then no one else will. If they choose to spend their time focused at the detail level of project execution, then they will squander a more effective use of their time on those vital activities that engage the higher thinking levels of the organization that cannot be reasonably delegated to others for effective action. Management must work on a long-term planning horizon in order to deliver sustained organizational strength. It must also review current actions to assure that short-term profitability is being achieved. But, whenever management spends more time on the short-term issues than it does on the long-term ones, it sacrifices future strength in favor of current results – and displays to the entire organization its lack of trust in the ability of the organization to perform its daily work. This behavior signals to the entire organization that a crisis exists and reinforces stagnation as the workers wait for the top management to intervene and make the decisions that they

should more properly make. A very important benefit of an effective policy deployment system is delegation of appropriate decision rights to the proper place in the organization where the best information exists and where action will be taken to implement that decision.

Some of the questions that are addressed during policy deployment include the following:

- What are the consequences of not doing this project?
- What are the risks inherent in this project?
- What will happen to the business if this project does not succeed?
- What will success in this project commit the business to?
- How does this project add to the total economic results of performance?
- Have we assigned our best people to work on breakthrough opportunities?
- Have we communicated clearly and taken into consideration all objections before chartering the project?

Policy Implementation

Policy implementation consists of the execution of the project plan – both the actions taken by the team involved in the change and the in-process management reviews. All change is implemented on a project-by-project basis according to management's priorities and the logical sequence for attacking each opportunity for improvement. The project plan typically will use a Gantt chart to assign clear responsibility for each improvement item in the implementation plan and to record its activity progress in accomplishing the project subtasks. Senior managers should also conduct regular progress reviews of each change project to monitor team progress in improvement, assure that the projects advance, and eliminate any possible barriers, roadblocks, or bottlenecks that restrict the project's advancement. Senior management should also monitor the execution of the change projects that they sponsor to assure that these projects will deliver the desired gain in the daily management system. If the project review indicates that insufficient progress is being made, then they can develop some countermeasures or reallocate resources so that appropriate corrective actions are taken to assure continued progress.

Another activity that occurs during the policy implementation is publication of information about the ongoing change projects so that the entire organization is informed of the actions being taken to improve performance. This communication can help the organization to align other activities with progress being made on these strategically focused change projects. As a guideline for communication, the management should inform all involved parties of any changes to the change project team's mission, vision of the outcome, guiding principles, or objectives. If the management team communicates effectively and often, then it will translate the planning rhetoric into action realities. Peter F. Drucker observed that "The most time consuming step in the process is not making the decision, but putting it into effect. Unless a decision has degenerated into work it is not a decision; it is at best a good intention" [14].

Some of the questions addressed during policy implementation include the following:

- Have we placed the right people in the right jobs to give the project the best opportunity to succeed?
- Does this project team have everything that it needs to get the job done?
- Are all the people who need to know about this project being informed?
- Are all the right actions across the organization being taken to assure success?
- Is this project implementation the best utilization of the knowledge and ability of the organization's people?
- Does this project implementation make the best overall contribution from use of the organization's limited resources (people, time, and money)?

Policy Review

As the annual cycle of change projects nears completion, results of project implementation efforts should be evaluated to determine performance against targets, shortfalls from expected performance, completion rate, and causes for both under- and overachievement. Specific action must be identified to compensate for performance deficiencies and prevent recurrence of such problems in future change management projects. Diagnosis of the performance of the policy planning process is conducted to drive improvements in planning systems. "Feedback has to

be built into the decision to provide a continuous testing, against actual events, of the expectations that underlie a decision” [15].

Policy review is accomplished in two ways: through management self-assessment (by senior managers as well as by local managers evaluating their activities to determine where they have opportunities for improvement: either performance enhancements or problem resolution) and through operating reviews of the results produced by the local organization where senior managers identify areas where results are not aligned with expectations for performance. Policy review applies two subprocesses to perform these duties: performance review and business measurement.

Aligning Objectives through Performance Review. The review process seeks to identify conformance to plans (e.g., is there any shortfall or overachievement in targets?). Once nonconformity is identified, then the root cause of the deviation is discerned to determine an appropriate response to the out-of-control type of condition. Both corrective actions and countermeasures are identified to realign the process and assure that process integrity and stability are achieved in the business control system. The following actions may be taken in response to an out-of-control condition:

- emergency countermeasures to alleviate the immediate issue, concern, or problem;
- short-term corrective action to prevent the specific problem from recurring; or
- long-term preventive action to remove problem root causes and mistake-proof the process making a permanent solution and preventing the problem from recurring.

The review step facilitates organizational learning by examining problem areas and critical success factors to discover what directional shifts need to be accomplished in order to achieve the desired end state or vision of the business. Strategy is the persistence of the vision – achieved one project at a time through exercising constancy of purpose in the business planning process. These project reviews are conducted to assess achievement relative to the following planning elements:

- change project objectives
- business planning objectives and corporate commitments

- business improvement plans
- economic plans and projections
- customer requirements and expectations
- competitive performance analysis
- business excellence self-assessment.

Questions addressed during this policy review include the following:

- What results have been demonstrated from this project?
- Which results were expected and which results were unexpected?
- What will this project’s outcome do for customers?
- What have we done well that our competitors have done poorly?
- What have we done poorly that our competitors seem to have no problem with?

Business Control and Management Responsibility. The ultimate objective of MBP is to establish a reliable organization – one that creates predictable results through the effective coordination of value-adding work that customers perceive as meeting their needs. In this environment, all employees are aware of their personal contribution to the objectives of the entire organization and are able to make local choices that are aligned with the strategic direction because they understand how the strategy affects their work and *vice versa*. To assure that these local decisions are aligned with strategic direction, it is the responsibility of the management team to develop a measurement system that provides employees with the visible line of sight from their work activities to its contribution to strategic direction. In this measurement system, it is essential that causal linkages (e.g., built using a Six Sigma $Y = f(X)$ transfer function) be established so that effective control can be executed at the local operating level.

Benefits of Policy Deployment

A policy deployment system orchestrates continuous improvements with breakthroughs to assure that the organization attains its long-range goals. Those elements of long-range plans that can be achieved in a one-year period are identified and become the focus or “vital few” goals to be achieved during that year. Policy deployment plans the way that change is implemented in an organization’s daily management

process. Accomplishments of such a planning system include

- communicating the vision required for sustained success,
- identifying and choosing breakthrough activities or projects required for the vision,
- orchestrating the direction of an organization's change,
- developing plans and projects that support the business objectives,
- aligning the organization's change efforts both vertically and horizontally,
- ensuring that the plan is effectively and efficiently executed,
- reviewing the progress in executing plans,
- changing plans when it is proved necessary to achieve targets,
- learning from the experience of planning and executing.

"If you can think of new methods to preserve the core, then by all means put them in place. If you can invent powerful new methods to stimulate progress, then give them a try. Use the proven methods *and* create new methods. You must do both" [16]. The imperative for organizations that endure is to do both breakthrough improvement and evolutionary improvement – both change management and routine management – at the same time. This is what Collins and Porras call "the genius of AND" – an inclusive approach to planning and executing change that requires organizations to embrace both aspects of change simultaneously.

The most important thing about priority decisions that face a business is that they are made and communicated deliberately and conscientiously. In a system of MBP, all the important decisions are visible and there is an opportunity for dialog to guide these decisions into the direction that the organization, as a whole, is influenced by the knowledge of all its members. In such a system the key decisions that drive toward its common goals are not made haphazardly, but with the full awareness of the organization. Such open decision processes elicit cooperation of the entire organization in the implementation and review of its activities to assure that it will be able to meet its desired outcomes. The responsibility of the management is to put in place a system of decision-making that generates this degree of collaborative work toward the common end.

Criticisms of Policy Deployment

Despite their application in many leading companies, policy deployment systems have been criticized for their mechanistic use of forms and templates that some see as restricting individual creativity. Some also believe that these planning systems lack strategic emphasis and do not engage the full organization as participants in strategy formulation. Osada summarizes the shortcomings of policy deployment as observed in some Japanese companies.

1. It is difficult for those at the middle management level and below to understand the process of formulating strategic policy. Compared with [editor note: the process for] policy deployment, the process of policy formulation is unclear [editor note: poorly understood and communicated] an indication of management's view that such form of communication is of little value.
2. Strategic policy is ostensibly based on the long-term interests of the firm, but there is no way to judge whether a policy is appropriate, or even truly 'strategic' [editor note: in the essential nature of the policy itself].
3. Several problems in formulating a long-term strategic plan are not addressed, for instance:
 - 3a. Changes in operating environment and other uncertainties are not adequately accounted for; possible difficulties are therefore not foreseen.
 - 3b. Positioning of business is not perceived objectively. The question of whether business aims are optimum and clear is not addressed.
 - 3c. Only one part of the staff, at the top level, participates in strategic policy formulation; it is therefore difficult to judge whether a policy reflects the reality at the "front line" of operations [10].

It must be observed that not all of these objections are strongly negative. Organizations must ask if they must really involve frontline employees actively in formulating strategy. Nokia Mobile Phones uses a current state analysis for self-assessment of front-line operations and then rolls this data into their strategy-setting process. They also create a "strategic dialog" that builds participation of midlevel managers in conversations about strategy. Other organizations open communication lines through email forums and

internal surveys. In such instances, the objection is not critical to the total impact of policy deployment implementation. Additionally, any argument that says “Every employee should have an *interest* in matters of strategic policy” is a very different argument than saying that every employee should be actively involved in *formulation* of business strategy. Satisfying employee interest in strategy can be addressed by improving communications. Also, a broad involvement of employees in formulation of strategy increases risk of inappropriate public statements or inadvertent disclosure of the company’s strategy in venues where competitors may discover sensitive information that can be used against the organization. Whenever this occurs, a company loses its competitive advantage. The challenge for management is to build a strong consensus, without risking disclosure of their strategic direction to competition. It is another challenge of the management to balance these concerns in a way that is appropriate to their way of working and business culture.

Summary

Policy deployment, when coupled with a statistically based business measurement system, has been observed in several leading companies to create a robust management process that engages an entire organization in the strategic planning process. It assures line of sight from the strategic goals of the organization to the operational tasks that workers perform at the front line as they do the work that produces the organization’s goods or services. The nature of this process can be described using the term “robustness” – a statistical state in which a process is able to accept variation in its inputs, without influencing the variation of its outputs. Such a process is capable of performing consistently – delivering consistent results according to its design intent. Because policy deployment engages the workforce in achieving its common goal of sustained success, this methodology has become a strategic tool for assuring sustained competitive advantage over current and potential business rivals.

Sustained success must be “dynamic” to achieve its enduring state. That is, it must provide continuous advantage despite changes in the environment, regulatory shifts, technological breakthroughs, or competitive market. Anticipating potential actions by rivals is critical to delivering sustained success. To

enjoy such sustained success, an organization must master the skills of priority setting and project management to assure that they effectively define and deploy the right initiatives that result in sustained success. Advantage means staying ahead of rivals and this requires that organizations not only make continuous improvements but also use “breakthrough” opportunities to distinguish themselves in their marketplace as a differentiated provider of products and services. This type of management requires managerial competence in three areas: business vulnerability analysis, action planning administration, and operational excellence. The best implementations of policy deployment will engage its employees in the strategy-setting processes as well as the organization’s change management process.

End Note

^a It must be noted that Peter F. Drucker initially discussed MBO in Japan in the mid-1950s. Drucker taught management concepts to the Japanese along with Dr. Joseph M. Juran and Dr. W. Edwards Deming. At that time Dr. Juran and Dr. Deming worked in the Graduate Management School of New York University under the supervision of Dr. Drucker.

References

- [1] Drucker, P.F. (2002). Keynote Address, *56th Annual Quality Congress*, 20 May 2002.
- [2] Akao, Y. (ed) (1991). *Hoshin Kanri: Policy Deployment for Successful TQM*, Productivity Press, Portland, p. xxx.
- [3] Collins, J.C. & Porras, J.I. (1994). *Built to Last: Successful Habits of Visionary Companies*, Harper Business, New York, p. 156.
- [4] Collins, J.C. & Porras, J.I. (1994). *Built to Last: Successful Habits of Visionary Companies*, Harper Business, New York, p. 146.
- [5] Osada, H. (1998). Strategic management by policy in total quality management, *Strategic Change* 7, 277–287.
- [6] Osada, H. (1998). Strategic management by policy in total quality management, *Strategic Change* 7, 277.
- [7] Collins, J.C. & Porras, J.I. (1994). *Built to Last: Successful Habits of Visionary Companies*, Harper Business, New York, p. 215.
- [8] Drucker, P.F. (1964). *Managing for Results*, Harper & Row, New York, p. 193.
- [9] Collins, J.C. & Porras, J.I. (1994). *Built to Last: Successful Habits of Visionary Companies*, Harper Business, New York, p. 186.
- [10] Osada, H. (1998). Strategic management by policy in total quality management, *Strategic Change* 7, 278.

- [11] Osada, H. (1998). Strategic management by policy in total quality management, *Strategic Change* **7**, 279.
- [12] Osada, H. (1998). Strategic management by policy in total quality management, *Strategic Change* **7**, 281.
- [13] Drucker, P.F. (1985). *The Effective Executive*, Harper & Row, New York, p. 114.
- [15] Drucker, P.F. (1985). *The Effective Executive*, Harper & Row, New York, p. 139.
- [16] Collins, J.C. & Porras, J.I. (1994). *Built to Last: Successful Habits of Visionary Companies*, Harper Business, New York, p. 216.

GREGORY H. WATSON

Classification and Regression Tree Methods

Introduction

Classification and regression are two important problems in statistics. Each deals with the prediction of a response variable y given the values of a vector of predictor variables $\mathbf{x}$. Let $\mathcal{X}$ denote the domain of $\mathbf{x}$ and $\mathcal{Y}$ the domain of y . If y is a continuous or discrete variable taking real values (e.g., the weight of a car or the number of accidents), the problem is called *regression*. Otherwise, if $\mathcal{Y}$ is a finite set of unordered values (e.g., the type of car or its country of origin), the problem is called *classification*. In mathematical terms, the problem is to find a function $d(\mathbf{x})$ that maps each point in $\mathcal{X}$ to a point in $\mathcal{Y}$. The construction of $d(\mathbf{x})$ requires the existence of a *training sample* of n observations $\mathcal{L} = \{(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_n, y_n)\}$. In computer science, the subject is known as *supervised learning*. The criterion for choosing $d(\mathbf{x})$ is usually mean squared prediction error $E\{d(\mathbf{x}) - E(y|\mathbf{x})\}^2$ for regression, where $E(y|\mathbf{x})$ is the expected value of y at $\mathbf{x}$, and expected misclassification cost for classification.

If $\mathcal{Y}$ contains J distinct values, the classification solution (or *classifier*) may be written as a partition of $\mathcal{X}$ into J disjoint pieces $A_j = \{\mathbf{x} : d(\mathbf{x}) = j\}$ such that $\mathcal{X} = \cup_{j=1}^J A_j$. A *classification tree* is a special form of classifier in which each A_j is itself a union of sets, with the sets being obtained by recursively partitioning the $\mathbf{x}$ space. This permits the classifier to be represented as a *decision tree*. A *regression tree* is similarly a tree-structured solution in which a constant or a relatively simple regression model is fitted to the data in each partition.

To illustrate, consider a dataset of new cars for the 2004 model year from the *Journal of Statistics Education Data Archive* (www.amstat.org/publications/jse/jse_data_archive.html). After filling in missing values and correcting errors, we obtain a dataset with complete information on 15 variables for 428 vehicles. Table 1 gives the definitions of the variables and Figure 1 shows two classification trees for predicting the *Region* variable, that is, whether the vehicle manufacturer is from Asia, Europe, or the United States (there are 158 Asian, 120 European, and 150 US vehicles in the dataset). Each

intermediate node has a condition associated with it. If an observation satisfies the condition, it goes down the left branch; otherwise it goes down the right branch. Each leaf node is labeled with a predicted class, with cross-hatched, light gray, and dark gray denoting Asia, Europe, and the United States, respectively. Beneath each node is a fraction, giving the number of misclassified samples divided by the number of samples. It is easily seen from the trees that *Dcost* is an important predictor for European cars: 125 of 159 cars with *Dcost* greater than \$30,045 are European. Of cars with *Dcost* less than or equal to \$30,045, those with 3.5-l or larger engines are six times as likely to be American, while those with smaller engines are predominantly Asian.

A classification or regression tree algorithm has three major tasks: (a) How to partition the data at each step? (b) When to stop partitioning? and (c) How to predict the value of y for each $\mathbf{x}$ in a partition? There are many approaches to the first task. For ease of interpretation, a large majority of algorithms employ *univariate* splits of the form $x_i \leq c$ (if x_i is noncategorical) or $x_i \in B$ (if x_i is categorical). The variable x_i and the split point c or the split set B are often found by an exhaustive search that optimizes a *node impurity criterion* such as entropy (for classification) or sum of squared residuals (for regression). There are also several ways to deal with the second task, such as stopping rules and tree pruning. The third task is the simplest: the predicted y value at a leaf node is the class that minimizes the estimated misclassification cost (for classification) or the fitted value from a model estimated at the node (for regression).

We review and compare below a few of the more established and widely available algorithms. First, we introduce some notation. Let N_j be the number of class j training samples and let $\pi(j)$ be the prior probability of class j . At each node t , let $N_j(t)$ be the number of class j training samples in t and let $p(j|t) = \pi(j)N_j(t)/N_j$ denote the estimated probability that an observation in t belongs to class j . It is easily verified that $p(j|t) \propto N_j(t)$ if $\pi(j) \propto N_j$.

CART and RPART

We first discuss CART [1] classification with the *Gini index* as node impurity criterion, $i(t) = 1 - \sum_{j=1}^J p^2(j|t)$. Suppose a split divides the data in

2 Classification and Regression Tree Methods

Table 1 Variables in the car dataset

Name	Definition	Values (#unique values in parentheses)
Region	Manufacturer region	Asia, Europe, or United States (3)
Make	Make of car	Acura, Audi, etc. (38)
Type	Type of car	Car, minivan, pickup, sports car, sport utility vehicle, or wagon (6)
Drive	Drive type	Front, rear, or four-wheel drive (3)
Rprice	Suggested retail price	US dollars (410)
Dcost	Dealer cost	US dollars (425)
Engnsz	Size of engine	Liters (43)
Cylin	Number of cylinders	−1 for the rotary-engine Mazda RX-8 (8)
Hp	Horsepower	Hp (110)
City	City miles per gallon	Miles (29)
Hwy	Highway miles per gallon	Miles (32)
Weight	Weight of car	Pounds (347)
Whlbase	Length of wheel base	Inches (40)
Length	Length of car	Inches (66)
Width	Width of car	Inches (18)

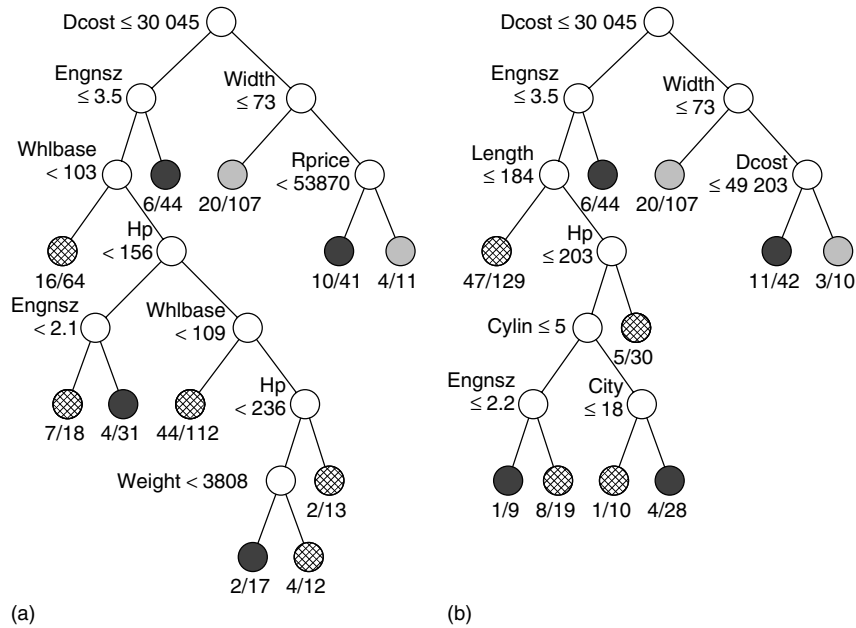


Figure 1 RPART (a) and QUEST (b) trees for Region. At each intermediate node, an observation goes to the left branch if and only if the condition shown at the node is satisfied. Leaf nodes that are cross-hatched, light gray, and dark gray are classified as Asia, Europe, and the United States, respectively. The fraction below a leaf node gives the number misclassified divided by the sample size. The total numbers misclassified are 119 for RPART and 106 for QUEST

t into a left node t_L and a right node t_R . Let p_L and p_R be the proportions of data in t_L and t_R , respectively. CART selects the split that maximizes

the decrease in impurity $i(t) - p_L i(t_L) - p_R i(t_R)$. Instead of employing stopping rules, CART generates a sequence of subtrees by growing a large tree

and pruning it back until only the root node is left. Then it uses cross-validation to estimate the misclassification cost of each subtree and chooses the one with the lowest estimated cost. The tree in Figure 1(a) is produced by RPART [2], a version of CART implemented in R [3].

Like the AID and THAID [4, 5] algorithms before it, CART has some undesirable properties. First, the splitting method is biased toward variables that have more distinct values. This is because a variable with m distinct values allows $(m - 1)$ splits if it is noncategorical and $(2^{m-1} - 1)$ splits if it is categorical. The bias was first noted in [6] and is discussed in [7]. Second, the exponential number of splits for categorical variables causes serious computational difficulties when m is large and y takes more than two values (a computational trick reduces the number of splits searched to $(m - 1)$ if y takes two values). Finally, the algorithm is also biased toward variables with more missing values. This is due to the impurity function depending only on the sample proportions, not the sample sizes. During split selection, if a variable x_i has missing values, only the observations nonmissing in both x_i and y are used in computing the decrease in impurity. Thus it is easier to “purify” a node by splitting on a variable with more missing values [8].

For regression, CART uses the sum of squared **residuals** as the impurity function and builds a piecewise-constant model with each leaf node fitted by the training sample mean. Everything else remains the same as in classification. For example, Figure 2 displays a RPART tree for predicting H_p with variables Region, Type, Drive, Engnsz, Cylin, Weight, Whlbase, Length, and Width. The predicted values beneath the leaf nodes show that H_p tends to increase with Cylin and tends to be higher for sports cars.

Piecewise-constant regression trees tend to have lower prediction accuracy compared to other regression methods that have more smoothness. In fact, CART regression trees typically have lower accuracy than even the classical multiple linear model – see, for example, [1, p. 227] and [9]. In the example here, the RPART model has an R^2 value of 75%, compared to 81% for the multiple linear model. One way to increase the accuracy of the piecewise-constant model is to use a larger tree, that is, with less pruning. The correct amount of pruning is, however, usually difficult to determine. Besides, a

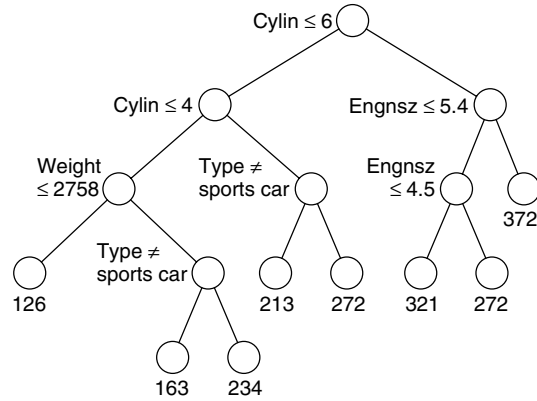


Figure 2 RPART tree for predicting H_p . The number in italics beneath each leaf node is the predicted value

large tree is undesirable because it is more difficult to interpret.

QUEST and CRUISE

QUEST [7] is a classification method. It avoids the selection bias and categorical variable computational problems of CART by first selecting the variable x_i and then selecting its split point or split set. This saves a significant amount of computation, because the algorithm does not have to search for the split points or split sets of the other variables. QUEST uses hypothesis tests to choose the split variables, namely, analysis of variance F tests for noncategorical variables and χ^2 tests for categorical variables. The variable with the smallest significance probability is selected to split a node. If the dataset has more than two classes, they are grouped into two superclasses prior to split point or split set selection. This superclass grouping allows application of the CART two-class computational shortcut for splits on categorical variables. The QUEST tree for our dataset is shown in Figure 1(b). It has the same splits at the top two levels as the RPART tree.

CRUISE [8] is another classification algorithm with unbiased variable selection. It differs from QUEST in three important respects: (a) each node can be split into as many as J branches, with one class forming a majority in each branch, (b) the significance tests include interactions between pairs of variables, and (c) it can fit linear discriminant models in each leaf node [10]. Being able to split

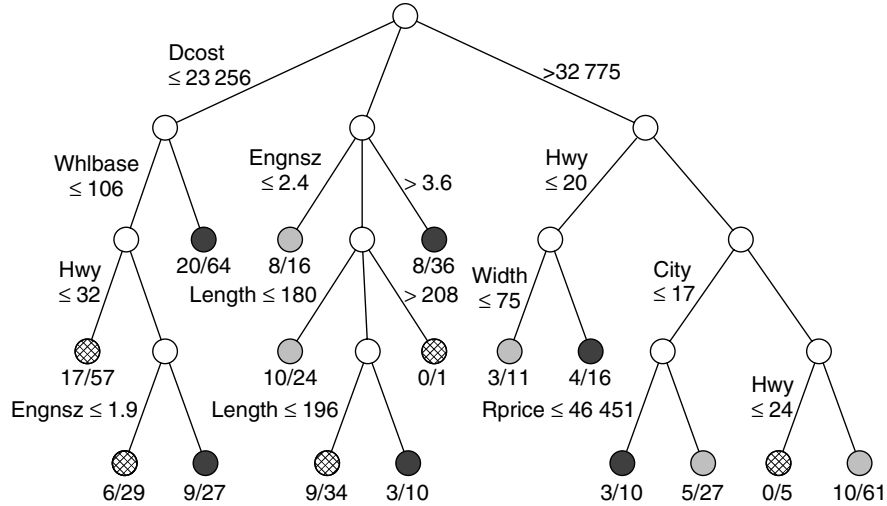


Figure 3 CRUISE tree for predicting Region. Cross-hatched, light gray, and dark gray leaf nodes are classified as Asia, Europe, and United States, respectively. The fraction below a leaf node gives the number of misclassified divided by the sample size. Total number of misclassified is 115

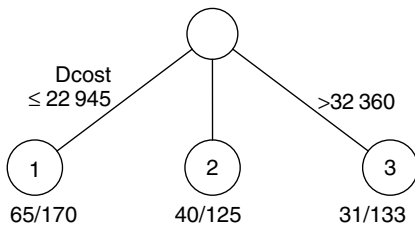


Figure 4 CRUISE tree with linear discriminant node models for predicting Region. The fraction below a leaf node gives the number of misclassified divided by the sample size

a node into J branches can be advantageous if J is large, because it increases the probability that at least one leaf node is assigned to each class. For example, Figure 3 shows the CRUISE tree for predicting Region. The vehicles in the left branch ($\text{Dcost} \leq 23\,256$) are mostly from Asian manufacturers and those in the right branch ($\text{Dcost} > 32\,775$) are mostly European.

Figure 4 shows the corresponding CRUISE tree where a linear discriminant model is fitted to the data in each leaf node. Notice how compact it is. This is due to part of the model complexity being absorbed by the discriminant models whose boundaries and data points are displayed in Figure 5. The plot for

node 3 reveals a very expensive car with Dcost more than \$170 000. It is a Porsche.

C4.5 and J48

C4.5 [11] is also a classification tree algorithm. If the variable x_i chosen to split a node is noncategorical, the node is divided into two using a split of the usual form $x_i \leq c$. On the other hand, if the variable is categorical taking m values, the node splits into m branches, with one branch for each categorical value. As a consequence, C4.5 has no difficulty dealing with categorical variables regardless of the size of m . The algorithm chooses the splits as follows. Suppose a node t is split into subnodes $t_1, t_2, \dots, t_k$. Let $e(t) = -\sum_j p(j|t) \log\{p(j|t)\}$ be the entropy at t and define the *gain* of the split as $e(t) - N(t)^{-1} \sum_i e(t_i)N(t_i)$. Define the *gain ratio* as the gain divided by $N(t)^{-1} \sum_i N(t_i) \{\log N(t_i) - \log N(t)\}$. C4.5 chooses the split that yields the highest gain ratio. C4.5 also grows a large tree and then prunes it back to a smaller size. But instead of cross-validation, it uses a conservative estimate of the error at each node for pruning.

Figure 6 shows a C4.5 tree, obtained using the J48 implementation in WEKA [12], with *Drive* as the response variable and *Make* as one of the predictor

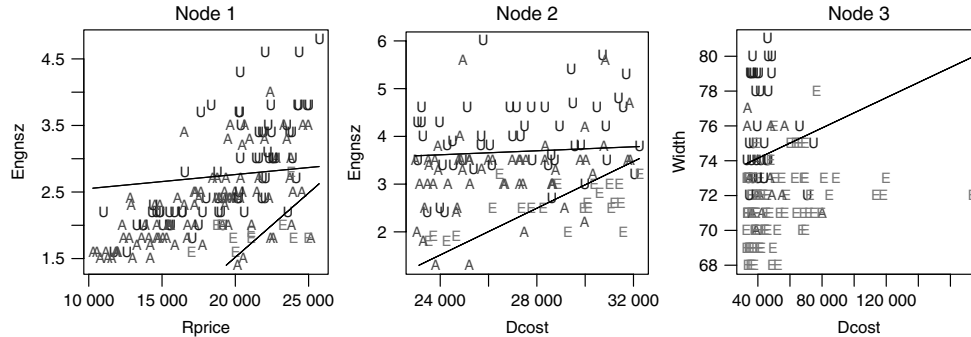


Figure 5 Data and linear discriminant boundaries for the nodes of the CRUISE tree in Figure 4. Asian, European, and US cars are denoted by the symbols A, E, and U, respectively. Node 3 has only one boundary: it serves to separate the European and US cars

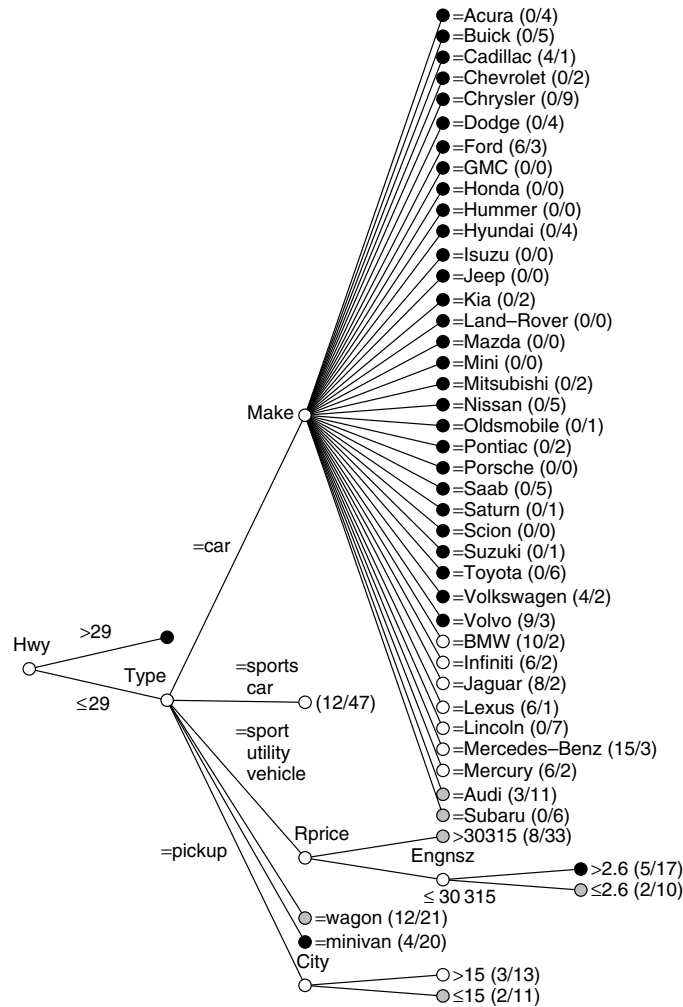


Figure 6 C4.5 (J48) tree for Drive, with black, white, and, gray leaf nodes denoting front-wheel, rear-wheel, and four-wheel drive, respectively

variables. As anticipated, when *Make* is chosen, the node splits into many branches. This is not a problem in itself, for we can manually merge the leaf nodes associated with the same predicted class. The problem is that each node before merging has too few samples to support further splitting. Hence the tree may be too short.

The avoidance of searching for binary splits on categorical variables and of cross-validation for pruning makes C4.5 one of the fastest classification tree algorithms. Its trees, however, tend to have more leaf nodes than those of other methods [13]. Like CART, it is biased toward selecting variables that allow more splits.

GUIDE

GUIDE [14] constructs piecewise-constant, multiple linear, and simple polynomial tree models for least-squares, quantile, Poisson, and proportional hazards regression. Like CRUISE and QUEST, its variable selection is unbiased. This is accomplished by recursively carrying out the following steps at each node: (a) fit a model to the training data there, (b) cross-tabulate the signs of the residuals with each predictor variable to find the one with the most significant χ^2 statistic, and (c) search for the best split on the selected variable, using the appropriate loss function. After a large tree is constructed, it is pruned with the cross-validation method of CART.

Figure 7 shows the GUIDE tree for predicting H_p when the best simple linear model is fitted to the data in each node (two cars with rotary engines are removed due to the *Cylin* variable being undefined for them). The sample mean value of H_p and the best linear predictor variable are given beneath each leaf node. The first two splits divide the tree into three main branches, with Asian, European, and American makes occupying the left, middle, and right branches, respectively. The most important linear predictors are *Cylin* and *Engnsz*. A major advantage of this model, compared to a piecewise multiple linear model, is that the data and the fitted regression functions can be visualized graphically, as demonstrated in Figure 8. We see from node 5 that European sports cars with large engines have among the highest H_p values. For American cars, H_p increases linearly with

Engnsz, irrespective of the other variables. This GUIDE model has an R^2 value of 84%, which is higher than that of the RPART and multiple linear models discussed in the section titled “CART and RPART”.

Conclusion

All the above algorithms can deal with datasets that have missing values. CART uses “surrogate splits” to pass observations with missing split values through its nodes. CRUISE follows a similar approach, but using an alternative set of splits. It is not known if one method is better than the other. C4.5 uses weights, passing an observation with a missing split value down every branch, where the weight is proportional to the number of observations with nonmissing values in the split variable. GUIDE and QUEST employ nodewise plug-in estimates for the missing values.

CART, CRUISE, and QUEST accept user-specified class prior probabilities and unequal misclassification costs. Class priors are useful if the training dataset is not a random sample. The three algorithms also allow splits on linear combinations of variables. Although such splits are hard to interpret, there is empirical evidence that they have much better prediction accuracy than trees with univariate splits [13]. C4.5 does not

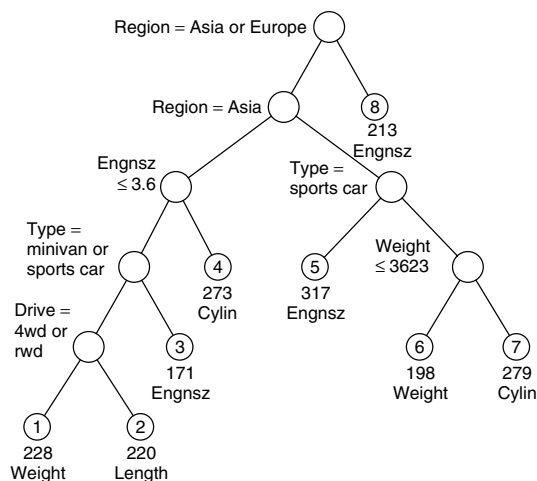


Figure 7 GUIDE piecewise simple linear model for predicting H_p . Beneath each leaf node are the sample mean of H_p and the name of the linear predictor

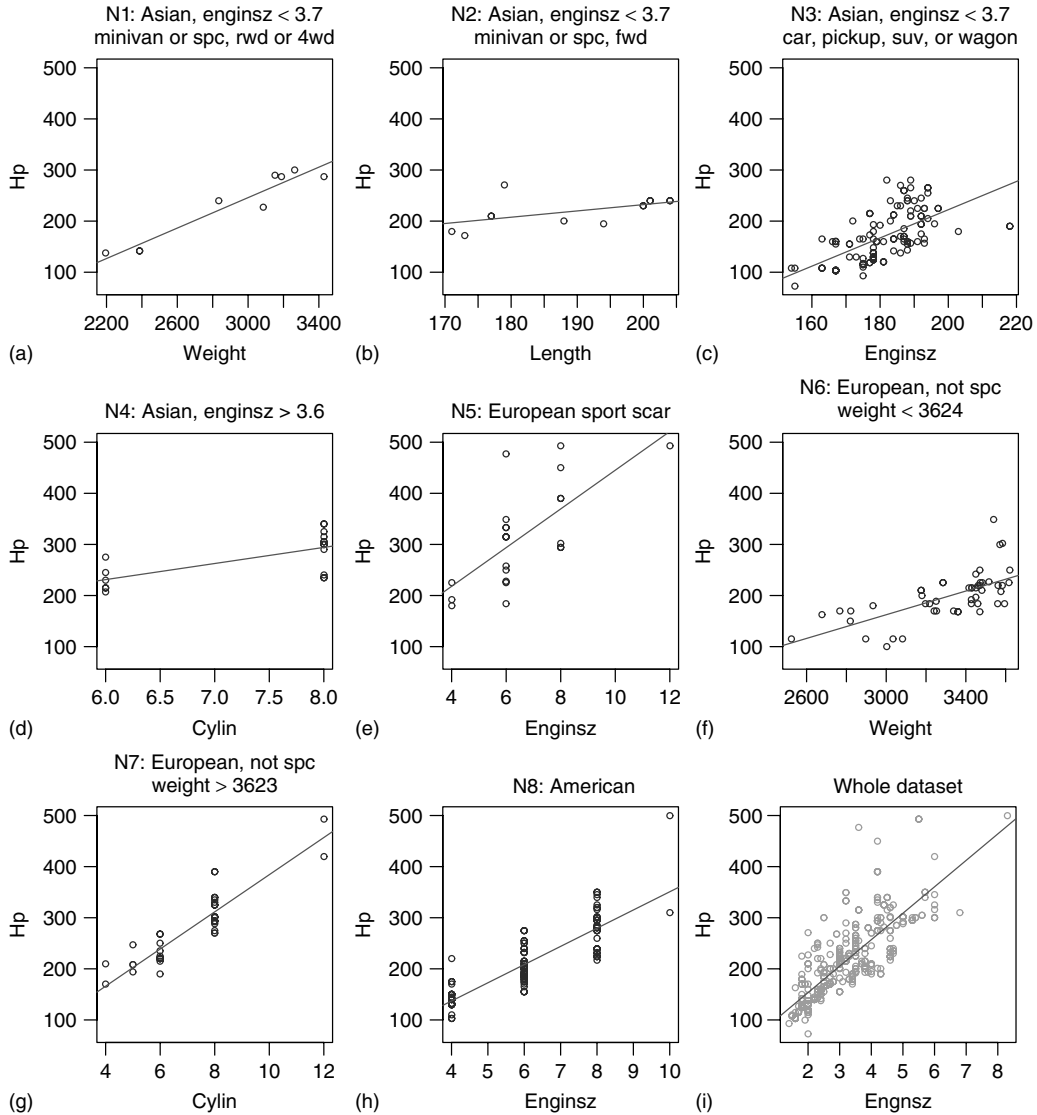


Figure 8 Data and simple linear regression models at the nodes of the GUIDE tree in Figure 7. The abbreviation *spc* stands for sports car. A plot of H_p versus $Engnsz$ for the whole dataset is given in (i) for comparison

have these capabilities. GUIDE is also restricted to univariate splits, but empirical studies with real datasets suggest that the prediction accuracy of its piecewise multiple linear models is, on average, comparable to that of the best methods, including spline models [9].

There are many other tree algorithms, but space restrictions prevent their discussion here. The interested reader is referred to CHAID [15, 16] for classification, RECPAM [17] for classification and

regression, M5 [18, 19] for regression, LOTUS [20] for logistic regression, and Bayesian CART [21, 22].

In terms of statistical theory, methods based on recursive partitioning are known to be *risk consistent* under certain regularity conditions. Specifically, the expected **mean squared error** of piecewise-constant regression trees and the expected misclassification cost of classification trees converge to the lowest possible values as the training sample size increases

[1]. Further, for piecewise linear regression trees, the predicted values converge to the unknown regression function values pointwise, at least for least-squares [23], logistic, Poisson [24], and quantile [25] regression.

Acknowledgments

CART is a registered trademark of California Statistical Software. This article was prepared with partial support from grants from the US Army Research Office and the National Science Foundation.

References

- [1] Breiman, L., Friedman, J.H., Olshen, R.A. & Stone, C.J. (1984). *Classification and Regression Trees*, Wadsworth, Belmont.
- [2] Therneau, T.M. & Atkinson, E.J. (1997). An introduction to recursive partitioning using the RPART routine, Technical Report 61, Mayo Clinic, Section of Statistics.
- [3] R Development Core Team. (2006). *R: A Language and Environment for Statistical Computing*, R Foundation for Statistical Computing, Vienna, ISBN 3-900051-07-0.
- [4] Morgan, J.N. & Sonquist, J.A. (1963). Problems in the analysis of survey data, and a proposal, *Journal of the American Statistical Association* **58**, 415–434.
- [5] Fielding, A. (1977). Binary segmentation: the automatic detector and related techniques for exploring data structure, in *The Analysis of Survey Data, Volume I, Exploring Data Structures*, C.A., O'Muircheartaigh & C., Payne, eds, John Wiley & Sons, New York.
- [6] Doyle, P. (1973). The use of automatic interaction detector and similar search procedures, *Operational Research Quarterly* **24**, 465–467.
- [7] Loh, W.-Y. & Shih, Y.-S. (1997). Split selection methods for classification trees, *Statistica Sinica* **7**, 815–840.
- [8] Kim, H. & Loh, W.-Y. (2001). Classification trees with unbiased multiway splits, *Journal of the American Statistical Association* **96**, 589–604.
- [9] Kim, H., Loh, W.-Y., Shih, Y.-S. & Chaudhuri, P. (2007). A visualizable and interpretable regression model with good prediction power, *IIE Transactions*. **39**, 565–579.
- [10] Kim, H. & Loh, W.-Y. (2003). Classification trees with bivariate linear discriminant node models, *Journal of Computational and Graphical Statistics* **12**, 512–530.
- [11] Quinlan, J.R. (1993). *C4.5: Programs for Machine Learning*, Morgan Kaufmann, San Mateo.
- [12] Witten, I. & Frank, E. (2005). *Data Mining: Practical Machine Learning Tools and Techniques with JAVA Implementations*, 2nd Edition, Morgan Kaufmann, San Francisco, <http://www.cs.waikato.ac.nz/ml/weka>.
- [13] Lim, T.-S., Loh, W.-Y. & Shih, Y.-S. (2000). A comparison of prediction accuracy, complexity, and training time of thirty-three old and new classification algorithms, *Machine Learning* **40**, 203–228.
- [14] Loh, W.-Y. (2002). Regression trees with unbiased variable selection and interaction detection, *Statistica Sinica* **12**, 361–386.
- [15] Kass, G.V. (1975). Significance testing in automatic interaction detection (A.I.D.), *Applied Statistics* **24**, 178–189.
- [16] Kass, G.V. (1980). An exploratory technique for investigating large quantities of categorical data, *Applied Statistics* **29**, 119–127.
- [17] Ciampi, A., Hogg, S.A., McKinney, S. & Thiffault, J. (1988). RECPAM: a computer program for recursive partition and amalgamation for censored survival data and other situations frequently occurring in biostatistics, I: methods and program features, *Computer Methods and Programs in Biomedicine* **26**, 239–256.
- [18] Quinlan, J.R. (1992). Learning with continuous classes, in *Proceedings of the Australian Joint Conference on Artificial Intelligence*, World Scientific, Singapore, pp. 343–348.
- [19] Wang, Y. & Witten, I. (1997). Inducing model trees for continuous classes, in *Proceedings of the Poster Papers of the European Conference on Machine Learning*, Prague.
- [20] Chan, K.-Y. & Loh, W.-Y. (2004). LOTUS: an algorithm for building accurate and comprehensible logistic regression trees, *Journal of Computational and Graphical Statistics* **13**, 826–852.
- [21] Chipman, H., George, E.I. & McCulloch, R.E. (1998). Bayesian CART model search (with discussion), *Journal of the American Statistical Association* **93**, 935–960.
- [22] Denison, D.G., Mallick, B.K. & Smith, A.F.M. (1998). A Bayesian CART algorithm, *Biometrika* **85**, 363–377.
- [23] Chaudhuri, P., Huang, M.-C., Loh, W.-Y. & Yao, R. (1994). Piecewise-polynomial regression trees, *Statistica Sinica* **4**, 143–167.
- [24] Chaudhuri, P., Lo, W.-D., Loh, W.-Y. & Yang, C.-C. (1995). Generalized regression trees, *Statistica Sinica* **5**, 641–666.
- [25] Chaudhuri, P. & Loh, W.-Y. (2002). Nonparametric estimation of conditional quantiles using quantile regression trees, *Bernoulli* **8**, 561–576.

Related Articles

Data Mining in Quality and Reliability; Generalized Linear Models; Mean Square Error; Residuals; Statistical Analysis of Survey Data; Statistical Methods for Counts, Rates, and Proportions; Sum of Squares.

WEI-YIN LOH

Optimal Reliability Design– Algorithms and Comparisons

Introduction

Optimal reliability design problems are known to be NP hard [1]. Finding efficient optimization algorithms is always a hot spot in this field. Classification of the literature by **reliability optimization** techniques is summarized in Table 1. This paper reviews and compares the recent developments of heuristic algorithms, metaheuristic algorithms, exact methods, and other optimization techniques in optimal reliability design. Owing to their robustness and feasibility, metaheuristic algorithms, especially genetic algorithms (GAs), have been widely and successfully applied (*see Genetic Algorithms in Reliability*). To improve computation efficiency or to avoid premature convergence, an important part of this work has been devoted in recent years to developing hybrid GAs, which usually combine a GA with heuristic algorithms, simulation annealing methods, **neural network** techniques, or other local search methods. Though more computation effort is involved, exact methods are particularly advantageous for small problems, and their solutions can be used to measure the performance of the heuristic or metaheuristic methods [2]. No obviously superior heuristic method has been proposed, but several of them have been well combined with exact or metaheuristic methods to improve their computational efficiency.

Metaheuristic Methods

Metaheuristic methods inspired by natural phenomena usually include the GAs, tabu search, simulated annealing algorithm, and ant colony optimization method. ant colony optimization (ACO) has been recently introduced into optimal reliability design, and it is proving to be a very promising general method in this field. Altıparmak *et al.* [10] provide a comparison of metaheuristics for the optimal design of computer networks.

Table 1 Reference classification by reliability optimization methods

	Metaheuristic algorithms
ACO	[3–8]
GA	[9–39]
HGA	[40–50]
TS	[51]
SA	[10, 52–55]
IA	[56]
GDA	[57]
CEA	[58]
	Exact methods
	[59–67]
	Max–min approaches
	[68, 69]
	Heuristic methods
	[70–72]
	Dynamic programming
	[73, 74]

Ant Colony Optimization Method

ACO is one of the adaptive metaheuristic optimization methods developed by M. Dorigo for traveling salesman problems in [75]. It is inspired by the behavior of real-life ants that consistently establish the shortest path from their nest to food. The essential trait of the ACO algorithm is the combination of *a priori* information about the structure of a promising solution with *posteriori* information about the structure of previously obtained good solutions [76].

Liang and Smith [3] first develop an ant colony metaheuristic optimization method to solve the reliability–**redundancy** allocation problem for a **k-out-of-n**: G series system. The proposed ACO approach includes four stages:

- Construction stage: construct an initial solution by selecting component j for subsystem i according to its specific heuristic η_{ij} and pheromone trail intensity τ_{ij} , which also sets up the transition probability mass function P_{ij} .
- Evaluation stage: evaluate the corresponding system reliability and penalized system reliability providing the specified penalized parameter.
- Improvement stage: improve the constructed solutions through local search.
- Updating stage: update the pheromone value online and offline given the corresponding penalized system reliability and the controlling parameter for pheromone persistence.

2 Optimal Reliability Design– Algorithms and Comparisons

Nahas and Noureldath [5] present an application of the ant system in a reliability optimization problem to maximize the system reliability subject to the system budget, with multichoice constraints incorporated at each subsystem. It also combines a local search algorithm and a specific improvement algorithm that uses the remaining budget to improve the quality of a solution.

The ACO algorithm has also been applied to a multiobjective reliability optimization problem [6] and to the optimal design of multistate series–parallel power systems [4].

Hybrid Genetic Algorithm

GA is a population-based directed random search technique inspired by the principles of evolution. Although it provides only heuristic solutions, it can be effectively applied to almost all complex combinatorial problems, and, thereby, it has been employed in a large number of references as shown in Table 1. Gen and Kim [77] provide a state-of-the-art survey of GA-based reliability design.

To improve computational efficiency, or to avoid premature convergence, numerous researchers have been inspired to seek effective combinations of GAs with heuristic algorithms, simulation annealing methods, neural network techniques, steepest-descent methods, or other local search methods. The combinations are generally called hybrid GAs, and they represent one of the most promising developmental directions in optimization techniques.

Considering a complex system with a known system structure function, Zhao and Liu [49] provide a unified modeling idea for both active and cold-standby redundancy optimization problems. The model prohibits any mixture of component types within subsystems. Both the lifetime and the cost of redundancy components are considered as random variables, so stochastic simulation is used to estimate the system performance, including the mean lifetime, percentile lifetime, and reliability. To speed up the solution process, these simulation results become the training data for training a neural network to approximate the system performance. The trained neural network is finally embedded into a GA to form a hybrid intelligent algorithm for solving the proposed model. Later, Wattanapongsakorn and Levitan [54] used random fuzzy lifetimes as the basic parameters and employed a random fuzzy simulation to generate the training data.

Lee *et al.* [42] develop a two-phase NN-hGA in which NN is used as a rough search technique to devise the initial solutions for a GA. By bounding the broad continuous search space with the NN technique, the NN-hGA derives the optimum robustly. However, in some cases, this algorithm may require too much computational time to be practical. To improve the computation efficiency, Lee *et al.* [43] present an NN-flcGA to effectively control the balance between exploitation and exploration, which characterizes the behavior of GAs. The essential features of the NN-flcGA include the following:

- combination with an NN technique to devise initial values for the GA;
- application of a fuzzy logic controller when tuning strategy GA parameters dynamically; and
- incorporation of the revised simplex search method.

Later, Lee *et al.* [44] proposed a similar hybrid GA called f-hGA for the redundancy allocation problem of a series–parallel system. It is based on

- application of a fuzzy logic controller to automatically regulate the GA parameters;
- incorporation of the iterative hill climbing method to perform local exploitation around the near-optimum solution.

Hsieh and Hsieh [40] considered the optimal task allocation strategy and hardware redundancy level for a cycle-free distributed computing system so that the system cost during the period of task execution is minimized. The proposed hybrid heuristic combines the GA and the steepest-descent method. Later, Hsieh [41] seeks similar optimal solutions to minimize system cost under constraints on the hardware redundancy levels.

Tabu Search

Although Kuo [2] described the promise of tabu search, Kulturel-Konak *et al.* [51] first developed an tabu search (TS) approach with the application of near feasible threshold (NFT) [78] for reliability optimization problems. This method uses a subsystem-based tabu entry and dynamic length tabu list to reduce the sensitivity of the algorithm to selection of the tabu list length. The definition of the moves in this approach offers an advantage in efficiency, since

it does not require recalculating the entire system reliability, but only the reliability of the changed subsystem. The results of several examples demonstrate the superior performance of this TS approach in terms of efficiency and solution superiority when compared to that of a GA.

Other Metaheuristic Methods

Some other adaptive metaheuristic optimization methods inspired by activities in nature have also been proposed and applied in optimal reliability design. Chen and You [56] developed an approach based on immune algorithms inspired by the natural immune system of all animals. It analogizes antibodies and antigens as the solutions and objection functions, respectively. Rocco *et al.* [58] proposed a cellular evolutionary approach combining the multi-member evolution strategy with concepts from cellular automata [79] for the selection step. In this approach, the selection of the parents is performed only in the neighborhood in contrast to the general evolutionary strategy that searches for parents in the whole population. And a great deluge algorithm is extended and applied to optimize the reliability of complex systems in [57]. When both accuracy and speed are considered simultaneously, it is proved to be an efficient alternative to ACO and other existing optimization techniques.

Exact Methods

Unlike metaheuristic algorithms, exact methods provide exact optimal solutions though much more computation complexity is involved. The development of exact methods, such as the branch-and-bound approach and lexicographic search, has recently been concentrated on techniques to reduce the search space of discrete optimization methods.

Sung and Cho [66] consider a reliability–redundancy allocation problem in which multiple choice and resource constraints are incorporated. The problem is first transformed into a bicriteria nonlinear integer programming problem by introducing 0–1 variables. Given a good feasible solution, the lower reliability bound of a subsystem is determined by the product of the maximal component reliabilities of all the other subsystems in the solution, while the upper bound is determined by the maximal amount of available sources of this subsystem. A branch-and-bound

procedure, based on this reduced solution space, is then derived to search for the global optimal solution. Later, Sung and Cho [67] even combine the lower and upper bounds of the system reliability, which are obtained by variable relaxation and Lagrangian relaxation techniques, to further reduce the search space.

Also, with a branch-and-bound algorithm, Djerdjour and Rekab [59] obtain the upper bound of series–parallel system reliability from its continuous relaxation problem. The relaxed problem is efficiently solved by the greedy procedure described in [80], combining heuristic methods to make use of some slack in the constraints obtained from rounding down. This technique assumes that the objective and constraint functions are monotonically increasing. Ha and Kuo [61] present an efficient branch-and-bound approach for **coherent systems** based on a one-neighborhood local maximum obtained from the steepest-ascent heuristic method. Numerical examples of a bridge system and a hierarchical series–parallel system demonstrate the advantages of this proposed algorithm in flexibility and efficiency.

Apart from the branch-and-bound approach, Prasad and Kuo [63] present a partial enumeration method for a wide range of complex optimization problems based on a lexicographic search. The proposed upper bound of system reliability is very useful in eliminating several inferior feasible or infeasible solutions as shown in either big or small numerical examples. It also shows that the search process described in [84] does not necessarily give an exact optimal solution due to its logical flows.

Lin and Kuo [62] develop a strong Lin and Kuo heuristic to search for an ideal allocation through the application of the reliability importance (*see Component Reliability Importance*). It concludes that, if there exists an invariant optimal allocation for a system, the optimal allocation is to assign component reliabilities according to *B*-importance ordering. This Lin and Kuo heuristic can provide an exact optimal allocation.

Assuming the existence of a convex and differential reliability cost function $C_i(y_{ij})$, $y_{ij} = \log(1 - r_{ij})$ for all components j in any subsystem i , Elegbede *et al.* [60] prove that the components in each subsystem of a series–parallel system must have identical reliability for the purpose of cost minimization. The solution of the corresponding unconstrained problem provides the upper bound of the cost, while a doubly minimization problem gives its

lower bound. With these results, the algorithm ECAY, which can provide either exact or approximate solutions depending on different stop criteria, is proposed for series–parallel systems.

Other Optimization Techniques

For series–parallel systems, Ramirez-Marquez *et al.* [69] formulate the reliability–redundancy optimization problem with the objective of maximizing the minimal subsystem reliability. Assuming linear constraints, this problem is equivalent to a linear formulation through an easy logarithm transformation [81], which can be solved by readily available commercial software. It can also serve as a surrogate for traditional reliability optimization problems by sequentially solving a series of max–min subproblems. Lee *et al.* [68] present a comparison between the Nakagawa and Nakashima method [82] and the max–min approach used by Ramirez-Marquez from the standpoint of solution quality and computational complexity. The experimental results show that the max–min approach is superior to the Nakagawa and Nakashima method in terms of solution quality in small-scale problems, but the analysis of its computational complexity demonstrates that the max–min approach is inferior to other greedy heuristics.

You and Chen [72] develop a heuristic approach inspired by the greedy method and a GA. The structure of this algorithm includes (a) randomly generating a specified population size number of minimum workable solutions, (b) assigning components either according to the greedy method or to the random selection method, and (c) improving solutions through an inner-system and intersystem solution revision process.

Kevin and Sancho [73] apply a hybrid dynamic programming/depth-first search algorithm to redundancy allocation problems with more than one constraint. Given the tightest upper bound, the knapsack relaxation problem is formulated with only one constraint, and its solution $f_1(b)$ is obtained by a dynamic programming method. After choosing a small specified parameter e , the depth-first search technique is used to find all near-optimal solutions with objectives between $f_1(b)$ and $f_1(b) - e$. The optimal solution is given by the best feasible solution among all of the near-optimal solutions. Also with dynamic programming, Yalaoui *et al.* [74] consider a reliability–redundancy allocation problem

in series–parallel systems with components chosen among a finite set. This pseudopolynomial YCC algorithm is composed of two steps: the solution of the subproblems, one for each subsystem, and the global resolution using previous results. It shows that the solutions converge quickly toward the optimum as a function of the required precision.

Comparisons and Discussions of Algorithms Reported in the Literature

In this section, we provide a comparison of several heuristic or metaheuristic algorithms reported in the literature. The compared numerical results are from the GA in Coit and Smith [78], the ACO in Liang and Smith [3], TS in Kulturel-Konak *et al.* [51], linear approximation in Hsieh [81], the immune algorithm (IA) in Chen and You [56] and the heuristic method in You and Chen [72]. The 33 variations of the Fyffe *et al.* problem, as devised by Nakagawa and Miyazaki [83], are used to test their performance, where different types are allowed to reside in parallel. In this problem set, the cost constraint is maintained at 130 and the weight constraint varies from 191 to 159.

As shown in literature, ACO [3], TS [51], IA [56], and heuristic methods [72] generally yield solutions with a higher reliability. When compared to GA [78],

- ACO [3] is reported to consistently perform well over different problem sizes and parameters and improve on GA's random behavior.
- TS [51] results in a superior performance in terms of best solutions found and reduced variability and greater efficiency based on the number of objective function evaluations required.
- IA [56] finds better or equally good solutions for all 33 test problems, but the performance of this IA-based approach is sensitive to value-combinations of the parameters, whose best values are case-dependent and only based upon the experience from preliminary runs and more CPU time is taken by IAs.
- The best solutions found by heuristic methods [72] are all better than, or as good as, the well-known best solutions from other approaches. With this method, the average CPU time for each problem is within 8 s.
- In terms of solution quality, the linear approximation approach in [81] is inferior, but very efficient,

and the CPU time for all of the test problems is within 1 s.

- If a decision-maker is considering the max–min approach as a surrogate for system reliability maximization, the max–min approach [69] is shown to be capable of obtaining a close solution (within 0.22%), but it is unknown whether this performance will continue as problem sizes become larger.

For all the optimization techniques mentioned above, it might be hard to discuss about which tool is superior because in different design problems or even in the same problem with different parameters, these tools will perform varyingly.

Generally, if computational efficiency is of most concern to designers, linear approximation or heuristic methods can obtain competitive feasible solutions within a very short time (few seconds), as in [72, 81]. The proposed linear approximation [81] is also easy to implement with any LP software. However, the main limitation of those reported approaches is that the constraints must be linear and separable.

Owing to their robustness and feasibility, metaheuristic methods such as GA and the recently developed TS and ACO could be successfully applied to almost all NP-hard reliability optimization problems. However, they cannot guarantee the optimality and sometimes can suffer from premature convergence because they have many unknown parameters and neither use a prior knowledge nor exploit local search information. Compared to traditional metaheuristic methods, a set of promising algorithms, hybrid GAs [40–44, 49, 50], are attractive since they retain the advantages of GAs in robustness and feasibility, but significantly improve their computational efficiency and searching ability in finding a global optimum by combining heuristic algorithms, neural network techniques, steepest-descent methods, or other local search methods.

For reliability optimization problems, exact solutions are not necessarily desirable because it is generally difficult to develop exact methods for reliability optimization problems that are equivalent to methods used for nonlinear integer programming problems [2]. However, exact methods may be particularly advantageous when the problem is not large. More importantly, such methods can be used to measure the performance of heuristic or metaheuristic methods.

Conclusions and Discussions

Heuristic, metaheuristic, and exact methods are significantly applied in optimal reliability design. Recently, many advances in metaheuristics and exact methods have been reported. Particularly, a new metaheuristic method called *ant colony optimization* has been introduced and demonstrated to be a very promising general method in this field. Hybrid GAs may be the most important recent development among the optimization techniques since they retain the advantages of GAs in robustness and feasibility but significantly improve their computational efficiency.

Optimal reliability design has attracted many researchers, who have produced hundreds of publications since 1960. Owing to the increasing complexity of practical engineering systems and the critical importance of reliability in these complex systems, this still seems to be a very fruitful area for future research. From the view of optimization techniques, there are opportunities for improved effectiveness and efficiency of reported ACO, TS, IA, and great deluge algorithm (GDA), while some new metaheuristic algorithms such as harmony search algorithm and particle swarm optimization may offer excellent solutions for reliability optimization problems. Hybrid optimization techniques are another very promising general developmental direction in this field. They may combine heuristic methods, NN or some local search methods with all kinds of metaheuristics to improve computational efficiency or with exact methods to reduce search space. We may even be able to combine two metaheuristic algorithms such as GA and SA or ACO in future research on optimal reliability design.

Acknowledgments

This work was partly supported by NSF projects DMI-0429176. It is republished with the permission of IEEE *Transactions on Systems, Man, and Cybernetics, Part A: Systems and Humans*.

References

- [1] Chen, M.S. (1987). On the computational complexity of reliability redundancy allocation in a series system, *Operations Research Letters* **SE-13**, 582–592.

- [2] Kuo, W. & Prasad, V.R. (2000). An annotated overview of system-reliability optimization, *IEEE Transactions on Reliability* **49**, 487–493.
- [3] Liang, Y.C. & Smith, A.E. (2004). An ant colony optimization algorithm for the redundancy allocation problem, *IEEE Transactions on Reliability* **53**, 417–423.
- [4] Massim, Y., Zeblah, A., Meziane, R., Benguediab, M. & Ghouraf, A. (2005). Optimal design and reliability evaluation of multi-state series-parallel power systems, *Nonlinear Dynamics* **40**, 309–321.
- [5] Nahas, N. & Nourelfath, M. (2005). Ant system for reliability optimization of a series system with multiple-choice and budget constraints, *Reliability Engineering and System Safety* **87**, 1–12.
- [6] Shelokar, P.S., Jayaraman, V.K. & Kulkarni, B.D. (2002). Ant algorithm for single and multiobjective reliability optimization problems, *Quality and Reliability Engineering International* **18**, 497–514.
- [7] Yin, P. & Wang, J. (2006). Ant colony optimization for the nonlinear resource allocation problem, *Applied Mathematics and Computation* **174**, 1438–1453.
- [8] Zhao, J., Liu, Z. & Dao, M. (2007). Reliability optimization using multiobjective ant colony system approaches, *Reliability Engineering and System Safety* **92**, 109–120.
- [9] Agarwal, M. & Gupta, R. (2006). Genetic search for redundancy optimization in complex systems, *Journal of Quality in Maintenance Engineering* **12**, 338–353.
- [10] Altiparmak, F., Dengiz, B. & Smith, A.E. (2003). Optimal design of reliable computer networks: a comparison of meta heuristics, *Journal of Heuristics* **9**, 471–487.
- [11] Busacca, P.G., Marseguerra, M. & Zio, E. (2001). Multi-objective optimization by genetic algorithms: application to safety systems, *Reliability Engineering and System Safety* **72**, 59–74.
- [12] Coit, D.W. & Smith, A.E. (1998). Redundancy allocation to maximize a lower percentile of the system time-to-failure distribution, *IEEE Transactions on Reliability* **47**, 79–87.
- [13] Coit, D.W. & Smith, A.E. (2002). Genetic algorithm to maximize a lower-bound for system time-to-failure with uncertain component Weibull parameters, *Computers and Industrial Engineering* **41**, 423–440.
- [14] Eleghbede, C. & Adjallah, K. (2003). Availability allocation to repairable systems with genetic algorithms: a multi-objective formulation, *Reliability Engineering and System Safety* **82**, 319–330.
- [15] Hsieh, Y.C., Chen, T.C. & Bricker, D.L. (1998). Genetic algorithm for reliability design problems, *Microelectronics Reliability* **38**, 1599–1605.
- [16] Knoak, A., Coit, D.W. & Smith, A.E. (2006). Multi-objective optimization using genetic algorithms: a tutorial, *Reliability Engineering and System Safety* **91**, 992–1007.
- [17] Levitin, G., Lisnianski, A., Ben-Haim, H. & Elmakis, D. (1998). Redundancy optimization for series-parallel multi-state systems, *IEEE Transactions on Reliability* **47**, 165–172.
- [18] Levitin, G. & Lisnianski, A. (1998). Structure optimization of power system with bridge topology, *Electric Power Systems Research* **45**, 201–208.
- [19] Levitin, G. & Lisnianski, A. (1999). Joint redundancy and maintenance optimization for multistate series-parallel systems, *Reliability Engineering and System Safety* **64**, 33–42.
- [20] Levitin, G. & Lisnianski, A. (2000). Survivability maximization for vulnerable multi-state systems with bridge topology, *Reliability Engineering and System Safety* **70**, 125–140.
- [21] Levitin, G. (2001). Redundancy optimization for multi-state system with fixed resource-requirements and unreliable sources, *IEEE Transactions on Reliability* **50**, 52–59.
- [22] Levitin, G. & Lisnianski, A. (2001). Reliability optimization for weighted voting system, *Reliability Engineering and System Safety* **71**, 131–138.
- [23] Levitin, G. & Lisnianski, A. (2001). Structure optimization of multi-state system with two failure modes, *Reliability Engineering and System Safety* **72**, 75–89.
- [24] Levitin, G. & Lisnianski, A. (2001). Optimal separation of elements in vulnerable multi-state systems, *Reliability Engineering and System Safety* **73**, 55–66.
- [25] Levitin, G. & Lisnianski, A. (2001). A new approach to solving problems of multi-state system reliability optimization, *Quality and Reliability Engineering International*, **17**, 93–104.
- [26] Levitin, G. (2002). Optimal allocation of multi-state retransmitters in acyclic transmission networks, *Reliability Engineering and System Safety* **75**, 73–82.
- [27] Levitin, G. (2002). Optimal allocation of elements in a linear multi-state sliding window system, *Reliability Engineering and System Safety* **76**, 245–254.
- [28] Levitin, G. (2002). Optimal series-parallel topology of multi-state system with two failure modes, *Reliability Engineering and System Safety* **77**, 93–107.
- [29] Levitin, G. (2003). Optimal allocation of multi-state elements in linear consecutively connected systems with vulnerable nodes, *European Journal of Operational Research* **150**, 406–419.
- [30] Levitin, G. (2003). Optimal allocation of multistate elements in a linear consecutively-connected system, *IEEE Transactions on Reliability* **52**, 192–199.
- [31] Levitin, G. & Lisnianski, A. (2003). Optimizing survivability of vulnerable series-parallel multi-state systems, *Reliability Engineering and System Safety* **79**, 319–331.
- [32] Levitin, G. (2003). Optimal multilevel protection in series-parallel systems, *Reliability Engineering and System Safety* **81**, 93–102.
- [33] Levitin, G., Dai, Y., Xie, M. & Poh, K.L. (2003). Optimizing survivability of multi-state systems with multilevel protection by multi-processor genetic algorithm, *Reliability Engineering and System Safety* **82**, 93–104.
- [34] Levitin, G. (2005). Uneven allocation of elements in linear multi-state sliding window system, *European Journal of Operational Research* **163**, 418–433.

-
- [35] Levitin, G. (2005). Optimal structure of fault-tolerant software systems, *Reliability Engineering and System Safety* **89**, 286–295.
 - [36] Lisnianski, A., Levitin, G. & Ben-Haim, H. (2000). Structure optimization of multi-state system with time redundancy, *Reliability Engineering and System Safety* **67**, 103–112.
 - [37] Marseguerra, M., Zio, E., Podofilini, L. & Coit, D.C. (2005). Optimal design of reliable network systems in presence of uncertainty, *IEEE Transactions on Reliability* **54**, 243–253.
 - [38] Yamachi, H., Yamachi, H., Tsujimura, Y., Kambayashi, Y. & Yamamoto, H. (2006). Multi-objective genetic algorithm for solving N -version program design problem. *Reliability Engineering and System Safety* **91**, 1083–1094.
 - [39] Yang, J.E., Hwang, M.J., Sung, T.Y. & Jin, Y. (1999). Application of genetic algorithm for reliability allocation in nuclear power plants, *Reliability Engineering and System Safety* **65**, 229–238.
 - [40] Hsieh, C.C. & Hsieh, Y.C. (2003). Reliability and cost optimization in distributed computing systems, *Computers and Operations Research* **30**, 1103–1119.
 - [41] Hsieh, C.C. (2003). Optimal task allocation and hardware redundancy policies in distributed computing system, *European Journal of Operational Research* **147**, 430–447.
 - [42] Lee, C.Y., Gen, M. & Kuo, W. (2002). Reliability optimization design using hybridized genetic algorithm with a neural network technique, *IEICE Transactions on Fundamentals of Electronics, Communications and Computer Sciences* **E84-A**, 627–637.
 - [43] Lee, C.Y., Yun, Y. & Gen, M. (2002). Reliability optimization design using hybrid NN-GA with fuzzy logic controller, *IEICE Transactions on Fundamentals of Electronics, Communications and Computer Sciences* **E85-A**, 432–447.
 - [44] Lee, C.Y., Gen, M. & Tsujimura, Y. (2002). Reliability optimization design for coherent systems by hybrid GA with fuzzy logic controller and local search, *IEICE Transactions on Fundamentals of Electronics, Communications and Computer Sciences* **E85-A**, 880–891.
 - [45] Sasaki, M. & Gen, M. (2003). A method of fuzzy multi-objective nonlinear programming with GUB structure by hybrid genetic algorithm, *International Journal of Smart Engineering Design* **5**, 281–288.
 - [46] Sasaki, M. & Gen, M. (2003). Fuzzy multiple objective optimal system design by hybrid genetic algorithm, *Applied Soft Computing* **2**, 189–196.
 - [47] Taguchi, T., Yokota, T. & Gen, M. (1998). Reliability optimal design problem with interval coefficients using hybrid genetic algorithms, *Computers and Industrial Engineering* **35**, 373–376.
 - [48] Taguchi, T. & Yokota, T. (1999). Optimal design problem of system reliability with interval coefficient using improved genetic algorithms, *Computers and Industrial Engineering* **37**, 145–149.
 - [49] Zhao, R. & Liu, B. (2003). Stochastic programming models for general redundancy-optimization problems, *IEEE Transactions on Reliability* **52**, 181–191.
 - [50] Zhao, R. & Liu, B. (2004). Redundancy optimization problems with uncertainty of combining randomness and fuzziness, *European Journal of Operation Research* **157**, 716–735.
 - [51] Kulturel-Konak, S., Smith, A.E. & Coit, D.W. (2003). Efficiently solving the redundancy allocation problem using tabu search, *IIE Transactions* **35**, 515–526.
 - [52] Kim, H., Bae, C. & Park, D. (2006). Reliability-redundancy optimization using simulated annealing algorithms, *Journal of Quality in Maintenance Engineering* **12**, 354–363.
 - [53] Suman, B. (2003). Simulated annealing-based multi-objective algorithm and their application for system reliability, *Engineering Optimization* **35**, 391–416.
 - [54] Wattanapongsakorn, N. & Levitan, S.P. (2004). Reliability optimization models for embedded systems with multiple applications, *IEEE Transactions on Reliability* **53**, 406–416.
 - [55] Zafiropoulos, E.P. & Dialynas, E.N. (2004). Reliability and cost optimization of electronic devices considering the component failure rate uncertainty, *Reliability Engineering and System Safety* **84**, 271–284.
 - [56] Chen, T.C. & You, P.S. (2005). Immune algorithms-based approach for redundant reliability problems with multiple component choices, *Computers in Industry* **56**, 195–205.
 - [57] Ravi, V. (2004). Optimization of complex system reliability by a modified great deluge algorithm, *Asia-Pacific Journal of Operational Research* **21**, 487–497.
 - [58] Rocco, C.M., Moreno, J.A. & Carrasquero, N. (2000). A cellular evolution approach applied to reliability optimization of complex systems, in Proceedings Annual Reliability and Maintainability Symposium, Los Angeles, California. pp. 210–215.
 - [59] Djerdjour, M. & Rekab, K. (2001). A branch and bound algorithm for designing reliable systems at a minimum cost, *Applied Mathematics and Computation* **118**, 247–259.
 - [60] Elegbede, C., Chu, C., Adjallah, K. & Yalaoui, F. (2003). Reliability allocation through cost minimization, *IEEE Transactions on Reliability* **52**, 106–111.
 - [61] Ha, C. & Kuo, W. (2006). Reliability redundancy allocation: an improved realization for nonconvex nonlinear programming problems, *European Journal of Operational Research* **171**, 24–38.
 - [62] Lin, F. & Kuo, W. (2002). Reliability importance and invariant optimal allocation, *Journal of Heuristics* **8**, 155–172.
 - [63] Prasad, V.R. & Kuo, W. (2000). Reliability optimization of coherent systems, *IEEE Transactions on Reliability* **49**, 323–330.
 - [64] Prasad, V.R., Kuo, W. & Kim, K.O. (2001). Maximization of a percentile life of a series system through component redundancy allocation, *IIE Transactions* **33**, 1071–1079.

- [65] Ruan, N. & Sun, X. (2006). An exact algorithm for cost minimization in series reliability systems with multiple component choices, *Applied Mathematics and Computation* **181**, 732–741.
- [66] Sung, C.S. & Cho, Y.K. (1999). Branch-and-bound redundancy optimization for a series system with multiple-choice constraints, *IEEE Transactions on Reliability* **48**, 108–117.
- [67] Sung, C.S. & Cho, Y.K. (2000). Reliability optimization of a series system with multiple-choice and budget constraints, *European Journal of Operational Research* **127**, 158–171.
- [68] Lee, H., Kuo, W. & Ha, C. (2003). Comparison of max-min approach and NN method for reliability optimization of series-parallel system, *Journal of System Science and Systems Engineering* **12**, 39–48.
- [69] Ramirez-Marquez, J.E., Coit, D.W. & Konak, A. (2004). Redundancy allocation for series-parallel systems using a max-min approach, *IIE Transactions* **36**, 891–898.
- [70] Coit, D.W. & Konak, A. (2006). Multiple weighted objective heuristic for the redundancy allocation problem, *IEEE Transactions on Reliability* **55**, 551–558.
- [71] Ramirez-Marquez, J.E. & Coit, D.W. (2004). A heuristic for solving the redundancy allocation problem for multi-state series-parallel system, *Reliability Engineering and System Safety* **83**, 341–349.
- [72] You, P.S. & Chen, T.C. (2005). An efficient heuristic for series-parallel redundant reliability problems, *Computer and Operations Research* **32**, 2117–2127.
- [73] Kevin, Y.K.N.G. & Sancho, N.G.F. (2001). A hybrid dynamic programming/depth-first search algorithm with an application to redundancy allocation, *IIE Transactions* **33**, 1047–1058.
- [74] Yalaoui, A., Châtelet, E. & Chu, C. (2005). A new dynamic programming method for reliability & redundancy allocation in a parallel-series system, *IEEE Transactions on Reliability* **54**, 254–261.
- [75] Dorigo, M. (1992). *Optimization learning and natural algorithm*, Ph.D. Thesis, Politecnico di Milano, Italy.
- [76] Maniezzo, V., Gambardella, L.M. & De Luigi, F. (2004). Ant colony optimization, in *New Optimization Techniques in Engineering*, G.C. Onwubolu & B.V. Babu, eds, Springer-Verlag Berlin Heidelberg, pp. 101–117.
- [77] Gen, M. & Kim, J.R. (1999). GA-based reliability design: state-of-the-art survey, *Computers and Industrial Engineering* **37**, 151–155.
- [78] Coit, D.W. & Smith, A.E. (1996). Reliability optimization of series-parallel systems using a genetic algorithm, *IEEE Transactions on Reliability* **45**, 254–260.
- [79] Wolfram, S. (1984). Cellular automata as models of complexity, *Nature* **311**, 419–424.
- [80] Djerdjour, M. (1997). An enumerative algorithm framework for a class of nonlinear programming problems, *European Journal of Operation Research* **101**, 101–121.
- [81] Hsieh, Y.C. (2002). A linear approximation for redundant reliability problems with multiple component choices, *Computers and Industrial Engineering* **44**, 91–103.
- [82] Kuo, W., Hwang, C.L. & Tillman, F.A. (1978). A note on heuristic method for in optimal system reliability, *IEEE Transactions on Reliability* **27**, 320–324.
- [83] Nakagawa, Y. & Miyazaki, S. (1981). Surrogate constraints algorithm for reliability optimization problems with two constraints, *IEEE Transactions on Reliability* **R-30**, 175–180.
- [84] Misra, K.B. (1991). An algorithm to solve integer programming problems: an efficient tool for reliability design, *Microelectronics and Reliability* **31**, 285–294.

Further Reading

- Gen, M. & Yun, Y. (2006). Soft computing approach for reliability optimization: state-of-art survey, *Reliability Engineering and System Safety* **91**, 1008–1026.
- Ha, C. & Kuo, W. (2005). Multi-path approach for reliability-redundancy allocation using a scaling method, *Journal of Heuristics* **11**, 201–237.
- Ha, C. & Kuo, W. (2006). Multi-paths iterative heuristic for redundancy allocation: the tree heuristic, *IEEE Transactions on Reliability* **55**, 37–43.
- Kim, K.O. & Kuo, W. (2003). Percentile life and reliability as a performance measures in optimal system design, *IIE Transactions* **35**, 1133–1142.
- Kumar, U.D. & Knezevic, J. (1998). Availability based spare optimization using renewal process, *Reliability Engineering and System Safety* **59**, 217–223.
- Kuo, W., Chien, K. & Kim, T. (1998). *Reliability, Yield and Stress Burn-in*, Kluwer.
- Kuo, W., Prasad, V.R., Tillman, F.A. & Hwang, C.L. (2001). *Optimal Reliability Design: Fundamentals and Applications*, Cambridge University Press.
- Kuo, W. & Zuo, M.J. (2003). *Optimal Reliability Modeling: Principles and Applications*, John Wiley & Sons.
- Levitin, G. & Lisnianski, A. (1999). Importance and sensitivity analysis of multi-state systems using the universal generating function method, *Reliability Engineering and System Safety* **65**, 271–282.
- Lin, F., Kuo, W. & Hwang, F. (1999). Structure importance of consecutive-k-out-of-n systems, *Operations Research Letters* **25**, 101–107.
- Marseguerra, M. & Zio, E. (2004). System design optimization by genetic algorithm, *IEEE Transactions on Reliability* **53**, 424–434.
- Misra, K.B. (1986). On optimal reliability design: a review, *System Science* **12**, 5–30.
- Tillman, F.A., Hwang, C.L. & Kuo, W. (1977). Optimization techniques for system reliability with redundancy- a review, *IEEE Transactions on Reliability* **R-26**, 148–152.

WAY KUO AND RUI WAN

Monitoring of Safety in EU Railways

ERA, the Ambitious Projects on Safety of a Newcomer, under Directive 2004/49/EC

Railway is the safest land transport mode and one of the overall safest transport modes. Directive 2004/49/EC, on railway safety, has been put in place to support the creation of an integrated European railway area by facilitating market opening through the harmonization of safety management and regulation. Directive 2004/49/EC will also contribute to improve safety by: a cultural shift from a deterministic to a risk-based approach, the establishment of safety management systems, the introduction of a system for independent accident investigation, as well as by moving from self-regulation to regulation by independent authorities (National Safety Authorities – NSAs). This aims to create a basis for mutual trust between member states (MSs), mainly through the development of common and transparent methods to monitor safety performance, to set targets and to manage the introduction of significant changes in railway systems with a risk-based systematic approach.

Collection and Publication of Information on Safety

Purpose

Safety information is collected mainly to provide input to the development of an objective tool for setting common safety targets (CSTs) and assessing their achievement. CSTs will be developed in subsequent sets, on the basis of safety performances in MSs and their possible economic impact in terms of costs and benefits. The European Railway Agency (ERA) is in charge of these projects, carried out in conjunction with industry and MSs in working groups of experts.

Development of Common Safety Indicators (CSIs) and Accidents Investigation

The working group on common safety indicators (CSIs) is developing and defining indicators structured as follows:

Safety performances

- type of accidents
- fatalities and injuries classified in categories of people
- precursors to accidents.

Economic aspects of safety performances

- societal benefits and costs of improving and worsening safety performances.

CSIs are expressed in total of events and consequences (fatalities and injuries) per unit of time (1 year), normalized mainly by traffic performances (train*km, passenger*km).

The flow of information and data exchange on safety issues between the concerned parties is completed by the activity of accidents investigation, carried out by National Investigation Bodies (NIBs) and communicated to ERA; the aim of this is to better understand the causes of accidents, which, together with the aforementioned statistics on accidents, precursors to accidents, and economic aspects of safety, provides the necessary information to ERA to set CSTs at EU level (see Figure 1).

Data Sources of CSIs

ERA relies on data from two sources, Eurostat and NSAs. Eurostat has been collecting EU-wide harmonized data on “type of accidents, related consequences, and traffic performances”, since 2004; NSAs have been publishing data on CSIs every year, starting from 2006 and according to a common format developed with ERA. Therefore, both Eurostat and NSAs will collect data on accidents and related consequences. In addition, NSAs will also report on precursors to accidents and economic aspects of safety. This double collection of accidents data should be temporary and it is a consequence of the further development, introduced by Directive 2004/49/EC, of the type of accidents classification and categorization of people; ERA and Eurostat are working together to eliminate these differences and to harmonize **data collection** processes.

Consistency between Eurostat and NSAs. ERA has informally collected 2004 and 2005 data on CSIs from NSAs, with the view to improving data reliability through the appraisal of consistency between

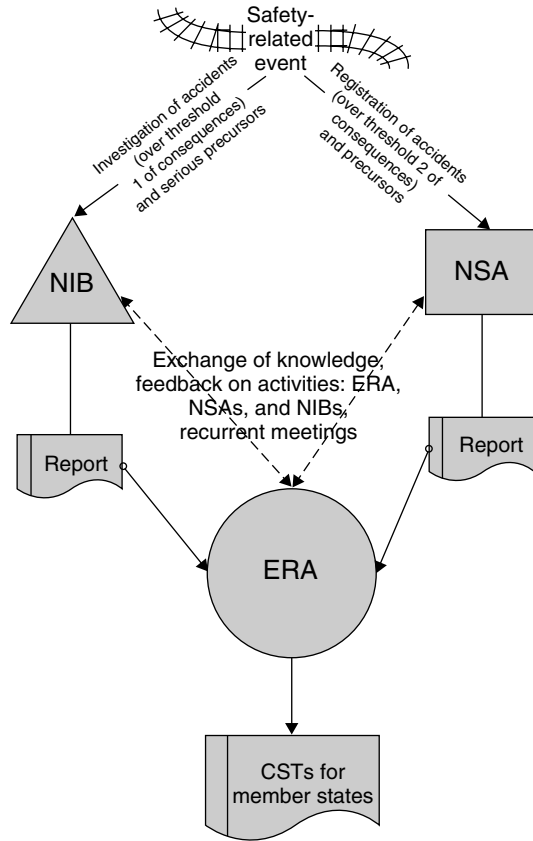


Figure 1 Flow of information to set CSTs

NSAs and Eurostat statistics. The differences identified between NSAs and Eurostat statistics have been registered and sent to both sources, inviting them to further discuss with each other, to eliminate existing inconsistencies.

Collection and Analysis of Time Series of National Data on Railway Safety. In the context of ERA's work to develop CSTs, it has been necessary to complement data delivered by Eurostat and data on CSIs with longer **time series** of national data on railway safety. This in order to provide appropriate input to the analysis aimed at identifying for each MS the

current safety performance of national railway systems in terms of risk. The need to look at long time series of data derives from the high degree of oscillation which data on accidents and related consequences (i.e., fatalities and serious injuries) typically show from year to year, making difficult the identification of the current risk levels attributable to national railway systems. Being, in fact, the risk calculated as a product of consequences of accidents and their respective frequencies, the oscillation of data both on frequency and consequences is systematically transmitted to risk estimates. This oscillation is the cumulative effect of rarity and statistical randomness of the events under consideration. Accidents with multiple fatalities and injuries represent rare events, therefore they are considered as high-impact and low-frequency events and generate outliers in time series.

Different techniques have been considered for smoothing down the outliers due to multiple fatality accidents and to calculate the underlying, intrinsic risk levels of national railway systems. These techniques use moving averages or moving **percentiles** (i.e., moving median or upper quartile), which respectively average and cut out isolated outliers in the same moving n -year time window.

The random behavior of accident statistics over time, cumulated with the intrinsic rarity of significant accidents, may also lead to recursive zero values in time series, especially for small, modern railway networks with low traffic volumes. This makes it difficult to estimate the underlying risk levels of these systems through robust statistical methods, especially when adequate time series of data are not available. Empirical solutions are needed in these cases for identifying and targeting risk levels, such as the adoption of empirical tolerance intervals $[0, \sigma]$ for the risk levels of the system, where σ is a predetermined standard deviation from a mean which is assumed to be equal to zero.

ANGELO PIRA, ROBERTO PIAZZA AND
JANE RAJAN

Bias in Modeling and Monitoring Health Outcomes

A Brief Background on Control Charts

Suppose we wish to monitor the rate of a particular health outcome in a specified population or health facility, there are several **control charts** that can be used, including Shewhart, cumulative sum (CUSUM), **exponentially weighted moving average** (EWMA), and variable life-adjusted display (VLAD). Each of these is now a standard textbook procedure and is described elsewhere in this publication.

For charts such as the Shewhart chart, signals occur on the basis of the extremity, or unusualness, of the most recent observation only. For other charts, such as the EWMA and CUSUM, where past data are taken into account, the occurrence of clusters of extreme events can also cause signals. The Shewhart chart can also be adapted to incorporate past data, by use of runs rules such as those suggested by the Western Electric Company [1].

Monitoring the Incidence of a Disease

Consider a scenario where we wish to monitor the incidence of a disease in a population. At any point where a signal occurs on the control chart used, we would most likely want to revise our original estimate of the rate of infection, from which we were monitoring for change, so that the control chart could be recalibrated and monitoring continue anew. The question as to what is the new infection rate is part of the larger question: “why was there a signal?”

The example data: re-estimating an infection rate following a signal

Results from a Shewhart and an EWMA chart. Figure 1 shows infection rates of methicillin-resistant *Staphylococcus aureus* (MRSA) in a UK NHS Trust at 6-monthly intervals in the period April 2001 to September 2005 [2]. The observed rates are the number of blood samples taken that tested positive

for MRSA divided by the number of bed-days in thousands in a 6-monthly interval. The observed rates are shown in black, accompanied by the outer pair of dotted lines, which are 99.8% Shewhart limits. These Shewhart limits are calculated using the rate across the first four 6-monthly intervals as a “target” rate (central dotted line) and assuming counts follow a Poisson distribution where the denominator is bed-days in thousands. An EWMA of the data is shown in gray, as are the EWMA chart limits (the inner pair of dotted lines). The smoothing parameter value used is $\kappa = 0.5$. The chart limits are 95% gamma limits centered on the target rate, with precision (or effective number of observations) equal to $(1 - \kappa)^{-1} = 2$ [3].

In the **quality control** literature, the state of the process being monitored is often described as either “in control” or “out of control” at any one time, and signals indicating that the process has moved away from an in-control state, then, are described as either false or true. In practice, however, it may be difficult to know whether a signal was truly warranted. In Figure 1, suppose that we have just observed the signal in interval 6 (October 2003–March 2004). How might we go about re-estimating the true underlying infection rate at this juncture?

Results from a CUSUM, and general concerns about bias. Figure 2 shows a Poisson CUSUM for the same data, testing for a 20% increase in the rate from the level over the period April 2001–September 2002 (intervals 1–4), and thus testing whether a rate of 0.26 is more plausible than a rate of 0.22. There is a signal in period 6 for this chart too. We might

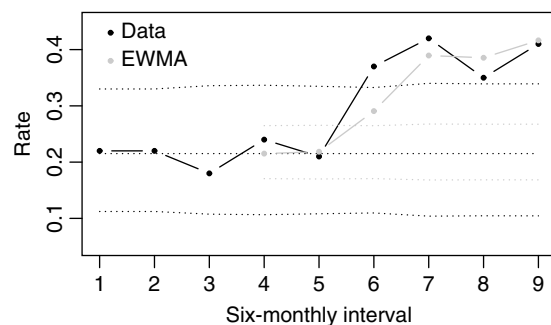


Figure 1 MRSA rates in a UK NHS Trust. Central dotted line: rate over first four periods; outer lines: transformed 99.8% Poisson limits, given bed-days; inner lines: 95% gamma limits, given bed-days

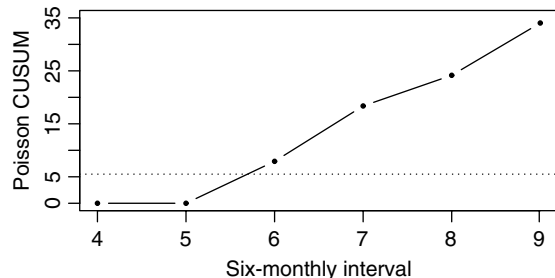


Figure 2 Poisson CUSUM testing for a 20% increase in MRSA infection rate in a UK. NHS Trust from the level in April 2001–September 2002, with boundary set at $h = 5.5$

re-estimate the underlying rate following the signal using only the data since the chart was last at zero, because this is where the chart indicates that the process was last “in control”. For the example, this gives an estimate of 0.37. Alternatively, we could use all the available data up to signal (giving an estimate of 0.24), but may be mixing data from different underlying sources and hence biasing the estimate of the rate from the “true” source. We might attempt to exclude outlying data points. For instance, if it were decided that the observation in interval 6 was anomalous for some reason, the estimate would be equal to the original estimate, 0.22. In any case, it remains difficult to determine whether the underlying infection rate has changed, and likewise whether the original estimate was reasonable. In this example, the subsequent data points are consistent with the observed rate upon signal (the rate over the last four periods is 0.39) and the CUSUM continues to accumulate in favor of the alternative with respect to the null.

Estimation following an early signal (one before the planned study end) of benefit of the treatment is of particular interest in the application of sequential **clinical trials** [4, 5]. A quick decision about whether the apparent benefit is real (that is, whether patients would benefit from future treatment) is desirable, but the estimate of the effect of the treatment may be biased if it is based on insufficient data, or a chance run of unlikely data under a null, or negligible effect. To be careful that the estimate of the effect of the treatment is not overstated, it might be **bias-adjusted** [6] or made more conservative according to some prior or external information [7].

More about the EWMA as both a chart and an estimator

In contrast to the CUSUM, which accumulates the log-likelihood ratio in favor of one hypothesis over another, the EWMA is a direct estimator for the rate [3]. As illustrated in Figure 1, it is formed by a weighted average on all past data (in the dataset used), where the data are weighted exponentially according to their contemporaneity.

It appears as though the use of the EWMA or other estimators in a testing scheme provides a way of simultaneously estimating and testing. Charts based on smoothers like the EWMA have inertia [8], because they do not give full weight to the most recent observation, and so are robust to **outliers** yet at the same time less adaptive to new process conditions. If we are primarily interested in estimating the rate directly after signals of a change in rate, inertia, combined with the nature of the chart stopping rule, restricts the estimate.

In the event of the EWMA chart signalling, the chart statistic is necessarily larger than the threshold. So, for an endemic disease, EWMA values following signal may be upwardly biased in terms of predicting future typical behaviour of the disease if the signal is caused by an isolated cluster of observations that, in retrospect, is difficult to account for. If the underlying infection rate suddenly increases, the EWMA value just after signal may underestimate the increase. In the example, the EWMA value in period 6 (0.29), if biased, could be an overestimate of the rate if the true rate is closer to the observed rate over the first four periods (0.22), or an underestimate if it is closer to that over the last four periods (0.39).

Though the EWMA value is already smoothed (or, results from an assumption that the underlying true process is smooth), we might wish to adjust it toward the initial target or an external value from peer processes. In that way, the re-estimate of the infection rate would be conservative and protect against subsequent regression to the mean if the true underlying rate were unchanged. However, it may happen that after adjustment the value is below the signaling threshold. For example, the EWMA value in period 6 in the example might be adjusted halfway toward the initial target if there were external information, with weight equal to the observed data, that the rate was unchanged. This

would lead to a supported value for the rate, at that point, of 0.25.

If the adjusted value is taken to be the current estimate for the infection rate and is, as in our example, lower than the threshold for a signal (in the example, this is 0.26), we might ask whether we consider that a signal has still occurred. In both cases, inferences on the estimation and testing front would seem to contradict each other, because we either have a signal and an estimate lower than the prespecified threshold, or do not have a signal but have an estimate higher than the threshold (if adjustment is only carried out given the presence of a signal). A solution would perhaps be to adjust the estimate no further than the threshold, but this would again present as a restriction on the parameter space of the estimate.

Conclusion

To conclude, combined estimation and **hypothesis testing** through the use of control charts may produce results that are difficult to analyze because of the issue of bias, and lead to possible contradictions in the inferences made. Hawkins and Olwell note that “even a somewhat biased estimate is better than no estimate at all.” This is reminiscent of Box’s words, “All models are wrong, but some are useful”. Perhaps faith in tangible models should override concerns of bias from truth.

References

- [1] Western Electric Company. (1956). *Statistical Quality Control Handbook*, Western Electric Corporation, Indianapolis.
- [2] The Department of Health. (2006). MRSA surveillance data April 2001–September 2005. =<http://www.dh.gov.uk/assetRoot/04/12/79/13/04127913.pdf> (accessed on 08/12/06).
- [3] Grigg, O.A. & Spiegelhalter, D.J. (2007). A simple risk-adjusted exponentially weighted moving average, *Journal of the American Statistical Association* **102**(477), 140–152.
- [4] Whitehead, J. (1997). *The Design and Analysis of Sequential Clinical Trials*, 3rd Edition, Ellis Horwood, Chichester.
- [5] Emmerson, S.S. & Fleming, T.R. (1990). Parameter estimation following group sequential hypothesis testing, *Biometrika* **77**, 875–892.
- [6] Todd, S., Whitehead, J. & Facey, K.M. (1996). Point and interval estimation following a sequential clinical trial, *Biometrika* **83**, 453–461.
- [7] Pocock, S.J. & Hughes, M.D. (1989). Practical problems in interim analyses, with particular regard to estimation, *Controlled Clinical Trials* **10**, 209–221.
- [8] Woodall, W.H. & Mahmoud, M.A. (2005). The inertial properties of quality control charts, *Technometrics* **47**, 425–436.

Further Reading

Grigg, O.A. (2004). *Risk-adjusted monitoring and smoothing in medical contexts*, Ph.D. thesis, King’s College, Cambridge.

OLIVIA A. J. GRIGG

Mixture Models, Multivariate: Aging and Dependence

Introduction

In this article, we maintain the same language and notation of the article **Aging and Positive Dependence** and consider vectors of lifetimes $(T_1, \dots, T_n)$ with joint survival function of the specific form

$$\bar{F}(t_1, \dots, t_n) = \int_{R^m} \prod_{j=1}^n \bar{G}_j(t_j | \theta_1, \dots, \theta_m) \times dH(\theta_1, \dots, \theta_m) \quad (1)$$

where, for some $m \geq 1$, H is an m -dimensional probability distribution and, for any vector $(\theta_1, \dots, \theta_m)$, $\bar{G}_j(\cdot | \theta_1, \dots, \theta_m)$ is a one-dimensional survival function ($j = 1, 2, \dots, n$). We can think of H as the joint distribution of a nonobservable random vector $\Theta = (\Theta_1, \dots, \Theta_m)$ such that $T_1, \dots, T_n$ are conditionally independent, given Θ . Generally, neither $\Theta_1, \dots, \Theta_m$ nor $T_1, \dots, T_n$ are independent random variables.

Models of the form in equation (1) are known as *multivariate mixture models*; they arise in several different fields and in a great variety of situations of applied interest. In particular, they can be used to model heterogeneity; for a discussion about the concept of heterogeneity in the field of reliability (see e.g., [1] the references indicated therein).

It is quite natural to analyze multivariate mixture models in the frame of a Bayesian approach. In such a frame, we will discuss here some aspects of aging and positive dependence.

$T_1, \dots, T_n$ are conditionally independent, generally nonidentically distributed, given Θ . Special multivariate mixture models of interest are defined by the conditions:

1. $\bar{G}_j(t | \theta_1, \dots, \theta_m)$, as a function of $(\theta_1, \dots, \theta_m)$, depends only on θ_j and
2. $\bar{G}_j(t | \theta_1, \dots, \theta_m)$, as a function of t and θ_j , has a general form that does not depend on j .

Such conditions can be summarized by writing that $\bar{G}_j(t | \theta_1, \dots, \theta_m)$ is of the form

$$\bar{G}_j(t | \theta_1, \dots, \theta_m) = \rho(t | \theta_j) \quad (2)$$

and describe special *frailty models*.

In several cases of interest, as an additional condition on these frailty models, we have that $\Theta_1, \dots, \Theta_m$ are exchangeable. As a remarkable feature, in these cases, we have heterogeneity and similarity at the same time: also $T_1, \dots, T_n$ turn out to be exchangeable. However, $T_1, \dots, T_n$ generally do not have an identical conditional distribution, given Θ . We point out then that we obtain in this way exchangeable variables, for which conditional independence is not, however, paired with identical conditional distributions.

Specially relevant cases of exchangeability are met when $\Theta_1, \dots, \Theta_m$ are independent and identically distributed, or when $P\{\Theta_1 = \dots = \Theta_m\} = 1$. The latter case is the basic one encountered in Bayesian statistics: $T_1, \dots, T_n$ are conditionally independent and identically distributed, given a same scalar parameter Θ ; in the former case, instead, $T_1, \dots, T_n$ are themselves independent and identically distributed. Several other models of interest, however, can also be met in between these two extreme cases.

For these models we also notice that, irrespective of the special form of dependence among $\Theta_1, \dots, \Theta_m$, the one-dimensional marginal distribution of T_j ($j = 1, 2, \dots, n$) is a mixture distribution with a survival function of the form

$$\bar{F}(t) = \int_R \rho(t | \theta) dH^{(1)}(\theta) \quad (3)$$

$H^{(1)}$ being the one-dimensional marginal of H .

In a general multivariate mixture model, as in equation (1) (not necessarily exchangeable), the marginal distributions of T_j ($j = 1, 2, \dots, n$), are univariate mixture distributions with survival functions given by

$$\bar{F}_j(t) = \int_{R^m} \bar{G}_j(t | \theta_1, \dots, \theta_m) dH(\theta_1, \dots, \theta_m) \quad (4)$$

In applications, there are different kinds of situations that give rise to univariate mixture distributions. The difference between different situations is not, sometimes, very clear. From a probabilistic point of view, however, one can simply say that these

2 Mixture Models, Multivariate: Aging and Dependence

situations are described by distributions obtained as marginals of different multivariate mixture models, just with different types of mixing distributions $H(\theta_1, \dots, \theta_m)$.

Aspects of Ageing

A fact of recurrent interest in quality and reliability is related with aging properties of (univariate) mixtures as in equation (4) and can be roughly summarized as follows: *whereas negative aging tends to be maintained by mixtures, positive aging tends to be lost under mixtures*. This fact has been made precise by formulating different probabilistic results; at the same time, it has been explained in terms of several heuristic interpretations or empirical analysis. It is, in particular, well-known that the property of decreasing failure rate (DFR) is maintained under mixtures (see e.g., [2]); on the other hand, several interpretations have been given of the fact that increasing failure rate (IFR) may, on the contrary, be lost under mixtures. In a reliability framework, well-known papers on this topic are [3, 4], the latter being in particular based on a Bayesian approach.

Here we want to illustrate the basic idea behind all that, still maintaining a Bayesian approach. To this purpose, we assume simplifying conditions of stochastic monotonicity that more easily allow us to explain the reasons why DFR is maintained by mixture, while IFR can be lost (see also [5–8]). Same type of arguments can also be, respectively, worked out for other notions of positive and negative aging.

Consider a lifetime T and a single positive random variable Θ ; T is the (observable) lifetime of a component and Θ is a nonobservable factor that affects the probability distribution of T ; Θ can have e.g., a meaning such as the quality or the frailty of the component. For the sake of technical simplicity, we consider cases where the joint probability distribution of (T, Θ) admits a joint probability, described as follows: the marginal distribution of Θ , that is concentrated over $[0, +\infty)$, admits a density $\pi(\cdot)$ and the conditional distributions of T , given $\{\Theta = \theta\}$ admits conditional failure rates $r(t|\theta)$, for $\theta \geq 0$; denote by $R(t|\theta)$ the conditional cumulative failure rate of T given $\{\Theta = \theta\}$:

$$R(t|\theta) = \int_0^t r(u|\theta) du \quad (5)$$

The marginal distribution of T , obtained by integrating out θ , has then survival function, density function, and failure rate, respectively given by

$$\begin{aligned} \bar{F}(t) &= \int_0^\infty \exp\{-R(t|\theta)\} \pi(\theta) d\theta, \\ f(t) &= \int_0^\infty r(t|\theta) \exp\{-R(t|\theta)\} \pi(\theta) d\theta \end{aligned} \quad (6)$$

and

$$\lambda(t) = \frac{f(t)}{\bar{F}(t)} = \int_0^\infty r(t|\theta) \pi_t(\theta) d\theta \quad (7)$$

where we set

$$\pi_t(\theta) = \frac{\exp\{-R(t|\theta)\} \pi(\theta)}{\int_0^\infty \exp\{-R(t|\theta)\} \pi(\theta) d\theta} \quad (8)$$

As an immediate application of Bayes formula, $\pi_t(\theta)$ can be interpreted as the conditional density of Θ , given the event $\{T > t\}$. In many practical situations, it can be reasonable to assume that $r(t|\theta)$ is, for any fixed t , an increasing function of θ : the bigger the value taken by θ , the stochastically smaller is T , in the hazard order sense; it can be easily shown, in such a case, that the distribution with density π_t is stochastically decreasing in t .

Compare now, for fixed t and for two different densities π', π'' , the expressions

$$\int_0^\infty r(t|\theta) \pi'(\theta) d\theta, \int_0^\infty r(t|\theta) \pi''(\theta) d\theta \quad (9)$$

Since, for fixed t , $r(t|\theta)$ is increasing in θ , we have

$$\int_0^\infty r(t|\theta) \pi'(\theta) d\theta \leq \int_0^\infty r(t|\theta) \pi''(\theta) d\theta \quad (10)$$

whenever the distribution with density π' is smaller than the distribution with density π'' , in the sense of the ordinary stochastic order.

Put now, for two values $0 < t' < t''$, $\pi'(\theta) = \pi_{t'}(\theta)$, $\pi''(\theta) = \pi_{t''}(\theta)$. For a same fixed value $\tau > 0$, we can then write, since π_t is stochastically decreasing in t and $r(t|\theta)$ is increasing in θ ,

$$\int_0^\infty r(\tau|\theta) \pi_{t'}(\theta) d\theta \geq \int_0^\infty r(\tau|\theta) \pi_{t''}(\theta) d\theta \quad (11)$$

Suppose T is conditionally DFR given θ , i.e. $r(t|\theta)$ is decreasing in t . By fixing e.g., $\tau = t'$, we

can immediately conclude that T is DFR; in fact

$$\lambda(t') = \int_0^\infty r(t'|\theta)\pi_{t'}(\theta) d\theta \geq \int_0^\infty r(t''|\theta)\pi_{t''}(\theta) d\theta = \lambda(t'') \quad (12)$$

Suppose now that T is IFR in the conditional distribution. Even if, for $t' < t''$, it is then $r(t'|\theta) \leq r(t''|\theta)$, we cannot generally conclude that $\lambda(t') \leq \lambda(t'')$, since in any case $\pi_{t'}$ is stochastically larger than $\pi_{t''}$.

This explains why the distribution of T , which is a mixture of IFR distributions and has the density function given by

$$f(t) = \int_0^\infty r(t|\theta) \exp\{-R(t|\theta)\} \pi(\theta) d\theta \quad (13)$$

may potentially be not IFR!

In passing we have to notice however that, under special conditions, a mixture of IFR distribution can still be IFR (see [9] and [10]).

This specific phenomenon that a mixture of IFR distribution is not generally IFR (and that the mixture can, however, be IFR, under specific conditions), has also been discovered and extensively analyzed in several other fields, related with reliability, such as demography, population dynamics, survival analysis, and medicine, (see e.g., [11–13]); it has also been remarked that the phenomenon is somewhat related to the *Simpson's paradox* (e.g., [13, 14]).

More generally, the aging behavior of univariate distributions obtained by mixtures has often been found to be relevant in different applications and related statistical analysis; the main problem is that many different types of behavior are possible when failure rates, of the distributions which are mixed, are not decreasing. Therefore, studies on the aging properties of mixtures, in an analytical viewpoint, have quite a long history (see e.g., reviews in [15–18]). Specific aspects about this theme concern the asymptotic (or tail) behavior (see e.g., [15, 19]) and conditions that give rise to *bathtub* distributions (see [20] and references therein).

An important problem in reliability is to understand what are the situations under which **burn-in** procedures are justified. From a Bayesian viewpoint, deciding whether it is convenient to adopt some *burn-in* procedure is just a decision problem (see

e.g., the introductory discussion in [21]) and establishing the duration of burn-in is just an optimal stopping problem (see [22, 23] and references therein).

Roughly speaking, burn-in is useful in cases of *infant mortality*. For this reason a special interest has been devoted to the theme of negative aging; the issues of concern are understanding what negative aging really means, which analytical conditions imply it, and what real reliability models can manifest such a behavior. In this respect, there are several ways to interpret the fact that, in some practical situations where mixtures appear, burn-in procedures are actually convenient (see also e.g., [6, 23, 24] and references therein). From a mathematical point of view, one can say that infant mortality and univariate negative aging often arise in the case of mixtures.

One can imagine, however, that multivariate mixture models also manifest some kind of negative multivariate aging. Multivariate properties of the Bayesian type (see in particular [18, 23, 25, 26] and the brief sketch in **Aging and Positive Dependence**) have been dealt with in [27] for these models in the exchangeable case; sufficient conditions that guarantee some multivariate negative aging have been proved therein, under suitable assumptions of positive dependence for the vector Θ and stochastic monotonicity of the conditional distributions of $T_1, \dots, T_n$, given Θ .

As one conclusion of main interest for this article, we point out that in cases of (univariate) mixture distributions and of multivariate mixture models one can encounter some form of negative aging.

Positive Dependence Properties

Let us now pass to analyze positive dependence for multivariate mixture models. In this analysis, we have to distinguish between the two cases where the unobservable factor Θ , w.r.t. which $T_1, \dots, T_n$ are conditionally independent, is a random scalar or a random vector. It is a general fact, however, that stochastic monotonicity of the different lifetimes $T_1, \dots, T_n$ with respect to Θ , combined with suitable positive dependence properties of Θ , gives rise to positive dependence properties for $\mathbf{T} = (T_1, \dots, T_n)$. This fact admits a rather direct heuristic interpretation in a Bayesian view (see also **Aging and Positive**

Dependence, the section titled “Notions of Positive Dependence”). When Θ reduces to a scalar Θ (i.e., if $m = 1$), no assumption is needed for Θ .

A basic result in this direction was given by Jogdeo in [28] (see also [29]).

Theorem 1 *Let $T_1, \dots, T_n$ be conditionally independent, given Θ and each of them stochastically increasing (or decreasing) in Θ , in the ordinary stochastic order. Then $T_1, \dots, T_n$ are associated.*

By postulating stronger types of stochastic monotonicity for $T_1, \dots, T_n$, one can obtain stronger conditions of positive dependence. Examples of that are the following results, proved in [30].

Theorem 2 *Let $T_1, \dots, T_n$ be conditionally independent, given Θ and each of them stochastically decreasing (or increasing) in Θ , in the hazard rate order. Then $\mathbf{T} = (T_1, \dots, T_n)$ is weakened by failure.*

For the following result, we need to assume that the support $\mathcal{D}$ of the distribution of $\mathbf{T}$ is a lattice, i.e., $\mathbf{u}, \mathbf{v} \in \mathcal{D} \implies \mathbf{u} \wedge \mathbf{v} \in \mathcal{D}, \mathbf{u} \vee \mathbf{v} \in \mathcal{D}$.

Theorem 3 *Let $T_1, \dots, T_n$ be conditionally independent, given Θ and each of them stochastically decreasing (or increasing) in Θ , in the likelihood ratio order. Then $\mathbf{T}$ is multivariate totally positive of order 2.*

Let us consider now cases in which $T_1, \dots, T_n$ are conditionally independent, given a vector Θ . In these cases, the space R^m of possible values of Θ is to be thought of as (partially) ordered in the component-wise sense. Theorem 1 above can be generalized as follows (see again [28]).

Theorem 4 *Let $T_1, \dots, T_n$ be conditionally independent, given Θ and each of them stochastically decreasing (or increasing) in Θ , in the ordinary stochastic order; furthermore, let Θ be associated. Then $T_1, \dots, T_n$ are associated.*

Theorem 1 is clearly a special case of Theorem 4, in fact, a scalar random variable Θ can be seen as a special case of an associated random vector. Analogously, Theorems 2 and 3 have also been respectively generalized in suitable ways to the case of a vector Θ ; in such results, Θ is to be assumed multivariate totally positive of order 2 (see [30]).

For some other aspects of dependence properties for mixture models, see also [31].

A general concept of positive dependence for a random vector $(X_1, \dots, X_r)$ is that of *dependence by total positivity with degree $(k_1, \dots, k_r)$* (see in particular [32, 33]). This concept is important since several different notions of positive dependence can just be obtained as special cases of it. In the paper [34], Kirmani and Kochar extended the aforementioned results in order to prove *dependence by total positivity with degree $(k_1, \dots, k_r)$* for vectors of lifetimes $T_1, \dots, T_n$, jointly distributed according to multivariate mixture modes.

Conclusions

We mentioned that multivariate mixture models are of remarkable interest for applications and that, in particular, they can be used to model heterogeneity. Some other examples of applications are mentioned, e.g., in [30].

We want to conclude this article, however, by stressing that multivariate mixture models are also relevant for the following aspect of conceptual character: they can likely manifest positive dependence along with some form of negative aging.

References

- [1] Arjas, E. & Bhattacharjee, M. (2004). Modelling heterogeneity in repeated failure time data: a hierarchical Bayesian approach, *Mathematical Reliability: An Expository Perspective*, Kluwer Academic Publishers, Boston, pp. 71–86.
- [2] Barlow, R.E. & Proschan, F. (1975). *Statistical Theory of Reliability and Life Testing*, Holt, Rinehart and Winston, New York.
- [3] Proschan, F. (1963). Theoretical explanation of observed decreasing failure rate, *Technometrics* **5**, 375–383.
- [4] Barlow, R.E. (1985). A Bayesian explanation of an apparent failure rate paradox, *IEEE Transactions on Reliability* **R34**, 107–108.
- [5] Spizzichino, F. (1992). Reliability decision problems under conditions of ageing, *Bayesian Statistic 4*, J. Bernardo, J. Berger, A.P. Dawid & A.F.M. Smith, eds, Clarendon Press, Oxford, pp. 803–811.
- [6] Block, H.W. & Savits, T.H. (1997). Burn-in, *Statistical Science* **12**, 1–19.
- [7] Lynn, N.J. & Singpurwalla, N.D. (1997). “Burn-in” makes us feel good, Comment to [6].
- [8] Finkelstein, M.S. & Esaulova, V. (2001). Modeling a failure rate for a mixture of distribution functions, *Probability in the Engineering and Informational Sciences* **15**, 383–400.

- [9] Lynch, J.D. (1999). On conditions for mixtures of increasing failure rate distributions to have an increasing failure rate, *Probability in the Engineering and Informational Sciences* **13**, 33–36.
- [10] Block, H.W., Li, Y. & Savits, T.H. (2003). Preservation of properties under mixture, *Probability in the Engineering and Informational Sciences* **17**, 205–212.
- [11] Vaupel, J.W. & Yashin, A.I. (1985). Heterogeneity's ruses: some surprising effects of selection on population dynamics, *The American Statistician* **39**, 176–185.
- [12] Cohen, J.E. (1986). An uncertainty principle in demography and the unisex issue, *The American Statistician* **40**, 32–39.
- [13] Gurland, J. & Sethuraman, J. (1995). How pooling failure data may reverse increasing failure rates, *Journal of the American Statistical Association* **90**, 1416–1423.
- [14] Scarsini, M. & Spizzichino, F. (1999). Simpson-type paradoxes, dependence and aging, *Journal of Applied Probability* **36**, 119–131.
- [15] Shaked, M. & Spizzichino, F. (2001). *Mixtures and Monotonicity of Failure Rate Functions*, *Handbook of Statistics*, N. Balakrishnan & C.R. Rao, eds, Elsevier Science, Amsterdam, Vol. 20, pp. 185–198.
- [16] Block, H.W. & Savits, T.H. (2004). The shape of failure rate mixtures, *Mathematical Reliability: An Expository Perspective*, Kluwer Academic Publishers, Boston, pp. 197–206.
- [17] Finkelstein, M.S. (2003). A model of aging and a shape of the observed force of mortality, *Lifetime Data Analysis* **9**, 93–109.
- [18] Lai, C.-D. & Xie, M. (2006). *Stochastic Ageing and Dependence for Reliability*, Springer-Verlag, New York.
- [19] Block, H.W. & Joe, H. (1997). Tail behavior of the failure rate functions of mixtures, *Lifetimes Data Analysis* **3**, 269–288.
- [20] Lai, C.D., Xie, M. & Murthy, D.N.P. (2001). *Bathtub-Shaped Failure Rate Life Distributions*, *Handbook of Statistics*, N. Balakrishnan & C.R. Rao, eds, Elsevier Science, Amsterdam, Vol. 20, pp. 69–104.
- [21] Clarotti, C.A. & Spizzichino, F. (1990). Bayesian burn-in decision procedures, *Probability in the Engineering and Informational Sciences* **4**, 437–445.
- [22] Jensen, U. & Spizzichino, F. (2004). The burn-in problem - a discussion of sequential stop and go strategies *Mathematical Reliability: An Expository Perspective*, Kluwer Academic Publishers, Boston, pp. 207–229.
- [23] Spizzichino, F. (2001). *Subjective Probability Models for Lifetimes*, Chapman and Hall/CRC, Boca Raton.
- [24] Block, H.W., Mi, J. & Savits, T.H. (1993). Burn-in and mixed populations, *Journal of Applied Probability* **30**, 692–702.
- [25] Barlow, R.E. & Mendel, M.B. (1992). de Finetti-type representations for life distributions, *Journal of the American Statistical Association* **87**, 1116–1122.
- [26] Bassan, B. & Spizzichino, F. (1999). Stochastic comparison for residual lifetimes and Bayesian notions of multivariate ageing, *Advances in Applied Probability* **31**, 1078–1094.
- [27] Spizzichino, F. & Torrisi, G. (2001). Multivariate negative aging in an exchangeable model, *Statistics and Probability Letters* **55**, 71–82.
- [28] Jogdeo, K. (1978). On a probability bound of Marshall and Olkin, *Annals of Statistics* **6**, 232–234.
- [29] Szekli, R. (1995). *Stochastic Ordering and Dependence in Applied Probability Lecture Notes in Mathematics* 97, Springer-Verlag, New York.
- [30] Shaked, M. & Spizzichino, F. (1998). Positive dependence properties of conditionally independent random lifetimes, *Mathematics of Operation Research* **23**, 944–959.
- [31] Belzunce, F. & Semeraro, P. (2004). Preservation of positive and negative orthant dependence concepts under mixtures and applications, *Journal of Applied Probability* **41**, 961–974.
- [32] Shaked, M. (1977). A concept of positive dependence for exchangeable random variables, *Annals of Statistics* **5**, 505–515.
- [33] Lee, M.L.T. (1985). Dependence by total positivity, *Annals of Probability* **13**, 572–582.
- [34] Khaledi, B.E. & Kochar, S. (2001). Dependence properties of multivariate mixture distributions and their applications, *Annals of the Institute of Statistical Mathematics* **53**, 620–630.

Related Articles

Aging and Positive Dependence; Bayesian Reliability Analysis; Bayesian Reliability Demonstration; Stochastic Orders and Aging; Subjective Life Models.

FABIO SPIZZICHINO

Repair Data, Sets of: How to Graph, Analyze, and Compare

Introduction

Purpose

Many products accumulate repeated repairs and repair costs over time. Analysis of such recurrence data requires special statistical models and methods not covered in basic reliability books. This article presents a simple and informative model and plot for analyzing data on numbers or costs of repeated repairs of a random sample of units. The plot is illustrated with transmission repair data from preproduction cars on a track test. This article also presents a method for comparing two such datasets, illustrated with automatic and manual transmissions. Computer programs that calculate and make the plots and comparisons with confidence limits are surveyed.

Information

The plot given here is a nonparametric graphical estimate of the population mean cumulative number or cost of repairs per unit *versus* population age. This estimate can be used to

1. evaluate whether the population repair (or cost) rate increases or decreases with age; this is useful for product retirement and **burn-in** decisions;
2. predict future numbers and costs of repairs;
3. compare two random samples from different designs, production periods, maintenance policies, environments, operating conditions, and so on;
4. reveal unexpected information and insight, an important advantage of plots.

Life Data

Most models and analysis methods for product reliability data are suitable for life data. Life data on a unit consists of a *single* failure time, the unit's end of life. Some sample units may be still unfailed. Then each such unit has a failure time that is beyond its current running time. Such lifetimes are modeled with

a life distribution, such as the Weibull distribution, which is fitted to the mix of failures and nonfailures. For example, Nelson [1] and Abernethy [2] present such models and data analyses, which are unsuited to repair data, which can have any number of repair events for a unit, not just one.

Overview

The next section describes typical repair data. The subsequent section "The Population Model and Its Mean Cumulative Function" presents the basic population model and its mean cumulative function (MCF) for the number or cost of repairs. The section "Estimate and Plot of the MCF" shows how to calculate and plot a sample estimate of the MCF from data from units with a mix of ages. "How to Interpret and Use a Plot" explains how to use and interpret such plots. The section "Comparison of Sets of Repair Data" provides a graphical comparison and confidence limits for the difference of two sample MCFs. The last section references further information on analysis of recurrence data.

Repair Data

Transmission Data

This section describes typical repair data from a sample of units. Table 1 displays a small set of typical repair data on a sample of 34 preproduction cars with an automatic transmission run on a severe test track. For each car the data consist of the car's mileage at each transmission repair and the latest observed mileage. For example, for car 024, the data are a repair at 7068 mi and its latest mileage 26 744 + mi; here + indicates the distance the car has been observed. Nelson [3–6] gives repair data on blood analyzers, residential heat pumps, window air conditioners, power supplies, and turbines. He also gives applications from other fields with repeated events data, including recurrent disease episodes, childbirths, auto accidents, and so on. Information sought from the transmission data includes

1. the mean cumulative number of repairs per car by 24 000 track miles (equivalent to 132 000 customer miles, the design life);

2 Repair Data, Sets of: How to Graph, Analyze, and Compare

Table 1 Automatic transmission repair data^a

Car	Mileage (+ latest)	
024	7068	26 744+
026	28	13 809+
027	48	1440
029	530	25 660+
031	21 762+	
032	14 235+	
034	1388	21 133+
035	21 401+	
098	21 876+	
107	5094	18 228+
108	21 691+	
109	20 890+	
110	22 486+	
111	19 321+	
112	21 585+	
113	18 676+	
114	23 520+	
115	17 955+	
116	19 507+	
117	24 177+	
118	22 854+	
119	17 844+	
120	22 637+	
121	375	19 607+
122	19 403+	
123	20 997+	
124	19 175+	
125	20 425+	
126	22 149+	
129	21 144+	
130	21 237+	
131	14 281+	
132	8250	21 974+
133	19 250	21 888+

^a © John Wiley & Sons Ltd, 1998

- the behavior of the population repair rate – increasing or decreasing as the population ages?

Censoring

A sample unit's latest observed age is called its (right) *censoring age*, because the unit's repair history to the right beyond that age is censored (unknown) at the time of the data analysis. Usually, such right-censoring ages differ and complicate the data analysis and require the method described below. A unit may have no failures; in such cases its data is just its censoring age. Other units may have one, two, three, or more repairs before its censoring age.

Age

Here “age” (or “time”) means any useful measure of unit usage, e.g., mileage, days, cycles, months, and so on.

Data Display

Figure 1 is a useful display of the data but not an analysis. Here each car history is displayed as a horizontal line segment whose length is the observed car mileage, and its censoring age is denoted by a vertical line. For a car, each transmission repair is marked with an × at the corresponding mileage on its line. This simple display suggests that earlier production has more repairs, which is expected, since production is being improved during start-up. A more informative analysis is given in the sections titled “Estimate and Plot of the MCF” and “How to Interpret and Use a Plot”.

The Population Model and its Mean Cumulative Function

Model

Most needed information on the repair behavior of a population of units is given by the population mean cumulative function (MCF) *versus* age t . The MCF (defined below) is a feature of the following population model, which has no censoring. Censoring is a feature of the sample data, not of the population model. At a particular age t , population unit i has accumulated a total cost (or number) of repairs $Y_i(t)$. These cumulative unit totals $Y_i(t)$ usually differ. Figure 2 depicts the population of such uncensored unit *cumulative history functions* for cost; they are shown as smooth curves for easy viewing and extend to any age of interest. In reality, such a history function $Y_i(t)$ is a staircase function where the rise of each step is the unit's repair cost or number of repairs at that age. However, staircase functions are hard to view in such a plot. In contrast, the data from a sample unit is from just the observed early portion of its history function, subsequent history being censored, that is, yet unobserved.

MCF

At a particular age t , there is a population distribution of cumulative cost of repairs, which is depicted

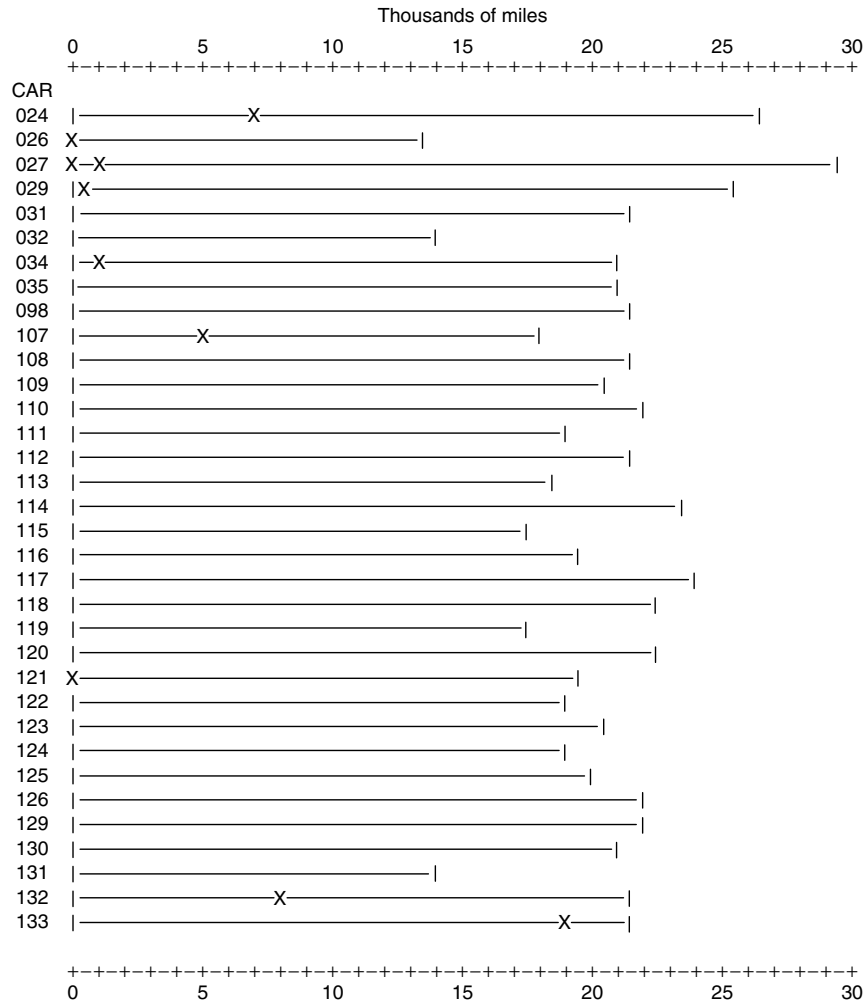


Figure 1 Display of automatic transmission data in production order [© SIAM, 2003.]

vertically in Figure 2 as a continuous density (which should be imagined perpendicular to the page). The distribution at age t has a mean value $M(t)$. In Figure 2, the heavy curve depicts $M(t)$ as a function of t . $M(t)$ is called the *mean cumulative function* (MCF) for the cost of repairs. It is simply the population mean curve for all the unit cumulative history functions. It provides most information sought from repair and other recurrence data. In practice, $M(t)$ is usually regarded as a smooth continuous function. Nelson [3, 6, 7] presents an unbiased estimate $M^*(t)$ for $M(t)$ from censored data; it

appears in the section titled “Estimate and Plot of the MCF”.

Number of Repairs

In Figure 2, a cost distribution is depicted as continuous. In contrast, for a count of the number of repairs, the corresponding vertical distribution for the cumulative numbers of repairs at age t is discrete and has the integers values 0, 1, 2, 3, and so on. Such discrete distributions are depicted in Figure 3 and should be imagined perpendicular to the page. In

4 Repair Data, Sets of: How to Graph, Analyze, and Compare

this case, $M(t)$ is the population mean curve for all unit staircase history functions.

Repair Rate

When $M(t)$ is the mean cumulative *number* of repairs, its derivative

$$m(t) \equiv \frac{dM(t)}{dt} \quad (1)$$

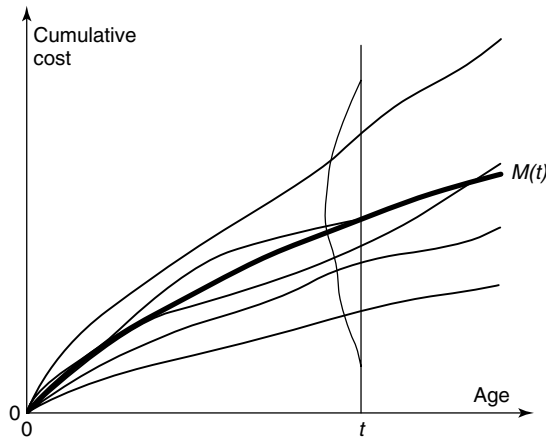


Figure 2 Population cumulative cost histories (uncensored) and cost distribution [Taken from TIS Report 89CRD239 © 1989.]

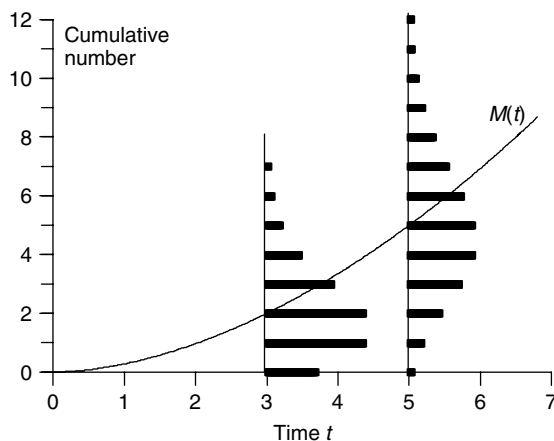


Figure 3 Discrete distributions for the cumulative number of repairs and MCF [Taken from TIS Report 89CRD239 © 1989.]

is assumed to exist and is called the population *instantaneous repair rate* at age t . It is also called the *recurrence rate* or *intensity function* for the number of other repeating occurrences. It is expressed in occurrences per unit time per population unit, e.g., transmission repairs per 1000 mi per car. Some mistakenly call $m(t)$ the *failure rate*, which causes confusion with the quite different failure rate (hazard function) of a life distribution for nonrepaired units (such as components). The failure rate for a life distribution has an entirely different definition, meaning, and use, as explained by Ascher and Feingold [8], Nelson [1], and Abernethy [2]. When $M(t)$ is the mean cumulative *cost* of repairs, $m(t)$ is the mean cost rate at age t , for example dollars per car per 1000 mi.

Estimate and Plot of the MCF

Purpose

This section shows how to estimate the population MCF.

Motivation

If all N sample units have been observed up to time t , the estimate of the population MCF is simple. Suppose that the cumulative number (or cost) of repairs on sample unit i by age t is $Y_i(t)$, $i = 1, 2, \dots, N$. Then the estimate of the MCF at age t is simply the sample average of the cumulative values at age t ; namely,

$$M^*(t) = [Y_1(t) + Y_2(t) + \dots + Y_N(t)]/N \quad (2)$$

When the time t is beyond any censoring ages, the MCF estimate is more complicated as follows.

Steps

The following steps yield a nonparametric estimate $M^*(t)$ of the population MCF $M(t)$ for the number of repairs from a simple random sample of N units; here $N = 34$ cars.

1. List all repair and censoring ages in order from smallest to largest, as in column 1 of Table 2. Denote each censoring age with a +. If a repair age of a unit equals its censoring age, put

Table 2 MCF calculations^a

1. Mileage	2. No. <i>I</i> in use	3. Mean No. 1/ <i>I</i>	4. Sample MCF
28	34	0.029	0.029
48	34	0.029	0.058
375	34	0.029	0.087
530	34	0.029	0.116
1388	34	0.029	0.145
1440	34	0.029	0.174
5094	34	0.029	0.203
7068	34	0.029	0.232
8250	34	0.029	0.261
13 809+	33		
14 235+	32		
14 281+	31		
17 844+	30		
17 955+	29		
18 228+	28		
18 676+	27		
19 175+	26		
19 250	26	0.038	0.299
19 321+	25		
19 403+	24		
19 507+	23		
19 607+	22		
20 425+	21		
20 890+	20		
20 997+	19		
21 133+	18		
21 144+	17		
21 237+	16		
21 401+	15		
21 585+	14		
21 691+	13		
21 762+	12		
21 876+	11		
21 888+	10		
21 974+	9		
22 149+	8		
22 486+	7		
22 637+	6		
22 854+	5		
23 520+	4		
24 177+	3		
25 660+	2		
26 744+	1		
29 834+	0		

^a© John Wiley & Sons Ltd, 1998

the repair age first. If two or more units have a common age, list them in a suitable order, possibly random.

- For each sample age, write the number I of units then in use ("at risk") in column 2 as follows.

If the earliest age is a censoring age, write $I = N - 1$; otherwise, write $I = N$. Proceed down column 2 writing the same I value for each successive repair age. At each censoring age, reduce the I value by one. The last age is a censoring age and for it, $I = 0$.

- For each repair, calculate the observed mean number of repairs per unit at that age as $1/I$. For example, for the repair at 28 mi, $1/34 = 0.029$, which appears in column 3. For a censoring age, the observed mean number of repairs per unit is zero, corresponding to a blank in column 3. The censoring ages are taken into account since they determine the I values of the repairs.
- In column 4, calculate the sample MCF $M^*(t)$ for each repair by cumulating as follows. For the earliest repair age, this is the corresponding mean number of repairs, namely, 0.029 in Table 2. For each successive repair age this is the corresponding mean number of repairs (column 3) plus the preceding mean cumulative number (column 4). For example, at 19 250 mi, this is $0.038 + 0.261 = 0.299$. Censoring ages have no mean cumulative number and are ignored in column 4.
- For each repair, plot on graph paper its mean cumulative number (column 4) against its age (column 1) as in Figure 4. This plot displays the nonparametric estimate $M^*(t)$ for the population $M(t)$. $M^*(t)$ is a staircase function and is called the *sample MCF*. Censoring times are not plotted but are taken into account in the MCF estimate. One interprets such a plot as explained in the section titled "How to Interpret and Use a Plot".

Plot

Figure 4 was plotted by Nelson and Doganaksoy's [9] program MCFLIM, which does the calculations above. Each data point "1" in the plot corresponds to a repair at the corresponding age in the sample. This plot shows a flat horizontal line segment to the right of each 1, which extends to the age of the next 1. This nonparametric sample MCF estimate is a staircase. In practice, one regards the population MCF as a smooth curve. Thus one usually imagines or fits a smooth curve through the plotted points as an MCF estimate. Often the horizontal segments are omitted, as in Figure 6.

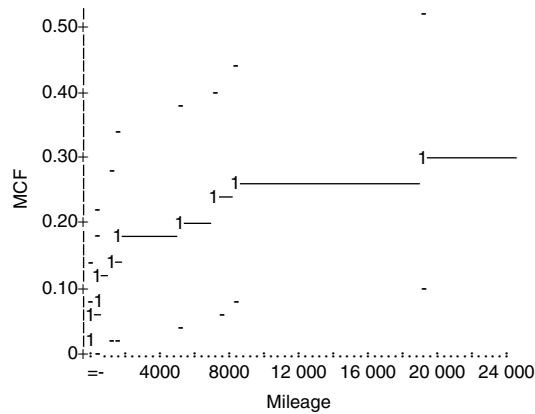


Figure 4 Automatic transmission sample MCF and 95% confidence limits “-” [© John Wiley & Sons Ltd, 1998.]

Confidence Limits

Figure 4 also displays vertical pairs of “-”, which are two-sided approximate 95% confidence limits for the corresponding population MCF at that age. Such a pair of pointwise limits encloses the population MCF with 95% confidence at any one age t . A pair of limits extends to the right to the next repair age. For example, the 95% confidence limits for the mean cumulative number of automatic transmission repairs at 24 000 track mi (read from Figure 4) are 0.1 and 0.5 repairs per car. Such pointwise limits do not simultaneously enclose the entire MCF curve over the range of the data with 95% confidence. Nelson’s [5, 6] complex calculations of these limits require a computer program like those described below. These limits employ a normal approximation to the sampling distribution of $M^*(t)$.

Assumptions

The nonparametric MCF estimate and approximate confidence limits above entail minimal assumptions described by Nelson [3, 5, 6]. In contrast, the parametric models and estimates presented by Englehardt [10] and Ascher and Feingold [8] apply to a single unit (not a sample of units) and require additional assumptions that may not be valid in applications. Rigdon and Basu [11] fit parametric non-homogeneous Poisson process models to repair data on a simple random sample of units; their models employ the often false assumption of independent increments.

Software

The following programs calculate and plot the MCF estimate and approximate confidence limits from data with exact ages and right censoring. They handle the number and cost (or value) of recurrences and allow for positive or negative values.

- MCFLIM of Nelson and Doganaksoy [9]. This software provided Figure 4.
- The RELIABILITY procedure in the SAS/QC software of the SAS Institute [12].
- The JMP software of the SAS Institute [13], pp. 565–572.
- SPLIDA features developed by Meeker and Escobar [14] for S-PLUS.
- A program developed for General Motors by Robinson [15].
- The ReliaSoft [16] Weibull++7 software has a recurrence data analysis (RDA) add-on.

How to Interpret and Use a Plot

MCF Estimate

The plot in Figure 4 displays a nonparametric estimate $M^*(t)$ of $M(t)$; that is, the estimate involves no assumptions about the form of $M(t)$ or the process generating the unit histories. This nonparametric estimate is a staircase function that is flat between repair ages, but the flat portions need not be plotted. The MCF of a large population is usually regarded as a smooth curve, and one usually imagines a smooth curve through the plotted points. Interpretations of such plots appear below. See Nelson [6] for more detail.

Mean Cumulative Number

An estimate of the population mean cumulative number of repairs by a specified age is read directly from such a staircase estimate or a curve through the plotted points. For example, from Figure 4, the estimate of this by 24 000 mi is 0.30 repairs per car, an answer to a basic question.

Repair Rate

The derivative of such a curve (imagined or fitted) estimates the repair rate $m(t)$. If the derivative increases with age, the population repair rate

increases as units age. If the derivative decreases, the population repair rate decreases with age. The behavior of the rate is used to determine burn-in, overhaul, and retirement policies. In Figure 4 the repair rate (derivative) decreases as the transmission population ages, providing the answer to a basic question.

Burn-In

Some products are subjected to a factory burn-in. Units typically are run and repaired until the instantaneous (population) repair rate decreases to a desired value m' . An estimate of the suitable length t' of burn-in is obtained from the sample MCF as shown in Figure 5. A straight line segment with slope m' is moved until it is tangent to the MCF. The corresponding burn-in age t' at the tangent point is suitable, as shown in Figure 5.

Other Information

Nelson [6] and Cook and Lawless [17] give other applications and show how to obtain other types of desired information.

Comparison of Sets of Repair Data

Purpose

This section presents a simple graphical comparison of the MCF estimates of two sets of repair data. This comparison also applies to recurrence

data on recurrent disease episodes, sociological and demographic applications such as childbirth, simulation of production processes, and other recurrent events. Nelson [6] and Cook and Lawless [17] give further examples of such comparisons.

Manual Transmission Data

Table 3 displays another set of repair data on a sample of 14 cars with a manual transmission. Information sought from these data include

1. the mean cumulative number of repairs per car by 24 000 mi (design life),
2. whether the population repair rate increases or decreases as the population ages, and
3. whether the automatic and manual sample MCFs differ statistically significantly.

Figure 6 is a plot of the sample MCF of the manual transmissions without the horizontal lines.

- A smooth curve through Figure 6 has a derivative $m^*(t)$ that decreases as the population ages. That is, the early repair rate decreases as the manual transmission population ages. This is typical of products with manufacturing defects.
- Figure 6 yields an estimate of 0.50 manual transmission repairs per car by 24 000 track mi. The corresponding estimate for automatic transmissions is 0.30.

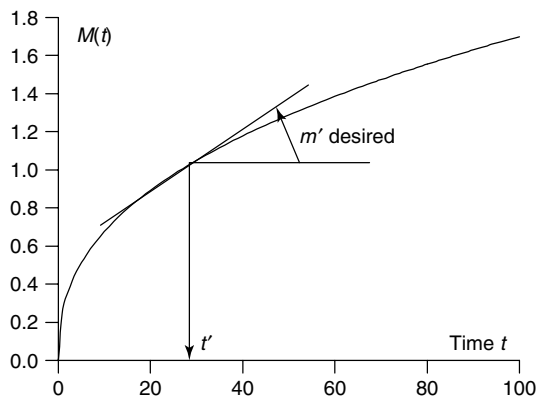


Figure 5 Burn-in age t' for repair rate m' [Taken from TIS Report 89CRD239, © 1989.]

Table 3 Manual transmission repair data^a

Car	Mileage (+ latest)		
025	27 099+		
028	21 999+		
030	11 891	27 583+	
097	19 966+		
099	26 146+		
100	3648	13 957	23 193+
101	19 823+		
102	2890	22 707+	
103	2714	19 275+	
104	19 803+		
105	19 630+		
106	22 056+		
127	22 940+		
128	3240	7690	18 965+

^a © John Wiley & Sons Ltd, 2000

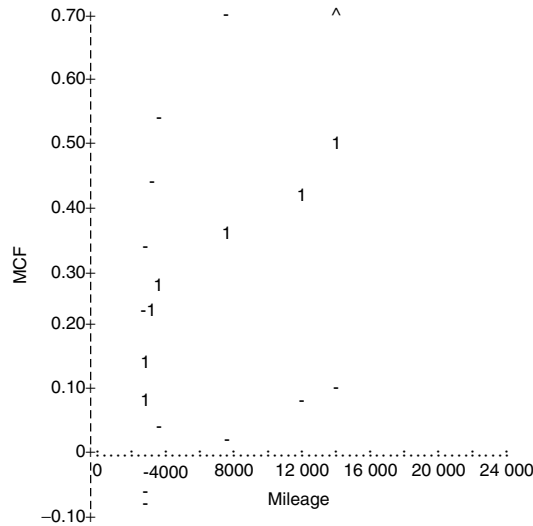


Figure 6 Manual transmission MCF and 95% confidence limits “-” [© John Wiley & Sons Ltd, 2000.]

Thus Figure 6 answers two of the questions listed above. Only the comparison of the automatic and manual transmission MCF estimates remains, which is described next.

Comparison

The following shows how to assess whether two sample MCFs differ statistically significantly, that is convincingly. A simple visual comparison of Figures 2 and 3 shows that the sample MCF of manual transmissions has a higher average cumulative number of repairs over time. The following method is like the one used to compare two sample means. However, here we are comparing two mean curves. To do this, we look at the difference between the two sample MCFs and the approximate two-sided 95% confidence limits for their difference. At an age t where those pointwise confidence limits do not enclose zero, that difference is statistically significant. Otherwise, the two sample MCFs at that age do not differ statistically significantly. Doganaksoy and Nelson [18] and Nelson [6] present the theory and calculations for these limits for the difference of two sample MCFs.

Transmission Comparison

Figure 7 displays the difference between the sample MCFs of the automatic and manual transmissions.

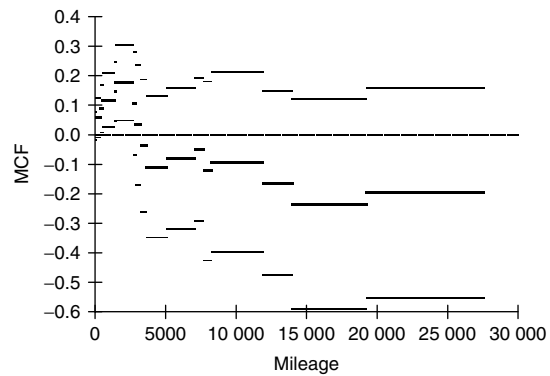


Figure 7 Difference of the MCFs (automatic minus manual) [© John Wiley & Sons Ltd, 2000.]

The plot includes pointwise 95% confidence limits for the difference. At essentially every age, the limits enclose zero. Thus the sample MCFs do not differ significantly over the range of the data. Of course, the two samples contain few cars and few repairs, and consequently the confidence limits are wide and the comparison is insensitive. In view of no significant difference, one might pool the two datasets to estimate a common MCF with greater accuracy, if that would be consistent with engineering knowledge.

Entire MCFs

The preceding comparison of two sample MCFs is pointwise. That is, the 95% confidence level is correct only for a single age. The limits do not enclose the *entire* difference of the true MCFs with 95% probability. Comparisons of the entire MCFs are given by Cook *et al.* [19], Cook and Lawless [17], and Nelson [20].

Software

All the previously referenced software calculate and plot the differences of two sample MCFs and corresponding pointwise confidence limits.

Further Information

The references, especially Nelson [6] and Cook and Lawless [17], and their references provide other applications and information on

1. predicting future numbers or costs of repairs for a fleet;
2. analyzing repair *cost* data or other numerical values associated with repairs;
3. analyzing such data on cost, downtime, and other measures of repair “value”;
4. analyzing history functions that increase and decrease (negative values) or are continuous;
5. analyzing availability data, which includes downtime for repairs;
6. analyzing data with more complex censoring, such as gaps in unit histories with missing repair data, Nelson [4];
7. analyzing data with a mix of types of repairs;
8. the minimal assumptions for the nonparametric estimate $M^*(t)$ and confidence limits;
9. data analyses and plots for samples other than simple random ones, such as cluster and stratified samples; and
10. estimating relationships, such as the **Cox model**, for the dependence of $M(t)$ on process, environmental, materials, and other engineering variables; see Cook and Lawless [17].

Literature

Many models and data analysis methods for repair data are parametric and involve more assumptions, which often are unrealistic. For example, Englehardt [10] and Ascher and Feingold [8] present models, analyses, and assumptions for a single unit, not for a sample of units. The simplest such parametric models are the **Poisson process** and **nonhomogeneous Poisson process**, presented in detail by Nelson [6] and Rigdon and Basu [11]. In contrast, the methods above extend to costs, downtimes, and other values associated with repairs, whereas previous methods apply only to simple counts of repairs.

Concluding Remarks

The simple plot here of the sample MCF is informative and widely useful. It requires minimal assumptions and is simple to make, interpret, and present to others.

Acknowledgment

This updated revision of Nelson [7, 21] appears here with the kind permission of Wiley, publisher of *Quality and*

Reliability Engineering International. The author gratefully thanks Mr Richard J. Rudy of Daimler–Chrysler, who generously granted permission to use the transmission data. The author also thanks Mr. Fred Faltin, the editor-in-chief of this encyclopedia, for his encouragement to write this article.

References

- [1] Nelson, W. B. (2004). *Applied Life Data Analysis*, John Wiley & Sons, New York, pp. 634. www.wiley.com.
- [2] Abernethy, R.B. (2006). *The New Weibull Handbook*, 5th Edition, available from R.B. Abernethy, weibull@worldnet.att.net.
- [3] Nelson, W. (1988). Graphical analysis of system repair data, *Journal of Quality Technology* **20**, 2–35.
- [4] Nelson, W. (1990). Hazard plotting of left truncated life data, *Journal of Quality Technology* **22**, 230–238.
- [5] Nelson, W. (1995). Confidence limits for recurrence data—applied to cost or number of repairs, *Technometrics* **37**, 147–157.
- [6] Nelson, W.B. (2003). *Recurrent Events Data Analysis for Product Repairs, Disease Recurrences, and Other Applications*, ASA-SIAM Series on Statistics and Applied Probability, SIAM, www.siam.org/books/sa10.
- [7] Nelson, W. (1998). An application of graphical analysis of repair data, *Quality and Reliability Engineering International* **14**, 49–52.
- [8] Ascher, H. & Feingold, H. (1984). *Repairable Systems Reliability*, Marcel Dekker, New York.
- [9] Nelson, W. & Doganaksoy, N. (1994). *Documentation for MCFLIM and MCFDIFF – Programs for Recurrence Data Analysis*, available from WNconsult@aol.com.
- [10] Englehardt, M. (1995). Models and analyses for the reliability of a single repairable system, in *Recent Advances in Life-Testing and Reliability*, N. Balakrishnan, ed., CRC Press, Boca Raton, pp. 79–106.
- [11] Rigdon, S.E. & Basu, A.P. (2000). *Statistical Methods for the Reliability of Repairable Systems*, John Wiley & Sons, New York.
- [12] SAS Institute. (2004). *SAS/QC 9.1 User's Guide*, SAS Publishing, Cary, NC, Vols. 1, 2, and 3, <http://support.sas.com/91doc/docMainpage.jsp>.
- [13] SAS Institute. (2005). *JMP Statistical Discovery Software: Statistics and Graphics Guide*, Version 6, SAS Publishing, Cary, NC, pp. 565–572.
- [14] Meeker, W.Q. & Escobar, L.A. (2004). *SPLIDA (S-PLUS Life Data Analysis) Software – Graphical User Interface*, available from Prof William Q. Meeker, Statistics Department, Iowa State University, Ames, IA 50010 or www.public.iastate.edu/~splida.
- [15] Robinson, J.A. (1995). Standard errors for the mean cumulative number of repairs on systems from a finite population, in *Recent Advances in Life-Testing and Reliability*, N. Balakrishnan, ed., CRC Press, Boca Raton, pp. 195–217.

- [16] ReliaSoft Corporation. (2005). *Weibull++7 Life Data Analysis Reference*, ReliaSoft Publishing, Tucson, AZ www.reliasoft.com. RDA documentation at http://www.weibull.com/LifeDataWeb/recurrent_events_data_analysis.htm.
- [17] Cook, R.J. & Lawless, J.F. (2007). *The Statistical Analysis of Recurrent Events*, Springer, New York, pp. 420.
- [18] Doganaksoy, Necip & Nelson, Wayne. (1998). A method to compare two samples of recurrence data, *Lifetime Data Analysis* **4**, 51–63.
- [19] Cook, R.J., Lawless, J.F. & Nadeau, C. (1996). Robust tests for treatment comparisons based on recurrent event responses, *Biometrics* **52**, 557–571.
- [20] Nelson, W.B. (2008). Comparisons of sets of recurrence data, 2008 *IEEE Transactions on Reliability*.
- [21] Nelson, W. (2000). Graphical comparison of sets of repair data, *Quality and Reliability Engineering International* **16**, 235–241.

Related Articles

Age-Dependent Minimal Repair and Maintenance; Analysis of Recurrent Events from Repairable Systems; General Minimal Repair Models; Imperfect Repair; Nonparametric Methods for Analysis of Repair Data; Repairable Systems Reliability; Repairable Systems: Statistical Inference; Repairable Systems: Bayesian Analysis; Warranty Claims and Costs: Statistical Analysis of.

WAYNE B. NELSON

Cumulative Hazard Function

Consider a random variable T denoting a failure time and the associated hazard function $h(t)$ (*see **Hazard Function***). The cumulative hazard function is defined as $H(t) = \int_0^t h(u) \, du$.

Intensity Function

Let $N(t)$, $t \geq 0$, be a point process and let $N(s, t]$, $s < t$, denote the number of events (failures) that occur in the interval $(s, t]$. Then the intensity function of the point process $N(t)$ is defined as

$$\lambda(t) = \lim_{\Delta t \rightarrow 0} \frac{P(N(t, t + \Delta t] \geq 1)}{\Delta t} \quad (1)$$

If the probability of simultaneous failures is zero, then the intensity function $\lambda(t)$ coincides with the ROCOF $\mu(t)$ (see **Rate of Occurrence of Failures**).

Hazard Function

Consider a random variable T denoting a failure time. The hazard function is defined as

$$h(t) = \lim_{\Delta t \rightarrow 0} \frac{P(t < T \leq t + \Delta t)}{P(T > t)} \quad (1)$$

The hazard function denotes the (conditional) probability that an item functioning at time t fails right after that time.

If T has density function $f(t)$ and reliability function (see **Reliability Function**) $R(t)$, then it follows that

$$h(t) = \frac{f(t)}{R(t)} \quad (2)$$

Since $f(t) = -R'(t)$ in the absolutely continuous case, it follows that, knowing either the density function or the reliability function, it is possible to determine the hazard function univocally. A similar result holds for discrete distributions.

Given the hazard function, it is possible to obtain the density function $f(t)$. In the absolutely continuous case, it follows that $f(t) = h(t)e^{-\int_0^t h(u) du}$.

Suppose that T has an exponential distribution, with density $\lambda e^{-\lambda t}$ and reliability function $e^{-\lambda t}$, $\lambda > 0$, $t \geq 0$. Therefore the hazard function is given by $h(t) = \lambda$. Since the hazard function univocally determines the distribution function, the constant failure rate is another typical property of the exponential distribution, which justifies the “memoryless” property commonly associated with it.

Mean Residual Life

Consider a random variable T denoting the failure time of an item. Suppose that the item has been functioning up to time t , then the mean residual life at time t is defined as

$$m(t) = E(T - t | T \geq t) \quad (1)$$

Given the reliability function (see **Reliability Function**) $R(t)$, it follows that $m(t) = \int_t^\infty R(u) du / R(t)$. There is a one-to-one correspondence between

mean residual life and reliability function since it holds that

$$R(t) = \frac{m(0)}{m(t)} e^{-\int_0^t (1/m(u)) du} \quad (2)$$

Suppose that T has an exponential distribution, with density $\lambda e^{-\lambda t}$ and reliability function $e^{-\lambda t}$, $\lambda > 0$, $t \geq 0$. Using the previous result about the link with the reliability function, it follows that $m(t) = 1/\lambda$, as another consequence of the “memoryless” property.

Mean Time Between Failures

Consider a random variable T denoting a time between two failures. The mean time between failures (MTBF) is defined as the expected value of T , provided it exists.

If T has a density function $f(t)$, then the MTBF is given by $\int_0^\infty t f(t) dt$ or $\sum_i t_i f(t_i)$, depending on whether T is absolutely continuous or discrete, respectively.

Suppose that T has an exponential distribution, with density $\lambda e^{-\lambda t}$, $\lambda > 0$, $t \geq 0$. Then the MTBF is given by

$$\int_0^\infty t \lambda e^{-\lambda t} dt = \frac{1}{\lambda} \quad (1)$$

Mean Time to Failure

Consider a random variable T denoting a time to failure. The mean time to failure (MTTF) is defined as the expected value of T , provided it exists.

If T has a density function $f(t)$, then the MTTF is given by $\int_0^\infty t f(t) dt$ or $\sum_i t_i f(t_i)$, depending on whether T is absolutely continuous or discrete, respectively.

Given the reliability function (*see Reliability Function*) $R(t)$, it follows that the MTTF equals $\int_0^\infty R(t) dt$ in the absolutely continuous case.

Suppose that T has an exponential distribution, with density $\lambda e^{-\lambda t}$ and reliability function $e^{-\lambda t}$, $\lambda > 0$, $t \geq 0$. Then the MTTF is given by

$$\int_0^\infty t \lambda e^{-\lambda t} dt = \int_0^\infty e^{-\lambda t} dt = \frac{1}{\lambda} \quad (1)$$

Mean Time to Repair

Consider a random variable R denoting the repair (or down) time after the failure of an item. The mean time to repair (MTTR) is defined as the expected value of R , provided it exists.

If R has a density function $f(r)$, then the MTTR is given by $\int_0^\infty r f(r) \, dr$ or $\sum_i r_i f(r_i)$, depending on whether R is absolutely continuous or discrete, respectively.

Given the distribution function $F(r)$, it follows

that the MTTR equals $\int_0^\infty [1 - F(r)] \, dr$ in the absolutely continuous case.

Suppose that R has an exponential distribution, with density $\lambda e^{-\lambda r}$ and distribution function $1 - e^{-\lambda r}$, $\lambda > 0$, $r \geq 0$. Then the MTTR is given by

$$\int_0^\infty r \lambda e^{-\lambda r} \, dr = \int_0^\infty e^{-\lambda r} \, dr = \frac{1}{\lambda} \quad (1)$$

Proportional Hazard Model

The proportional hazard model relates the hazard function (*see* **Hazard Function**) at two different conditions, x and the baseline x_0 , by

$$h(t; x) = g(x)h(t; x_0) \quad (1)$$

with $t > 0$ and g a positive function, e.g. $g(x; \beta) = e^{x\beta}$, where β is a parameter.

The links between the reliability function (*see* **Reliability Function**) and the distribution function and the baselines ones are given, respectively, by $R(t; x) = [R(t; x_0)]^{g(x)}$ and $F(t; x) = 1 - [1 - F(t; x_0)]^{g(x)}$.

Reliability Function

Consider a random variable T denoting a failure time. The reliability function is defined as $R(t) = P(T > t)$, for $t \geq 0$.

If T has a density function $f(t)$, then the reliability function is given by $\int_t^\infty f(u) du$ or $\sum_{t_i > t} f(t_i)$, depending on whether T is absolutely continuous or discrete, respectively.

Given the distribution function $F(t)$, it follows that $R(t) = 1 - F(t)$.

It can be shown that the reliability function $R(t)$ is linked to the hazard function $h(t)$ (see **Haz-**

ard Function) and the cumulative hazard function $H(t)$ (see **Cumulative Hazard Function**) by $R(t) = e^{-\int_0^t h(u) du}$ and $R(t) = e^{-H(t)}$, respectively. For the converse, see **Hazard Function**. Therefore, knowing either the hazard function or the cumulative hazard function, it is possible to determine the reliability function univocally.

Suppose that T has an exponential distribution, with density $\lambda e^{-\lambda t}$ and distribution function $1 - e^{-\lambda t}$, $\lambda > 0$, $t \geq 0$. Then the reliability function is given by

$$\int_t^\infty \lambda e^{-\lambda u} du = e^{-\lambda t} \quad (1)$$

Rate of Occurrence of Failures

Let $N(t)$, $t \geq 0$, be a point process and let $N(s, t]$, $s < t$, denote the number of events (failures) that occur in the interval $(s, t]$. Consider the mean function of the process $N(t)$, defined as $\Lambda(t) = EN(t)$, which denotes the expected number of events in

the interval $(0, t]$. If $\Lambda(t)$ is differentiable, then the ROCOF of the point process $N(t)$ is defined as

$$\mu(t) = \frac{d}{dt} \Lambda(t) \quad (1)$$

If the probability of simultaneous failures is zero, then the ROCOF $\mu(t)$ coincides with the intensity function $\lambda(t)$ (*see **Intensity Function***).

Contingency Table, Square

An $R \times C$ contingency table – that is, a contingency table with R rows and C columns – is frequently seen in biomedical studies. When $R = C$, the contingency table is often referred to as a *square table*. Square tables usually arise in repeated measures experiments, in which a subject has his or her outcome variable observed repeatedly. Examples of repeated measures studies include longitudinal studies (observing the same subject over time); rater agreement studies (two investigators rate each subject in the study); and paired data studies (data obtained from a husband and wife, father and son, or case and control).

First, we introduce some notation. Suppose that two discrete random variables, Y_{i1} and Y_{i2} , are observed on each of n independent subjects, where Y_{i1} can take on values $1, \dots, R$, and Y_{i2} can take on values $1, \dots, C$. Let the probabilities of the multinomial joint distribution of Y_{i1} and Y_{i2} be denoted by

$$p_{jk} = \Pr(Y_{i1} = j, Y_{i2} = k) \quad (1)$$

for $j = 1, \dots, R$ and $k = 1, \dots, C$. Since all the probabilities must sum to 1, there are $(RC - 1)$ nonredundant multinomial cell probabilities. We let $\mathbf{p}$ denote the $(RC - 1) \times 1$ probability vector of p_{jk} 's; for simplicity, we delete p_{RC} , and we take $R = C$ below. The marginal probabilities of the contingency table are $p_{j+} = \Pr(Y_{i1} = j)$ and $p_{+k} = \Pr(Y_{i2} = k)$, where the plus signs denote summing over the subscript they replace. We let n_{jk} denote the number of subjects with response level j on Y_{i1} and level k on Y_{i2} , and let $n = n_{++}$ be the total number of subjects in the study.

Models for square tables fall into three general classes – marginal models, loglinear models, and conditional models. Marginal models describe the similarity between the row and column marginal distributions. With loglinear models, we model the cell probabilities of the table. For example, the probabilities may be symmetric about the main diagonal, and a loglinear model can be used to describe this relationship among the cell probabilities. Conditional models model the conditional probabilities of the column variable given the row variable and are popularly

used in longitudinal studies, letting the row variable represent a response at time 1 and the column variable represent the response at time 2. Often, in square tables, we are not always interested in modeling the data, but are interested in determining the association or agreement between the row and column variables.

Marginal Models

In a crossover study, a subject is given one treatment, has a washout period, and is then given a new treatment. Suppose that the response to each treatment is success, partial success, or failure. We are often interested in whether the marginal distribution of the outcome is the same for both treatments, often called *marginal homogeneity*. Under marginal homogeneity, the sum of the cell probabilities (i.e., the marginal probability) in the j th row of the contingency table equals the sum of the cell probabilities in the j th column; that is, $p_{j+} = p_{+j}$ for $j = 1, \dots, R$. Lipsitz *et al.* [1] show that the maximum-likelihood estimates (MLEs) for the p_{jk} 's under marginal homogeneity can be obtained *via* a linear model of the form

$$\mathbf{p} = \mathbf{X}\boldsymbol{\beta} \quad (2)$$

for the appropriate “design” matrix $\mathbf{X}$, which we now describe for $R = 3$, but which generalizes to any R . For $R = 3$, we can write the vector $\mathbf{p}$ in terms of the four cell probabilities, $\{p_{11}, p_{12}, p_{21}, p_{22}\}$ and the four marginal probabilities, $\{p_{1+}, p_{2+}, p_{+1}, p_{+2}\}$, as follows:

$$\begin{bmatrix} p_{11} \\ p_{12} \\ p_{13} \\ p_{21} \\ p_{22} \\ p_{23} \\ p_{31} \\ p_{32} \end{bmatrix} = \begin{bmatrix} p_{11} \\ p_{12} \\ p_{1+} - p_{11} - p_{12} \\ p_{21} \\ p_{22} \\ p_{2+} - p_{21} - p_{22} \\ p_{+1} - p_{11} - p_{21} \\ p_{+2} - p_{12} - p_{22} \end{bmatrix} \quad (3)$$

Under marginal homogeneity, the row and column marginal probabilities are equal – that is, $p_{j+} = p_{+j} = p_j$ – which leads to the following linear model

2 Contingency Table, Square

for the cell probabilities:

$$\begin{aligned}
 \begin{bmatrix} p_{11} \\ p_{12} \\ p_{13} \\ p_{21} \\ p_{22} \\ p_{23} \\ p_{31} \\ p_{32} \end{bmatrix} &= \begin{bmatrix} p_{11} \\ p_{12} \\ p_1 - p_{11} - p_{12} \\ p_{21} \\ p_{22} \\ p_2 - p_{21} - p_{22} \\ p_1 - p_{11} - p_{21} \\ p_2 - p_{12} - p_{22} \end{bmatrix} \\
 &= \begin{bmatrix} 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 \\ -1 & -1 & 0 & 0 & 1 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & -1 & -1 & 0 & 1 \\ -1 & 0 & -1 & 0 & 1 & 0 \\ 0 & -1 & 0 & -1 & 0 & 1 \end{bmatrix} \begin{bmatrix} p_{11} \\ p_{12} \\ p_{21} \\ p_{22} \\ p_1 \\ p_2 \end{bmatrix} \\
 &= \mathbf{X}\boldsymbol{\beta}. \tag{4}
 \end{aligned}$$

The MLEs can be obtained in any generalized linear modeling program, such as SAS *Proc GENMOD* [2], without any additional programming or iteration loops. In these programs, the outcome is the count n_{jk} , the covariates are the appropriate row of $\mathbf{X}$ in equation (4) without an intercept, and we specify a linear link with Poisson errors (we actually must include n_{33} as a value of the outcome, as described in [1]). Firth [3] also describes a method for obtaining the linear model for the p_{jk} 's using Latin squares. Another method found in the literature for obtaining the MLE under homogeneity uses Lagrange multipliers [4].

After the estimates of the cell probabilities under marginal homogeneity have been obtained, a likelihood-ratio statistic can be used to test whether marginal homogeneity holds. This statistic is approximately χ^2 with $R - 1$ **degrees of freedom** under the null. Alternately, a Wald statistic [5] can be used. It does not require the use of iterative techniques, since only the estimates of the cell probabilities for the saturated model are needed. The MLE of p_{jk} for the saturated model (with $R^2 - 1$ nonredundant probabilities) is

$$\hat{p}_{jk} = \frac{n_{jk}}{n} \tag{5}$$

and the MLEs of the marginal probabilities are $\hat{p}_{j+} = n_{j+}/n$ and $\hat{p}_{+j} = n_{+j}/n$. Suppose that we let the vector

$$\mathbf{U} = [\hat{p}_{1+} - \hat{p}_{+1}, \dots, \hat{p}_{R-1,+} - \hat{p}_{+,R-1}]' \tag{6}$$

contain the first $R - 1$ differences in the marginal probabilities (the R th difference is redundant). Under the null of marginal homogeneity, $E(\mathbf{U}) = \mathbf{0}$. If we let $\hat{\mathbf{V}}$ be the estimated **covariance** matrix of the vector $\mathbf{U}$ under the alternative (see [5]), then the Wald test statistic for homogeneity is the quadratic form

$$X^2 = \mathbf{U}'\hat{\mathbf{V}}^{-1}\mathbf{U} \tag{7}$$

which will be approximately χ^2 with $R - 1$ degrees of freedom under the null. If, instead of estimating the variance under the alternative in equation (7), we let $\hat{\mathbf{V}}$ be the estimated covariance of $\mathbf{U}$ under the null of homogeneity, and if there are only two rows and columns in the table ($R = C = 2$), then (7) reduces to

$$X^2 = \frac{(n_{12} - n_{21})^2}{n_{12} + n_{21}} \tag{8}$$

The test statistic in equation (8) is popularly known as *McNemar's test* for equality of correlated proportions [6].

In the crossover study discussed earlier, the outcomes "success", "partial success", and "failure" are examples of ordered categorical data. Suppose that we want to test for marginal homogeneity, taking this ordering into account. We can put an ordinal model on the margins, such as the proportional-odds model [7], and test whether the parameters of the proportional-odds models in the two margins are the same. The test statistics are similar to those given for the nominal case above. The Wald test for marginal homogeneity for proportional-odds marginal models can be obtained in SAS *Proc CATMOD* [8] using methods derived in [9]. The likelihood-ratio statistic can be obtained in a **generalized linear models** program, but will require some additional programming. More details on marginal modeling can be found in [10], Chapter 9 and in the article on marginal models.

Loglinear Models

The models just discussed involved testing whether the marginal probabilities of the table are equal. Now, we discuss models which describe the relationship among the cell probabilities of the $R \times R$ table. One such model for the cell probabilities is a symmetry model, in which

$$p_{jk} = p_{kj}, \quad \text{for } j \neq k \tag{9}$$

Note that symmetry implies marginal homogeneity since, under equation (9),

$$p_{j+} = \sum_{k=1}^R p_{jk} = \sum_{k=1}^R p_{kj} = p_{+j} \quad (10)$$

In fact, when $R = C = 2$, marginal homogeneity and symmetry are identical, and both imply that $p_{12} = p_{21}$.

For modeling purposes, symmetry is most easily written as a loglinear model. Loglinear models for contingency tables are usually written in terms of the expected cell counts

$$m_{jk} = E(n_{jk}) = np_{jk} \quad (11)$$

In terms of the expected cell counts, symmetry is written as

$$m_{jk} = m_{kj}, \quad \text{for } j \neq k \quad (12)$$

Then, the symmetry loglinear model is

$$\log m_{jk} = \mu + \beta_j + \beta_k + \beta_{jk} \quad (13)$$

where $\beta_{jk} = \beta_{kj}$, with the appropriate identifiability constraints on the parameters, such as $\beta_R = 0$ and $\beta_{Rj} = \beta_{jR} = 0$ for $j = 1, \dots, R$. Since $\beta_{jk} = \beta_{kj}$, it is easy to show that $m_{jk} = m_{kj}$, and symmetry holds for equation (13).

The MLEs for the expected cell counts under symmetry are

$$\hat{m}_{jj} = n_{jj} \quad \text{and} \quad \hat{m}_{jk} = \frac{n_{jk} + n_{kj}}{2} \quad (14)$$

The estimated cell probabilities are just $\hat{p}_{jk} = \hat{m}_{jk}/n$. Intuitively, since symmetry does not concern the diagonal cells, the estimated expected diagonal counts are just the observed diagonal counts. Also, an estimated off-diagonal count is just the average of the observed counts in cells jk and kj . Under the symmetry model, we only need to estimate $R(R-1)/2$ off-diagonal probabilities on one side of the diagonal, so we have placed $R(R-1)/2$ constraints on the cell probabilities under the null, and the likelihood-ratio test statistic, score test statistic, or Wald test statistic for symmetry are approximately χ^2 with $R(R-1)/2$ degrees of freedom. The score test

statistic (equivalent to Pearson's χ^2 for this problem) is the simplest [11]:

$$X^2 = \sum_{j < k} \frac{(n_{jk} - n_{kj})^2}{n_{jk} + n_{kj}} \quad (15)$$

Recall that, when $R = C = 2$, marginal homogeneity and symmetry are identical, so that equation (15) is also a test statistic for marginal homogeneity when $R = 2$, and is identical to the McNemar test statistic discussed earlier.

The symmetry loglinear model puts many constraints on the probabilities. A loglinear model that has fewer constraints (and more parameters) does not force the row and column main effects in equation (13) to be equal,

$$\log m_{jk} = \mu + \alpha_j + \beta_k + (\alpha\beta)_{jk} \quad (16)$$

where $(\alpha\beta)_{jk} = (\alpha\beta)_{kj}$. The loglinear model in equation (16) is called *quasisymmetry* [12]. If symmetry holds, then $\alpha_j = \beta_j$. The Bradley-Terry model [13] for *paired comparisons* is a special case of this quasisymmetry model.

Using the identifiability constraints $\alpha_R = 0$, $\beta_R = 0$, and $(\alpha\beta)_{Rj} = (\alpha\beta)_{jR} = 0$ for $j = 1, \dots, R$, the interaction parameter $(\alpha\beta)_{jk}$ is the log-odds ratio for the j th and R th rows and k th and R th columns; that is,

$$(\alpha\beta)_{jk} = \log \left(\frac{p_{jk} p_{RR}}{p_{jR} p_{Rk}} \right) \quad (17)$$

Then, in the quasisymmetry model, since $(\alpha\beta)_{jk} = (\alpha\beta)_{kj}$, the log-odds ratio for the j th and R th rows and k th and R th columns equals the log-odds ratio for the k th and R th rows and j th and R th columns. In particular, these log-odds ratios are symmetric.

One last loglinear model that we discuss is a *quasi-independence* model. If the row and column variables are independent, then

$$p_{jk} = p_{j+} p_{+k} \quad (18)$$

or, equivalently, in terms of the loglinear model,

$$\log m_{jk} = \mu + \alpha_j + \beta_k \quad (19)$$

In repeated measures studies, most of the agreement is often on the diagonal, so that the diagonal elements are longer than expected under independence, with independence holding in the off-diagonal

4 Contingency Table, Square

elements. A simple loglinear model that expresses such a relationship is the quasi-independence loglinear model,

$$\log m_{jk} = \mu + \alpha_j + \beta_k + \gamma_j I(j = k) \quad (20)$$

where $I(j = k)$ equals 1 for diagonal elements $j = k$ and 0 for off-diagonal elements. For off-diagonal elements, equation (20) is identical to equation (19), and, if $\gamma_j > 0$, then the diagonal elements are longer than expected under independence.

When the rows and columns are ordered, these loglinear models can be extended by assigning scores to the rows and columns (see, for example [14]). For more details on loglinear models, see [11].

Conditional Models

Often in longitudinal studies, we are interested in transitions or change in states, such as the response at time 2 given the response at time 1. In general, in crossover designs, polling studies, or employment studies (mover–stayer studies), investigators are interested in these conditional probabilities for the column variable given the row variable. When the row variable represents a response at time 1, and column variable represents the response at time 2, we model the probability of response k at time 2 given response j at time 1:

$$p_{k|j} = \Pr[Y_{i2} = k | Y_{i1} = j], \quad j = 1, \dots, R \quad (21)$$

Suppose that, in an clinical trial on subjects with arthritis for the effectiveness of a single treatment, the subjects are observed once before treatment, then are given the new treatment, and then observed again. Suppose that the possible outcomes at the two times are “no pain”, “mild pain”, and “severe pain”. Then, we are interested in the changes in pain, as modeled by the conditional probabilities of pain status at time 2 given the pain status at time 1. If there are two levels ($R = 2$), then we can apply logistic regression to the model given in equation (21), treating Y_{i1} as a **covariate**. If ($R > 2$) and the levels are not ordered, multinomial logistic regression can be used; if the levels are ordered (as in this example), an ordinal logistic model such as the proportional-odds model [7] can be used.

Measures of Association and Agreement

For general ($R \times C$) tables, we are often interested in the association between the row and column variables. For a (2×2) table, the odds ratio is a popular measure of association. In an ($R \times R$) square table, the log-odds ratios, as described above, can be used to measure association. For the row and column variables to be associated, you should be able to predict one from another. For example, if the odds ratio is zero, the row and column variables in a (2×2) are perfectly negatively associated. In repeated measures studies, and, in particular, interrater reliability studies, in which two investigators rate each subject in the study, we are interested in how well the row and column variables (the two ratings on each subject) agree. When two ratings agree, most of the observations in the contingency table will be on the diagonal (most ratings are similar). Two variables can be highly associated, but agreement between the two could be very low. Suppose that both diagonal elements in a (2×2) table are zero: then the raters completely disagree, and agreement (by any measure) is very low. However, the odds ratio is zero, and, as discussed above, the two variables are perfectly negatively associated. The κ coefficient [15] is a popular measure of agreement corrected for chance agreement.

References

- [1] Lipsitz, S.R., Laird, N.M. & Harrington, D.P. (1990). Finding the design matrix for the marginal homogeneity model, *Biometrika* **77**, 353–358.
- [2] SAS Institute Incorporated. (1993). SAS/STAT software: the GENMOD procedure, release 6.09, SAS Technical Report P-243, SAS Institute Incorporated, Cary.
- [3] Firth, D. (1989). Marginal homogeneity and the superposition of Latin squares, *Biometrika* **76**, 179–182.
- [4] Madansky, A. (1963). Tests of homogeneity for correlated samples, *Journal of the American Statistical Association* **58**, 97–119.
- [5] Bhapkar, V.P. (1966). A note on the equivalence of two test criteria for hypotheses in categorical data, *Journal of the American Statistical Association* **61**, 228–235.
- [6] McNemar, Q. (1947). Note on the sampling error of the difference between correlated proportions or percentages, *Psychometrika* **12**, 153–157.
- [7] McCullagh, P. (1980). Regression models for ordinal data (with discussion), *Journal of the Royal Statistical Society, Series B* **42**, 109–142.

-
- [8] SAS Institute Incorporated. (1993). *SAS/STAT User's Guide Volume 1: The CATMOD Procedure Version 6*, 4th Edition, SAS Institute Incorporated, Cary, pp. 405–517.
- [9] Koch, G.G., Landis, J.R., Freeman, J.L., Freeman, D.H. & Lehnen, R.G. (1977). A general methodology for the analysis of experiments with repeated measurement of categorical data, *Biometrics* **33**, 133–158.
- [10] Agresti, A. (1996). *An Introduction to Categorical Data Analysis*, John Wiley & Sons, New York.
- [11] Agresti, A. (1990). *Categorical Data Analysis*, John Wiley & Sons, New York.
- [12] Caussinus, H. (1965). Contribution a l'analyse statistique des tableaux de correlation, *Annales de la Faculté des Sciences Université de Toulouse* **29**, 77–182.
- [13] Bradley, R.A. & Terry, M.E. (1952). Rank analysis of incomplete block designs I. The method of paired comparisons, *Biometrika* **39**, 324–345.
- [14] Goodman, L.A. (1979). Simple models for the analysis of association in cross-classifications having ordered categories, *Journal of the American Statistical Association* **74**, 537–552.
- [15] Cohen, J. (1960). A coefficient of agreement for nominal scales, *Educational and Psychological Measurement* **20**, 37–46.

STUART R. LIPSITZ

Article originally published in Encyclopedia of Biostatistics, 2nd Edition (2005, John Wiley & Sons, Ltd). Minor revisions for this publication by Jeroen de Mast.

Contingency Table

We start by defining a two-way contingency table. Suppose that each of a sample of n individuals is classified according to each of two separate criteria, and each of these criteria can take only a finite number of distinct values. Let us take as an example a study to assess factors associated with women's attitudes toward mammography [1, p. 220]. Each of 309 women has been classified according to:

1. Her "Mammography Experience", which has as its possible values "Never", "Over one year ago", and "Within the past year".
2. Her response to the question "How likely is it that a mammogram could find a new case of breast cancer?", which has as its possible values "Not likely", "Somewhat likely", and "Very likely".

Let the mammography experience correspond to the *rows* of the table and the response to the question about the detection of breast cancer to the *columns*. Although it happens in this particular example that each of the rows and columns has a natural ordering, this is not necessary to the definition of a contingency table. Then the resulting 3×3 *table* of frequencies or counts is, say, (n_{ij}) , where

1. n_{ij} is the number of individuals in row i and column j for $i = 1, \dots, I$ and $j = 1, \dots, J$.
2. I and J are, respectively, the total numbers of rows and columns of the table, here 3, 3. Let n be the total sample size, so that n is the sum of entries of the table; we can then write $n = \sum \sum n_{ij}$.

Note that in our construction of this *cross tabulation* or contingency table, we assume that each individual of the sample of size n is classified into exactly one of the IJ categories of the table, these categories being mutually exclusive and exhaustive. Written less formally, this means that the IJ categories do not overlap in any way, and together they cover all possibilities.

Of course, exactly how the sample is collected for a particular study is of great importance in the subsequent statistical analysis. In classical contingency table analysis, it is assumed that the n individuals form a *random sample* from a population:

this means that we assume that each individual is classified independently of the others, and each has the same probability, say p_{ij} , of being classified into row i and column j . Thus $\sum \sum p_{ij} = 1$. Our assumption that we have a random sample has the consequence that the probability distribution of the variable (n_{ij}) is multinomial; that is, the frequency function $p[(n_{ij})|p]$ is

$$n! \prod_i \prod_j \left(\frac{p_{ij}^{n_{ij}}}{n_{ij}!} \right)$$

provided that $n_{ij} \geq 0$ and $\sum \sum n_{ij} = n$.

In this case we say that (n_{ij}) is multinomial, with parameters n and (p_{ij}) .

To continue with the above example of the attitudes toward mammography data, the results of our survey are shown in Table 1.

These data show a dramatic and self-evident association between rows and columns: in other words, the women's responses to the question about the detection of breast cancer depend strongly on their mammography experience. This can be seen even more clearly by presenting the data from this sample in terms of the *row percentages*. Rounded to the nearest whole number for clarity, these are shown in Table 2.

Presenting the data in this way shows more clearly how the response to the question depends on the mammography experience. Another illuminating way to show this effect would be by using three

Table 1 Mammography: the counts

Mammography experience	Detection of breast cancer		
	Not likely	Somewhat likely	Very likely
Never	13	77	144
Over one year ago	4	16	54
Within the past year	1	12	91

Table 2 Mammography: the row percentages

Mammography experience	Detection of breast cancer			
	Not likely	Somewhat likely	Very likely	Total
Never	6	33	61	100
Over one year ago	5	22	73	100
Within the past year	1	11	87	100

2 Contingency Table

bar charts, overlaid on the same graph, with each corresponding to a different row of the table. There are several possible analyses we can carry out on this particular table, such as a simple summary in terms of percentages, a graphical summary, or a statistical test of independence, which we describe below. None of these analyses can establish whether the association between rows and columns is a *causal* one. This is an inherent limitation of the purely statistical approach. In the current example, the column variable might be thought of as a response to the row variable, which is an explanatory variable, or **covariate**. But the simple two-way contingency table cannot directly tell us the mechanism by which the explanatory variable affects the response variable. In any case, in the current example, there could be a mixture of several complex mechanisms which act simultaneously.

The Hypothesis of Independence, H_0 , and the Chi-square Test

With the notation introduced above, that

$$p_{ij} = \Pr(\text{row} = i, \text{column} = j) \quad (1)$$

for $i = 1, \dots, I$, $j = 1, \dots, J$, we see that the hypothesis of independence of rows and columns may be written as

$$H_0 : p_{ij} = p_{i+}p_{+j} \quad (2)$$

for all i, j , where $p_{i+} = \sum_j p_{ij}$ is the probability that Row = i , and $p_{+j} = \sum_i p_{ij}$ is the probability that Column = j .

Observe that an equivalent way of writing H_0 is

$$H_0 : \frac{p_{ij}}{p_{i+}} = p_{+j} \quad (3)$$

for all i, j . Now p_{ij}/p_{i+} is the probability that the column is j , conditional on the row being i : we write this as $P(\text{Column} = j | \text{Row} = i)$. Thus the hypothesis of independence of rows and columns is equivalent to the statement that the distribution of the variable column, conditional on the value of the variable row being i , is the same for all values of i : in this case we say that these conditional distributions are *homogeneous*.

To test the null hypothesis of independence H_0 , we compute Pearson's χ^2 statistic, which is defined as

$$\chi^2 = \sum_i \sum_j \frac{(n_{ij} - e_{ij})^2}{e_{ij}} \quad (4)$$

where e_{ij} are conventionally called the *expected values*, and are defined by

$$e_{ij} = \frac{(n_{i+}n_{+j})}{n} \quad \text{for all } i, j \quad (5)$$

and n_{i+} and n_{+j} are the row and column totals, respectively, defined by $n_{i+} = \sum_j n_{ij}$ for all i and $n_{+j} = \sum_i n_{ij}$ for all j . (A pedantic but more accurate description of (e_{ij}) is “the maximum-likelihood estimates of the expected values of (n_{ij}) under H_0 ”.) We know that, under the null hypothesis H_0 , the distribution of χ^2 is, for large n , approximately χ_f^2 , where f , the **degrees of freedom**, are $(I - 1)(J - 1)$. Hence for a significance test of H_0 with approximate **significance level** α , we reject H_0 if

$$\chi^2 \geq \chi_f^2(\alpha), \quad (6)$$

the right-hand side being the upper α point of the χ_f^2 distribution, which we can find from statistical tables: typically we would take α as 0.01 or 0.05. Thus we have approximately arranged that, for large n , our test gives

$$\Pr(\text{reject } H_0 | H_0 \text{ true}) \leq \alpha \quad (7)$$

This is what we mean by saying that the test is of approximate *size* α .

χ^2 tests are described in depth in another article. The value of the χ^2 statistic for the numerical example given above is 24.15, which lies well above 18.47, the 0.001 point of the χ_4^2 distribution. Thus for these data, the null hypothesis of independence of the women's response to the question about the detection of breast cancer and their “Mammography Experiences” is rejected, as our preliminary look at the table suggested would be the case.

From the point of view of scientific enquiry, a statistical significance test is a crude and blunt instrument. If we found that the value of χ^2 was significant, we would usually want to investigate in some detail *why* the null hypothesis of independence fails to fit. It is good statistical practice to compare the *adjusted residuals*

$$\frac{(n_{ij} - e_{ij})}{[e_{ij}(1 - \hat{p}_{i+})(1 - \hat{p}_{+j})]^{1/2}} \quad (8)$$

where $\hat{p}_{i+} = (n_{i+}/n)$, $\hat{p}_{+j} = (n_{+j}/n)$, with the standard **normal distribution**. Thus any residual greater than about 1.5 in absolute value is “interesting”, and

the corresponding cell (i, j) of the table is special in some way and deserves investigation.

Under the null hypothesis H_0 , the sufficient statistics for the unknown parameters $((p_{i+}), (p_{+j}))$ are $((n_{i+}), (n_{+j}))$. One consequence of this is that the observed values of these statistics must agree exactly with the corresponding expected values: thus

$$n_{i+} = e_{i+} \quad \text{for each } i, \text{ and } n_{+j} = e_{+j} \quad \text{for each } j \quad (9)$$

This is a special case of a general result for exponential families, and we shall see examples of the same result when we consider multiway tables in the final section.

A test statistic which is equivalent to χ^2 for large sample size is the *deviance*, conventionally denoted by G^2 . It is derived from the ratio of maximized likelihoods. For the example of testing H_0 , G^2 is defined by

$$G^2 = 2 \sum \sum n_{ij} \log \left(\frac{n_{ij}}{e_{ij}} \right) \quad (10)$$

and for the numerical example given above G^2 has the value 26.80; recall that the value of χ^2 was 24.15.

Both χ^2 and G^2 are appropriate test statistics for sampling setups other than the straightforward single multinomial one introduced above. For example, suppose the contingency table (n_{ij}) were collected with the row totals

- (n_{i+}) fixed, and
- $(n_{ij}|n_{i+})$ independently and multinomially distributed, with
- (n_{ij}) multinomial parameters $n_{i+}, (\theta_{ij})$,
- where $\sum_j \theta_{ij} = 1$ for each i .

Then the statistics χ^2 and G^2 are the appropriate test statistics for testing $H: \theta_{ij} = \phi_j$ for all i, j , for some (unknown) ϕ_j such that $\sum \phi_j = 1$; in other words, the null hypothesis of homogeneity of row distributions.

Yet another way of writing H_0 is

$$H_0 : \log(p_{ij}) = \alpha_i + \beta_j \quad \text{for all } i, j \quad (11)$$

for some α, β such that $\sum p_{ij} = 1$. One reason for writing H_0 in this form is that it enables us to see that it is a *loglinear* hypothesis; that is, the log of the probabilities can be written as a function which is linear in the unknown parameters, which in this case are $(\alpha_i), (\beta_j)$.

Computational Aspects

Contingency table analysis, both by χ^2 tests and by small-sample “exact” methods, is easily achieved in many statistical software packages. Usually we can ask for direct computation of the appropriate χ^2 statistic, with associated residuals and significance tests. More sophisticated software will allow us to test complex hypotheses of independence by fitting models using one or both of iterative proportional fitting (IPF) or generalized linear modeling (GLM). IPF may be computationally more efficient than the GLM approach, since in using GLM for contingency table analysis we are ignoring the particular structure of the parameters. For example, in using GLM to test for independence in a two-way contingency table, we use an iterative technique (effectively the *Newton–Raphson*) to solve the maximum-likelihood equations when in fact five minutes with a pencil and paper would show us that iteration is unnecessary: there is a closed-form solution. But an advantage of the GLM approach is that it is more suited to the problem of choosing the simplest possible model consistent with a given dataset. Most statisticians are familiar with the use of GLM in any case, and so this advantage outweighs the possible gain in computational efficiency of IPF. The GLM approach may be rather more forbidding for the less mathematical user, but has definite advantages when we come to consider cross tabulations which are more complex than just the two-way ones. We discuss *multiway* contingency tables in the last section of this article. Although there are several possible probability distributions available for GLMs, there is no multinomial distribution. This is not a problem, provided that we are fitting a *log-linear* model, such as H_0 as written in its final form above. Thus for testing independence in a two-way contingency table we can use the Poisson distribution as a “surrogate” for the multinomial, so that we can compute the appropriate deviance for testing independence for the multinomial model by pretending that (n_{ij}) are observations on independent Poisson

4 Contingency Table

variables. This is a special case of a general result, relating Poisson and multinomial loglinear models.

Quasi-Independence

For some two-way contingency tables with a special structure, it may be obvious that the null hypothesis of independence H_0 cannot possibly be expected to fit, because it is far too “severe”. However, a weaker hypothesis which is chosen to represent a version of independence, or quasi-independence, may be helpful in the interpretation of the data. For example, Altham [2] discusses the data on initial and final ratings of 121 stroke patients in Table 3, taken from [3].

The rows of the 5×5 table correspond to the patient’s state on admission to hospital, and the columns to his state when discharged. These states have possible values A, B, C, D , and E , ranging from A as the least severe to E as the most severe. The resulting contingency table is *triangular*, because the patient’s state on being discharged is never worse than his initial state. Because of this constraint, the hypothesis of independence of rows and columns cannot possibly hold. But if we assume that the frequencies (n_{ij}) for $1 \leq j \leq i \leq 5$ come from a multinomial distribution with corresponding parameters (p_{ij}) , then it may still be helpful to fit a hypothesis of *quasi-independence* H_{qi} , say:

$$H_{qi} : p_{ij} = \alpha_i \beta_j, \quad \text{for all } j \leq i \quad (12)$$

for some (α_i) , (β_j) . This is another loglinear hypothesis, and so may be tested by treating the (n_{ij}) as independent Poisson variables. For the data given here, the deviance for testing H_{qi} has value 9.60, $df = 6$, indicating a reasonable fit. Similar models may be appropriate for the type of “triangular” data arising from capture–recapture studies in ecology, where, for example, a bird first ringed in 1975 may

be observed in any one of the 10 successive years, but of course a bird first ringed in 1977 could not have appeared in the 1975 count.

Quasi-independence models for square contingency tables (omitting the diagonal entries) are known as *mover–stayer* models in the sociological context. This hypothesis of quasi-independence is one of several models especially suitable for a square contingency table; others in this class are models of symmetry, marginal homogeneity, and quasisymmetry.

Exact Tests of Independence

For tables with small frequencies, the large-sample approximation of the distribution of χ^2 , or equivalently of the distribution of G^2 , to the χ^2 distribution may be doubtful. Most software will give appropriate “warning messages” if the “expected counts” e_{ij} are too small, less than five being a conservative interpretation of “too small” here. In this case the statistician will need to use an exact method. For a 2×2 table, Fisher’s exact test, which is based on the hypergeometric distribution, is appropriate, and it has been generalized, say, to $I \times J$ tables by using the multivariate hypergeometric distribution: this is the result of conditioning the multinomial distribution on both the row and column frequency totals, under the null hypothesis of independence of rows and columns.

For example, consider the 2×2 contingency table in Table 4.

In carrying out a χ^2 test on this table, we are essentially asking whether the proportion $8/16$ is different from the proportion $20/23$. The χ^2 statistic is 4.67, with corresponding P value = 0.0307. But the software warns that the expected counts may be too low for the large-sample approximation to be valid, and it is easy to see that e_{11} , which is $16 \times 11/39$, is 4.5. For the sake of comparison, we do a two-sided Fisher exact test, and find that the corresponding P value is 0.0272. So either method would lead us to reject the null hypothesis of independence of rows and columns.

Table 3 Stroke patients^(a)

	A	B	C	D	E	
A	5	–	–	–	–	5
B	4	5	–	–	–	9
C	6	4	4	–	–	14
D	9	10	4	1	–	24
E	11	23	12	15	8	69
	35	42	20	16	8	121

^(a) © Biometrics Society

Table 4 Example for Fisher’s exact test

8	8
3	20

Multiway Contingency Tables

An example of a four-way contingency table is given in Table 5. These data have been slightly adapted for pedagogic purposes from a dataset on adolescents supplied by Professor I.J. Goodyer of Cambridge University, Developmental Psychiatry. For further discussion of the original data, see [4].

One purpose in collecting such a dataset is to find a model that explains how the four variables “Gender”, “Depression status”, “Behavior”, and “Anxiety” are interrelated.

For example, we can see from the two-way marginal Table 6 that the prevalence of depression among girls is 17%, compared with only 8% as the corresponding figure for boys. (The χ^2 statistic for this 2×2 table is 4.82, which is significant when referred to χ^2_1 .) Similarly, the two-way marginal Table 7, for which the corresponding χ^2 statistic is 7.42, suggests that there is a strong positive association between anxiety and behavioral symptoms.

However, we can see that this piecemeal and *ad hoc* approach to the data analysis is rather unsatisfactory. Looking at the pairwise marginal tables can raise some interesting suggestions, but it may also

be misleading because it possibly conceals important features of the data, as we shall see later. Furthermore, even for a four-way contingency table there are six possible pairwise marginal tables to be examined, so that it is clear that for large tables, for example seven-dimensional, we need a more focused modeling strategy.

We return to the above practical example when we have discussed types of independence for multiway contingency in a formal way.

Suppose that the rows, columns, layers, etc., of the multiway table are labeled A, B, C , etc. There are many different types of independence between these variables that are possible. This makes analysis of multiway contingency tables interesting and complex. Fortunately, the relationship between the variety of types of independence and loglinear models fits naturally within the GLM framework. We will once again make use of the relationship between the Poisson and the multinomial in the context of loglinear models.

An example with only three variables, say A, B , and C , serves to illustrate the methods used in tables of dimension higher than two. Suppose that A, B , and C correspond, respectively, to the rows, columns, and layers of the three-way table. Let

$$p_{ijk} = \Pr(A = i, B = j, C = k) \quad \text{for } i = 1, \dots, I, \\ j = 1, \dots, J, k = 1, \dots, K \quad (13)$$

so that $\sum p_{ijk} = 1$, and let (n_{ijk}) be the corresponding observed frequencies, assumed to be observations from a multinomial distribution, parameters $n, (p_{ijk})$.

There are eight different hypotheses corresponding to types of independence between A, B , and C that we now consider. Assume in all of these that the parameters given are such that $\sum p_{ijk} = 1$.

$$H_0 : p_{ijk} = p_{i++}p_{+j+}p_{++k} \quad \text{for all } i, j, k \quad (14)$$

thus H_0 corresponds to A, B , and C independent.

$$H_1 : p_{ijk} = p_{i++}p_{+jk} \quad \text{for all } i, j, k \quad (15)$$

thus H_1 corresponds to A independent of (B, C) .

(Likewise, we could consider the hypothesis B independent of (A, C) , and the hypothesis C independent of (A, B) .)

$$H_2 : \frac{p_{ijk}}{p_{i++}} = \left(\frac{p_{ij+}}{p_{i++}} \right) \left(\frac{p_{i+k}}{p_{i++}} \right) \\ \text{for all } i, j, k \quad (16)$$

Table 5 2^4 table from psychiatry

		Behavioral symptoms		No		Yes	
		Anxiety symptoms		No	Yes	No	Yes
Girls	Depression = no			85	14	10	2
Girls	Depression = yes			8	4	7	4
Boys	Depression = no			107	13	8	3
Boys	Depression = yes			2	3	3	4

Table 6 A pairwise summary from Table 5

	Depression	
	No	Yes
Girls	111	23
Boys	131	12

Table 7 Another pairwise summary from Table 5

		Anxiety symptoms	
		No	Yes
Behavioral symptoms	No	202	34
Behavioral symptoms	Yes	28	13

6 Contingency Table

Thus H_2 is equivalent to

$$\Pr(B = j, C = k | A = i) = \Pr(B = j | A = i) \times \Pr(C = k | A = i) \quad \text{for all } i, j, k \quad (17)$$

and so H_2 corresponds to the hypothesis that, for each i , conditional on $A = i$, the variables B and C are independent. In this case we say that B and C are independent, conditional on A . (Likewise, we can define two similar hypotheses by interchanging A , B , and C .)

Independence and conditional independence between variables can be more clearly understood by showing suitable *association graphs*. In this graphical representation of the hypotheses, we represent the variables A , B , and C by the *vertices* of the graph, with links between these vertices representing dependence, and absence of links representing independence. Thus

H_0 corresponds to no links between any of A , B , and C .

H_1 corresponds to a link between B and C only, and A as an isolated point.

H_2 corresponds to a link between A and B , and a link between A and C : thus B and C are linked only through A ; see the diagram below.

The theory of association graphs is a useful and elegant one, particularly suitable for displaying and understanding the structure of high-dimensional tables.

Finally, we consider

$$H_3 : p_{ijk} = \alpha_{jk}\beta_{ik}\gamma_{jk} \quad \text{for all } i, j, k, \quad \text{for some } \alpha, \beta, \gamma \quad (18)$$

This hypothesis, which is symmetric in A , B , and C , cannot be given an interpretation in terms of conditional probability (see Figure 1). It corresponds to saying that the *interaction* between any two of the three factors, say A and B , given the level of the third factor, say C , is independent of the level of C . We say that H_3 corresponds to “no three-way interaction” between A , B , and C .

One way of writing this formally is to say that for each i, j the odds ratio describing the dependence

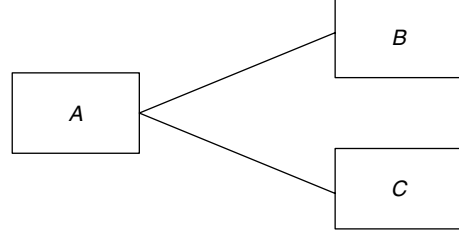


Figure 1 B and C conditionally independent, given A

between A and B

$$\frac{(p_{ijk}p_{Ijk})}{(p_{iJk}p_{Ijk})} \quad (19)$$

is the same for all k .

The eight hypotheses are related to one another as follows: for any given probabilities (p_{ijk}) ,

H_0 implies H_1 , which in turn implies H_3 , and
 H_0 implies H_2 , which in turn implies H_3 , and
 H_1 and H_2 are together equivalent to H_0 .

Thus H_0 is the strongest hypothesis and H_3 is the weakest.

All of the eight hypotheses above may be written as loglinear hypotheses and hence tested within the GLM framework with the Poisson distribution and log link function. Since the log is the *canonical link function* for the Poisson distribution, it is the *default* link function for the Poisson in GLM terms.

For example, we may rewrite H_2 as

$$H_2 : \log(p_{ijk}) = \phi_{ij} + \psi_{ik} \quad \text{for all } i, j, k \quad (20)$$

for some ϕ, ψ . In the GLM notation for interactions between factors, this corresponds to the model

$$A * B + A * C \quad \text{or, equivalently, } A * (B + C)$$

We pause to explain briefly the GLM notation for interactions between factors. For example, consider the model

$$A * B + A * C$$

In the context of loglinear models, this means that we can write $\log(p_{ijk})$ as

$$\log(p_{ijk}) = \mu + \alpha_i + \beta_j + \gamma_k + (\alpha\beta)_{ij} + (\alpha\gamma)_{ik} \quad (21)$$

for all i, j, k . Here the set of parameters $[(\alpha\beta)_{ij}]$ is termed $A.B$, the AB interaction, the set $[(\alpha\gamma)_{ik}]$ the AC interaction, and the sets (α_i) , (β_j) , (γ_k) are, respectively, the main effects of A , B , and C denoted by A , B , C .

For parameter identifiability in our computations we will need to impose constraints on (α_i) , and so on. We assume in the numerical examples that follow that the *corner-point* constraints are imposed; thus

$$\alpha_1 = 0, \quad \beta_1 = 0, \quad \gamma_1 = 0, \quad (\alpha\beta)_{1j} = 0 \quad (22)$$

and so on, so that any parameter with a subscript 1 anywhere is set to zero. Different software may impose a different set of constraints: this is always a confusing point for a beginner.

In the same GLM notation, H_0 , H_1 , and H_3 correspond, respectively, to $A + B + C$, $A + B * C$, and $(B * C + A * B + A * C)$.

For completeness, we will also define the *saturated* model H_s , say, which makes no independence statement about A , B , and C :

$$H_s: \log(p_{ijk}) = \rho_{ijk} \quad \text{for all } i, j, k \quad (23)$$

for some ρ : in GLM notation this is simply

$$A * B * C$$

In terms of our association graph, this corresponds to all three links being present between the vertices A , B , and C .

Example

Consider the $2 \times 2 \times 2$ table shown in Table 8.

For example, A might correspond to “agree with a controversial statement” (such as “women are not inherently better than men at looking after children”, B might be the educational level (below or above average), and C might be the gender (men/women).

Table 8 A 2^3 table

	$C = 1$		$C = 2$	
	$A = 1$	$A = 2$	$A = 1$	$A = 2$
$B = 1$	17	23	36	50
$B = 2$	29	14	59	24

As our baseline model, we first fit the saturated model

$$A * B * C$$

This is bound to give a perfect fit, with deviance and df both zero. So this first step in the model fitting might not appear to be a useful exercise. However, inspection of the resulting *parameter estimates* together with their *standard errors* suggests that the three-factor interaction term, $A.B.C$ can be dropped from the model, since the estimate of $A.B.C$ is -0.197 with $se = 0.561$. So our next step is to fit $(A + B + C) * (A + B + C)$, which gives a deviance of 0.12, with 1 df. Comparison with χ^2_1 shows that this model, which is H_3 , is a good fit. Thus there is no three-way interaction between A , B , and C , and so, for example, we can say that the way in which the response to the question A depends on educational level is the same for both men and women.

Again, we compare the parameter estimates with their standard errors, and find that the pairwise interaction $A.C$ is -0.068 , with $se = 0.28$. This suggests that $A.C$ should now be dropped from the model, and indeed the resulting model $(A + C) * B$ has deviance 0.18 (2 df). So this model is a good fit. It states that conditional on the educational level, the response to A is independent of the gender.

This model can be rewritten as $A * B + B * C$, and this has the consequence that the *observed and fitted frequencies* for the two-way marginal table $A * B$ must agree exactly, and similarly for the $B * C$ marginal table.

Inspection of the parameter estimates and their standard errors for this model shows that $B.C$ can also be dropped from the model, leading to the model $A * B + C$, which has deviance 0.34 (3 df). But this final step would not make sense if the data were collected with *fixed* totals in the $B.C$ marginal table, since under the model $A * B + C$ we would not necessarily have exact agreement between the observed and fitted frequencies of the $B.C$ table.

The Relationship between Binomial Logistic Regression and Loglinear Models in a Multiway Contingency Table

In a multiway contingency table, it may not be appropriate to treat the variables, say A , B , C , etc., symmetrically. For example, it may be more natural

8 Contingency Table

to treat A as a *response* variable, and B , C , etc., as *explanatory* variables.

In particular, if the number of levels of A is two, for example, corresponding to “yes, no”, then it may make the analysis easier to interpret if we carry out a binomial logistic regression of A on the factors B , C , etc.

Such an analysis is not essentially different from a loglinear analysis. We can see from the following considerations that there must be certain exact correspondences between the two approaches. To be specific, take the case in which (Y_{ijk}) is multinomial, parameters n , (p_{ijk}) and suppose $i = 1, 2$. Write y_{+jk} as $y_{1jk} + y_{2jk}$. Then $Y_{1jk}|y_{+jk}$ are independent binomial variables, parameters y_{+jk} , θ_{jk} , where

$$\theta_{jk} = \frac{p_{1jk}}{p_{+jk}} \quad (24)$$

So, for example, the model $A * B + B * C + C * A$ for (p_{ijk}) can be shown to be equivalent to the model for logit $\theta_{jk} = \log[\theta_{jk}/(1 - \theta_{jk})]$,

$$\text{logit } \theta_{jk} = \beta_j + \gamma_k \quad (25)$$

For example, we could use the data from Table 8 above, with A as the response variable, so that we use the binomial proportions 17/40, 29/43, 36/86, and 59/83 as the responses corresponding to the factors (B,C) as (1,1), (2,1), (1,2), and (2,2). In this case it may be seen that the deviance and the fitted frequencies for the logit model $B + C$ are *exactly* the same as those for loglinear model $A * B + B * C + C * A$ with data for the $2 \times 2 \times 2$ and the multinomial model, as above.

We now return to the 2^4 contingency table, Table 5. The simplest model consistent with these data is (gender + anxiety symptoms + behavioral symptoms) * depression. This model has deviance 4.91, with 8 df. For a relatively complex table, a good way to find the simplest reasonable model is to start by fitting the saturated model, and then “step down”, discarding the unnecessary parameters, using the Akaike Information Criterion (AIC) as the guide to when to stop the stepping down process.

The final model here shows the key role of depression in explaining the dependence in the four-way table: conditional on the depression status, the two variables anxiety and behavior are independent of each other and are also independent of gender.

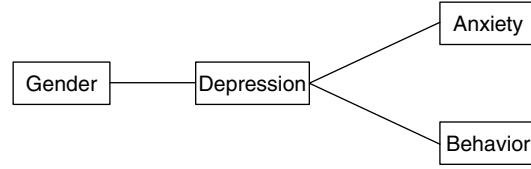


Figure 2 A graphic model for psychiatry data

Its graphical representation is a graph in which the three vertices “gender”, “anxiety symptoms”, and “behavioral symptoms” are linked to depression, and no other links are present (see Figure 2).

This final model also shows that summarizing the data by certain 2×2 tables, for example by *collapsing* over the variable gender and depression, may be a misleading way to investigate the association between variables. To put this more technically, for the current example with the final model, the two-way marginal table anxiety symptoms $\times$ behavioral symptoms is not among the sufficient statistics for the unknown parameters.

The danger in obtaining misleading results by collapsing a contingency table is known as *Simpson’s paradox* (also known as *Yule’s paradox*).

This brief article has introduced the topic of contingency tables, with particular discussion of types of independence for multiway tables. The reader will have realized that in none of the methods described above is any attention paid to the *order* of the categories of a variable: for example, if A takes as possible values 1, 2, 3, and 4, the above analyses would all remain the same if these levels were renamed as *apples*, *pears*, *oranges*, *bananas*. In this sense our treatment of the variables so far has been entirely *nominal*. This approach may not always be the most sensible one to follow. For example, both rows and columns of Table 1 are ordered. We might prefer our analysis to recognize this fact.

This article represents a brief introduction to the topic of contingency tables, and is not an exhaustive coverage of the available methods. For more comprehensive coverage, the reader is referred to the two textbooks by Agresti [5, 6].

References

- [1] Hosmer, D.W. & Lemeshow, S. (1989). *Applied Logistic Regression*, John Wiley & Sons, New York.

-
- [2] Altham, P.M.E. (1975). Quasi-independent triangular contingency tables, *Biometrics* **31**, 233–238.
- [3] Bishop, Y.M.M. & Fienberg, S.E. (1969). Incomplete two-dimensional contingency tables, *Biometrics* **25**, 119–128.
- [4] Altham, P.M.E. (1992). Fitting graphical models to multi-way contingency tables in GLIM, *GLIM Newsletter* **21**, 4–8.
- [5] Agresti, A. (2002). *Categorical Data Analysis*, 2nd Edition, John Wiley & Sons, New York.
- [6] Agresti, A. (1996). *An Introduction to Categorical Data Analysis*, John Wiley & Sons, New York.

P.M.E. ALTHAM

Article originally published in Encyclopedia of Biostatistics, 2nd Edition (2005, John Wiley & Sons, Ltd). Minor revisions for this publication by Jeroen de Mast.

Statistical Efficiency

Introduction

In 1922, Ronald Fisher [1], one of the founders of modern statistics, wrote his fundamental paper. *On the Mathematical Foundations of Theoretical Statistics*. In that paper he stated that “The object of statistical method is the reduction of data.” He then identified “three problems which arise in the reduction of data”. These are

1. *Specification* – choosing the right mathematical model for a population.
2. *Estimation* – methods to calculate, from a sample, estimates of the parameters of the hypothetical population.
3. *Distribution* – properties of statistics derived from samples.

Later, Colin Mallows, a distinguished researcher at AT&T Bell laboratories, added a “zeroth problem” – considering the *relevance of the observed data*, and other data that might be observed, to the substantive problem [2].

Kenett and Thyregod [3] present the statistical consulting cycle to solve real problems as a five-step cycle:

- (1) problem elicitation,
- (2) data collection,
- (3) data analysis,
- (4) findings formulation, and
- (5) findings presentation (see Figure 1).

Academia typically focuses only step (3), the data analysis step. The application of statistics requires the full cycle. The three problems addressed by Fisher focus on steps (2) and (3). The zeroth problem identified by Mallows relates to steps (1) and (2). Practical

statistical efficiency, introduced below, addresses the full life cycle.

Definition of Practical Statistical Efficiency

To expand the view of statistical methods, Godfrey [4, 5] and later, Kenett *et al.* [6, 7] suggest evaluating statistical work using a formula of practical statistical efficiency (PSE). The PSE formula of Kenett *et al.* [6] accounts for eight components and is computed as follows:

$$PSE = V\{D\} \times V\{M\} \times V\{P\} \times V\{PS\} \times P\{S\} \times P\{I\} \times T\{I\} \times E\{R\} \quad (1)$$

where

$V\{D\}$: value of the data actually collected,

$V\{M\}$: value of the statistical method employed,

$V\{P\}$: value of the problem to be solved,

$V\{PS\}$: value of the problem actually solved,

$P\{S\}$: probability level that the problem actually gets solved,

$P\{I\}$: probability level the solution is actually implemented,

$T\{I\}$: time the solution stays implemented, and

$E\{R\}$: expected number of replications.

These components can be assessed qualitatively, using expert opinions, or quantitatively, if the relevant data exists.

A straightforward approach to evaluate PSE is to use a scale of “1” for not very good to “5” for excellent. This method of scoring can be applied uniformly for all PSE components. Some of the PSE components can be also assessed quantitatively. $P\{S\}$ and

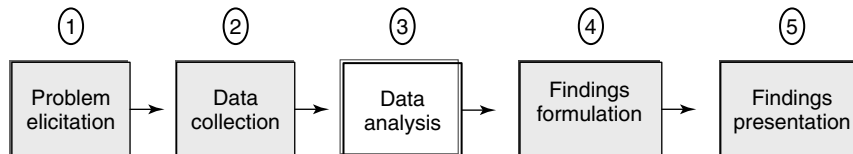


Figure 1 The statistical consulting cycle

2 Statistical Efficiency

$P\{I\}$ are probability levels, $T\{I\}$ can be measured in months, $V\{P\}$ and $V\{PS\}$ can be evaluated in euros, dollars, or pounds. $V\{PS\}$ is the value of the problem actually solved, as a fraction of the problem to be solved. If this is evaluated qualitatively, a large portion would be “4” or “5”, a small portion “1” or “2”. $V\{D\}$, the value of the data actually collected is related to Fisher’s zero problem presented in Mallows [2]. Whether PSE terms are evaluated quantitatively or qualitatively, PSE is a conceptual measure rather than a numerically precise one. A more elaborated approach to PSE evaluation can include differential weighing of the PSE components and/or nonlinear assessments.

A PSE Case Study

An old family-owned company decided, in 1998, to modernize their management of the company. Key performance indicators (KPIs) including efficiency, staff absence, and waste were adopted and the concept proved so useful that the original 39 KPIs were increased to 50. There were several early wins, for example a 90% reduction in out loading errors to a major customer and an improvement in performance of parallel shifts from 54 and 68 to 80% each. A questionnaire was given to staff to find out how they felt about the KPIs. It was found that staff did not like the staff absence KPIs because they felt powerless to influence them but did relate to waste KPIs because they are traditionally given in real terms rather than percentage. So “staff absence” was changed to staff attendance and KPIs focused on output using both quantities and money made and lost. Reporting changed from displaying the daily performance of 10 machines on a page to displaying weekly performance of one machine per page. This meant that charts could be placed near to machines and operators could observe their performance. It also led to a clearer understanding of causes of downtime. It was then decided to take the plunge and change from reporting efficiency to reporting utilization as a much more realistic and meaningful way to assess performance. Performance appeared to drop from 80 to 70% but it was felt that staff were ready for the psychological blow and eventually performance on this measure is gradually improved. The tea packers embraced the quality improvement project wholeheartedly and initially had a specific problem to solve: they had two

apparently identical packing lines and yet one produced 54% efficiency and the other produced 68% with the same staff. The problem was eventually located to the availability of the engineers to mend the line when problems occurred. The PSE components for this case study have been assessed qualitatively as

$V\{D\}$: value of the data actually collected = 4,

$V\{M\}$: value of the statistical method employed = 3,

$V\{P\}$: value of the problem to be solved = 4,

$V\{PS\}$: value of the problem actually solved = 4,

$P\{S\}$: probability level that the problem actually gets solved = 5,

$P\{I\}$: probability level the solution is actually implemented = 5,

$T\{I\}$: time the solution stays implemented = 3, and

$E\{R\}$: expected number of replications = 4.

so that $PSE = 57\,600$.

Kenett *et al.* [6] refer to five case studies (including the one above), and derive qualitatively scores for the eight components of PSE. The cases are referred to as *tea*, *drink*, *metal*, *utility*, and *filter* and the results are shown in Table 1.

A radar plot of these scores is presented as Figure 2.

It can be seen from Figure 2 that the scores are generally high for the longevity ($E\{R\}$ and $T\{I\}$) and probability components for all five case studies and the difference comes in the values. In particular the $V\{D\}$, value of the data. This value reflects the effort put into the project. A similar comment applies to the $V\{M\}$, value of the method used. When the problem is of low priority, resources are not made available for

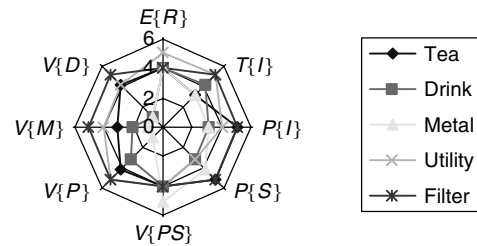


Figure 2 Practical statistical efficiency evaluation of the five case studies

Table 1

Case	$E\{R\}$	$T\{I\}$	$P\{I\}$	$P\{S\}$	$V\{PS\}$	$V\{P\}$	$V\{M\}$	$V\{D\}$	PSE
Tea	4	3	5	5	4	4	3	4	57 600
Drink	4	4	3	3	4	3	2	1	3456
Metal	4	3	3	4	5	1	1	1	720
Utility	5	5	4	3	4	5	4	4	96 000
Filter	4	5	5	5	4	5	5	5	250 000

methodology to be as good as possible. For example, in the heavy metal company, the time available to the operative involved and his expertise limited the methodology to the most basic charts.

In all of these case studies, the $V\{PS\}$ value of the problem actually solved is high. In the less successful case studies, the worth was not appreciated and the resources required to deal with the side effects were not available.

The Statistical Efficiency Conjecture

Kenett *et al.* [7] study empirically 21 case studies where PSE is evaluated along with the maturity level of the industrial organization from a management perspective. They formulate the *statistical efficiency conjecture* that links management maturity with the impact level of problem solving and improvements driven by statistical methods. The maturity level to the management of industrial organizations can be summarized and classified using a four-step *quality ladder* [8]. The four approaches are (a) fire fighting, (b) inspection, (c) process control, and (d) quality by design. To each management approach, corresponds a particular set of statistical methods and the quality ladder maps each management approach with appropriate statistical methods. The statistical efficiency conjecture states that organizations higher up on the quality ladder achieve higher values of PSE. The implication being that moving up the quality ladder is justified on practical grounds using effectiveness and efficiency criteria.

References

- [1] Fisher, R.A. (1922). On the mathematical foundations of theoretical statistics, *Philosophical Transaction of the Royal Society. Series A* **222**, 309–368.
- [2] Mallows, C. (1998). The zeroth problem, *The American Statistician* **52**, 1–9.
- [3] Kenett, R.S. & Thyregod, P. (2006). Aspects of statistical consulting not taught by academia, *Statistica Neerlandica* **60**(3), 396–412.
- [4] Godfrey, A.B. (1988). Statistics, quality and the bottom line, part 1, *ASQC Statistics Division Newsletter* **9**(2), 211–213.
- [5] Godfrey, A.B. (1989). Statistics, quality and the bottom line, part 2, *ASQC Statistics Division Newsletter* **10**(1), 14–17.
- [6] Kenett, R.S., Coleman, S. & Stewardson, D. (2003). Statistical efficiency: the practical perspective, *Quality and Reliability Engineering International* **19**, 265–272.
- [7] Kenett, R.S., de Frenne, A., Tort-Martorell, X. & McCollin, C. (2008). The statistical efficiency conjecture, in *Applying Statistical Methods in Business and Industry – The State of the Art*, S. Coleman, T. Greenfield, & D. Montgomery, eds, John Wiley & Sons.
- [8] Kenett, R.S. & Zacks, S. (1988). *Modern Industrial Statistics: Design and Control of Quality and Reliability*, Duxbury Press, San Francisco, Thomson learning 2nd edition 2003, Chinese edition 2004.

RON S. KENETT

Sigma Level of a Process

The sigma level is a measure of the quality of a characteristic of a process or product. It is based on defects per million opportunities (DPMO), and, for continuous metrics only, is also related to process

capability. Higher quality (that is, lower DPMO) corresponds to higher sigma level. For example, sigma levels of 2, 3, 4, 5, and 6, correspond to DPMO values of 308 538 66 807 6210 233 and 3.4, respectively.

Defects per Million Opportunities

Defects per million opportunities (DPMO) is a measure of the quality of a characteristic of a process or product. DPMO is calculated by determining the number of defects found in a sample according to the equation

$$DPMO = \frac{(\text{Defects} \times 1 \text{ million})}{(\text{Number of units in the sample} \times \text{Opportunities per unit})} \quad (1)$$

Here, opportunities per unit is defined to be the number of defects that could possibly occur in a single unit of a product or a single execution of a process.

C_p, C_{pk}

C_p and C_{pk} are indices used to express the capability to meet the specifications for a product or process characteristic. C_p and C_{pk} are calculated for stable, normally distributed data according to the equations

$$C_p = \frac{(\text{Upper specification limit} - \text{Lower specification limit})}{6 \times (\text{Short-term process standard deviation})} \quad (1)$$

$$C_{pk} = \frac{\text{Min}(\text{Upper specification limit} - \text{Mean}, \text{Mean} - \text{Lower specification limit})}{3 \times (\text{Short-term process standard deviation})} \quad (2)$$

Note that C_{pk} accounts for the location of the data (through its average or mean value), whereas C_p does not. C_{pk} is defined for either one- or two-sided specification limits, whereas C_p is defined only for

two-sided limits. For these reasons, C_{pk} is usually the preferred capability metric in practice. C_p is sometimes called the *potential process capability*, since, in the case of two-sided specifications, it provides the value that C_{pk} would attain for a perfectly centered process. Because these indices are calculated using a measure of short-term variability, they will generally

be higher than the corresponding long-term capability measures, P_p and P_{pk} .

P_p, P_{pk}

P_p and P_{pk} are indices used to express the capability to meet specifications for a product or process characteristic. Higher values denote better capability. P_p and P_{pk} are calculated for stable, normally distributed data according to the equations:

Note that P_{pk} accounts for the location of the data (through its average or mean value), whereas P_p does not. P_{pk} is defined for either one- or two-sided specification limits, whereas P_p is defined only for two-sided limits. For these reasons, P_{pk} is usually the preferred metric in practice. Because these indices are calculated using a measure of long-term variability, they will generally be lower than the corresponding short-term capability measures, C_p and C_{pk} (see C_p, C_{pk}).

$$P_p = \frac{(\text{Upper specification limit} - \text{Lower specification limit})}{6 \times (\text{Long-term process standard deviation})} \quad (1)$$

$$P_{pk} = \frac{\text{Min}(\text{Upper specification limit} - \text{Mean}, \text{Mean} - \text{Lower specification limit})}{3 \times (\text{Long-term process standard deviation})} \quad (2)$$

Defects per Unit

Defects per unit (DPU) is a measure of the quality of a characteristic of a process or product. DPU is calculated by determining the total number of defects

found in a sample according to the equation:

$$DPU = \frac{\text{Total defects}}{\text{Number of units observed}} \quad (1)$$

Accelerated Life Tests: Bayesian Models

Introduction and Overview

When dealing with high-reliability components/systems, it is not unusual to conduct tests in a more severe environment than the actual use environment. Such tests are called *accelerated life tests* (ALTs) and are usually conducted with the purpose of inducing early failures to reduce the time and cost of testing. The important statistical issue in ALTs is to make inferences about the failure behavior of the items at the user environment based on **failure data** obtained under more severe environments. In recent years there has been an increasing interest in developing Bayesian methods to make inferences from ALTs; see for example Mazzuchi and Soyer [1], van Dorp *et al.* [2], Mazzuchi *et al.* [3], and van Dorp and Mazzuchi [4, 5]. As a result of using complicated time transformation functions and/or ramping functions in step-stress ALTs, Bayesian inferences from these models have become extremely difficult. To deal with intractable predictive and posterior distributions, the previous approaches used large-sample approximations or approximate methods of inference such as the linear Bayes. Only a few of the earlier Bayesian work in ALTs, such as Blackwell and Singpurwalla [6] and Mazzuchi and Soyer [1], have considered models that allow time transformation functions that change with stress environment.

In this chapter we consider Bayesian models that allow dynamic time transformation functions, and propose different modeling strategies. We focus specifically on hierarchical Bayes models, similar to those described by Lindley and Smith [7], and dynamic models as in Mazzuchi and Soyer [1], and develop Bayesian computational methods for making inferences using **Markov chain Monte Carlo** (MCMC) approaches. The attractive feature of the MCMC methods is that they enable generation of samples from the posterior and predictive distributions without having to obtain the exact distributional forms. Thus, unlike the previous approaches, we can perform Bayesian analyses fully by computing all the posterior and predictive distributions of interest.

In the following section, we present a basic parametric accelerated **life testing** model assuming a Weibull **life distribution** and several time transformation functions. We discuss the difficulties involved in computing the posterior and predictive distributions using conventional techniques and present alternative inference based on MCMC methods. In the section titled “Extensions of the Basic Parametric Model”, we consider hierarchical and Markov models as novel extensions of the basic parametric model to allow for dynamic time transformation functions and discuss inference and reliability predictions using MCMC methods. In the section titled “Analysis of Accelerated Life Testing Data”, we illustrate our approaches using two real-life testing examples, discuss model comparison, and investigate the appropriateness of dynamic time transformation functions.

A Parametric ALT Model

An accelerated test environment is created typically by increasing the level of one or more stress variables (temperature, voltage, etc.) to values that are higher than those at normal operating conditions. For simplicity, in our development we assume that the environment is defined by a single stress variable. Extension to the multivariable case is straightforward. To introduce some notation, let S_i denote the level of the stress variable at the i th accelerated test environment and assume that test will be conducted at K accelerated levels of the stress variable that are specified in advance. As noted before, our main objective is to make predictive inferences about the failure behavior of items at the use-stress environment, S_u , based on the data from the K accelerated testing environments, where $S_i > S_u$ for $i = 1, 2, \dots, K$. It is also possible to incorporate testing at the use environment with the proposed procedure, but in our development we assume that no testing is conducted at S_u .

We assume that under the i th accelerated test environment, the failure behavior of the items can be described by a Weibull model with density

$$p(x_i|\lambda_i, \beta) = \beta\lambda_i x_i^{\beta-1} e^{-\lambda_i x_i^\beta} \quad (1)$$

with scale parameter, $\lambda_i > 0$ and shape parameter, $\beta > 0$. We will denote the above model as $(X_i|\lambda_i, \beta) \sim \text{Weib}(\lambda_i, \beta)$. The failure rate of the

2 Accelerated Life Tests: Bayesian Models

model is given by

$$m_i(x_i | \beta, \lambda_i) = \beta \lambda_i x_i^{\beta-1} \quad (2)$$

where the values $\beta > 1 (<1)$ imply an increasing (decreasing) failure rate distribution, and when $\beta = 1$ the exponential model arises as the special constant failure rate case. The above model implies that the scale parameter, λ_i , depends on the stress environment, but the shape parameter, β , does not, which is a common assumption in the literature (see for example Nelson [8] and Mazzuchi *et al.* [3]). It is common to assume a functional relationship between the failure rate and applied stress level. Such a relationship is referred to as a *time transformation function*. The Arrhenius law and the power law models are two of the most commonly used time transformation functions (Mann *et al.* [9]).

Assuming a power law model, the relationship between the scale parameter and the applied stress level in the i th accelerated test environment is given by

$$\lambda_i = \theta_1 S_i^{\theta_2} \quad (3)$$

where $\theta_1 > 0$ and θ_2 are unknown coefficients. Alternatively, an Arrhenius law model assumes that

$$\lambda_i = \exp \{ \theta_1 - \theta_2 / S_i \} \quad (4)$$

In both cases, the aging effect is captured by the shape parameter β , whereas the effect of the stress environment on the failure behavior is captured by functional parameters θ_1 and θ_2 .

Bayesian Inference for the ALT Model

Assume that n_i items are tested for t_i units of time at the i th accelerated environment with level S_i . Let X_{ij} denote the time to failure of the j th item tested under the i th environment with realization x_{ij} , $j = 1, 2, \dots, r_i \leq n_i$, where r_i denotes the number of failures observed during the observation period $(0, t_i]$. As noted by Mazzuchi *et al.* [3], using this notation, we can easily consider the complete sample (no censoring) case, where testing continues until all items have failed, as well as the type I and type II censoring cases. Specifically, the no censoring case is obtained with $t_i = \infty$; type I censoring case is obtained with $t_i < \infty$; and type II censoring case where testing continues until the r th failure is obtained with $t_i = x_{i(r_i)}$, where $x_{i(r_i)}$ is the r_i th failure

time at the i th accelerated environment. In what follows, we will assume that testing is conducted simultaneously at all stress levels, but the case of sequential testing from one stress level to another can also be accommodated.

Let D_i denote the test data from i th accelerated stress environment, that is,

$$D_i = \{n_i, r_i, x_{i1}, \dots, x_{ir_i}\} \quad (5)$$

Assuming conditional independence of the failure times $\{x_{ij}\}$ given the stress levels $\{S_i\}$, and the parameters β, θ_1 , and θ_2 , the likelihood function of λ_i and β from the i th accelerated stress environment is obtained as

$$\mathcal{L}_i(\lambda_i, \beta; D_i) = (\beta \lambda_i)^{r_i} \left\{ \prod_{j=1}^{r_i} (x_{ij})^\beta \right\} \exp \{-\lambda_i T_i(\beta)\}, \quad i = 1, \dots, K \quad (6)$$

where

$$T_i(\beta) = \begin{cases} \sum_{j=1}^{r_i} (x_{ij})^\beta + (n_i - r_i) t_i^\beta & \text{type I censoring} \\ \sum_{j=1}^{r_i} (x_{ij})^\beta + (n_i - r_i) (x_{i(r_i)})^\beta & \text{type II censoring} \\ \sum_{j=1}^{n_i} (x_{ij})^\beta & \text{no censoring} \end{cases} \quad (7)$$

and λ_i is defined either by equation (3) or (4). We note that in the case of no censoring $r_i = n_i$ in equation (6). Thus, the likelihood function of β, θ_1 and θ_2 given all the data $D = \{D_1, D_2, \dots, D_K\}$ is

$$\mathcal{L}(\theta_1, \theta_2, \beta; D) = \prod_{i=1}^K \mathcal{L}_i(\lambda_i, \beta; D_i) \quad (8)$$

The Bayesian approach requires specification of a prior distribution for β, θ_1 , and θ_2 . The joint posterior distribution can then be obtained as

$$p(\theta_1, \theta_2, \beta | D) \propto \mathcal{L}(\theta_1, \theta_2, \beta; D) p(\theta_1, \theta_2, \beta) \quad (9)$$

where $p(\theta_1, \theta_2, \beta)$ is the joint prior distribution. As noted by Mazzuchi *et al.* [3], the posterior distribution in equation (9) cannot be obtained in any

tractable form for any choice of the prior $p(\theta_1, \theta_2, \beta)$ in equation (9). Thus, to perform a fully **Bayesian analysis**, we will use the Gibbs sampler, an MCMC method, which is easy to implement in this case. The Gibbs sampler enables us to generate samples from the posterior distribution $p(\theta_1, \theta_2, \beta|D)$, which can then be used to make assessments about the failure behavior at the use stress environment in our problem. This is achieved without actually computing the distributional form by obtaining successive drawings from the full conditional distributions $p(\theta_1|\theta_2, \beta, D)$, $p(\theta_2|\theta_1, \beta, D)$, and $p(\beta|\theta_1, \theta_2, D)$. The process starts with a vector of arbitrary starting values of the parameters, $(\theta_1^{(0)}, \theta_2^{(0)}, \beta^{(0)})$, and continues by iteratively generating samples from the full conditional distributions (see Gelfand and Smith [10] for more details).

We note that the implementation of the Gibbs sampler requires the availability of the full conditional distributions. If the full conditional distributions are not of any known form or if they do not exist in closed form, some type of random number generation method must be employed to facilitate the implementation of the Gibbs sampler. If the full conditional distributions are log concave, one can use the efficient adaptive rejection sampling (ARS) method of Gilks and Wild [11]. Unlike standard rejection sampling methods, ARS avoids any maximization steps and therefore can be easily adopted at each iteration of the Gibbs sampler.

In our development for the power law model we will assume a gamma prior for θ_1 , denoted as $\Gamma[a, b]$, and we will denote the priors for θ_2 and β as $p(\theta_2)$ and $p(\beta)$, respectively. Furthermore, we assume that *a priori* θ_1, θ_2 , and β are independent. Under the power law model, with the given choice of priors, the full conditional distributions are given by

$$p(\theta_1|\theta_2, \beta, D) \propto \left\{ \theta_1^{\sum_{i=1}^K r_i} \exp \left[-\theta_1 \sum_{i=1}^K S_i^{\theta_2} T_i(\beta) \right] \right\} \times \theta_1^{a-1} \exp \{-b\theta_1\} \quad (10)$$

$$p(\theta_2|\theta_1, \beta, D) \propto \left\{ \prod_{i=1}^K S_i^{r_i \theta_2} \right\} \exp \left[-\theta_1 \sum_{i=1}^K S_i^{\theta_2} T_i(\beta) \right] p(\theta_2) \quad (11)$$

and

$$p(\beta|\theta_1, \theta_2, D) \propto \beta^{\sum_{i=1}^K r_i} \left\{ \prod_{i=1}^K \left(\prod_{j=1}^{r_i} (x_{ij})^\beta \right) \right\} \times \exp \left[-\theta_1 \sum_{i=1}^K S_i^{\theta_2} T_i(\beta) \right] p(\beta) \quad (12)$$

Note that equation (10) implies $(\theta_1|\theta_2, \beta, D) \sim \Gamma \left[a + \sum_{i=1}^K r_i, b + \sum_{i=1}^K S_i^{\theta_2} T_i(\beta) \right]$, whereas for any reasonable prior form for $p(\theta_2)$ and $p(\beta)$, the posterior conditionals in equations (11) and (12) cannot be obtained as a familiar form. However, they are both log concave if the priors $p(\theta_2)$ and $p(\beta)$ are log concave since both conditional log likelihoods are log concave. Thus the ARS method can be implemented to sample from the posterior densities in equations (11) and (12) at each iteration of the Gibbs sampler.

Once a sample $\{\theta_1^{(j)}, \theta_2^{(j)}, \beta^{(j)}\}_{j=1}^J$ is obtained from the posterior distribution $p(\theta_1, \theta_2, \beta|D)$, all the marginal posterior distributions and the posterior **moments** for the θ_1, θ_2 , and β can be computed from the sample. Given the data D from the K accelerated environments, inferences at the use-stress environment S_u can be easily made by using the posterior sample. For example, the predictive reliability at the use-stress environment is given by

$$R(x_u|D) = \int R(x_u|\theta_1, \theta_2, \beta) \times p(\theta_1, \theta_2, \beta|D) d\theta_1 d\theta_2 d\beta \quad (13)$$

where

$$R(x_u|\theta_1, \theta_2, \beta) = e^{-\lambda_u x_u^\beta} \quad (14)$$

and λ_u is given by equation (3) or (4). The integral in equation (13) can be computed by using a Monte Carlo average using the posterior sample as

$$R(x_u|D) \simeq \frac{1}{J} \sum_{j=1}^J R(x_u|\theta_1^{(j)}, \theta_2^{(j)}, \beta^{(j)}) \quad (15)$$

which is the expected reliability at mission time x_u . We can also make probability statements about the reliability at mission time x_u by using a histogram estimate based on $R(x_u|\theta_1^{(j)}, \theta_2^{(j)}, \beta^{(j)})$ $j = 1, \dots, J$. Similarly, expected time to failure at the use stress environment can be obtained. The above setup can

be easily modified to incorporate different priors. For example, for the power law model we may assume a lognormal prior for θ_1 and a normal prior for θ_2 .

Extensions of the Basic Parametric Model

Earlier work by Blackwell and Singpurwalla [6] and Mazzuchi and Soyer [1] considered time transformation functions that change with the stress environment. These authors claimed that such a change is likely to happen owing to changes in the basic failure mechanism with changes in the stress level, and introduced non-Gaussian Kalman filter type models for ALT, and used approximate Bayesian methods for inference.

In what follows we present two models that describe changes in the time transformation function. The first model is a hierarchical Bayes model in the sense of Lindley and Smith [7] and assumes exchangeability of the parameters of the time transformation function. The second model is a first-order Markov model similar to what is considered by Mazzuchi and Soyer [1]. We note that in both cases the likelihood functions (6) and (8) will still be valid with changes reflected in λ_i . We focus on the power law model, but extension to the Arrhenius law is straightforward.

A Hierarchical (Exchangeable) Bayesian Model

To allow for changes in the time transformation function, we rewrite the power law model (3) as

$$\lambda_i = \theta_{1i} S_i^{\theta_2} \quad (16)$$

where the parameter θ_{1i} depends on the stress as reflected by the index i . Assume θ_{1i} s are conditionally independent with $\theta_{1i} \sim \Gamma[a, b]$, where a and b are random quantities with hyperprior $p(a, b)$. This implies that $\{\theta_{1i}\}, i = 1, \dots, K$, is an exchangeable sequence. If we further assume *a priori* independence of β and θ_2 , and of θ_{1i} s, we can recast the ALT model into a hierarchical Bayesian model as follows:

$$(X_i | \lambda_i, \beta) \sim \text{Weib}(\lambda_i, \beta) \quad (17)$$

$$\lambda_i = \theta_{1i} S_i^{\theta_2} \quad (18)$$

$$(\theta_2, \beta) \sim p(\theta_2) p(\beta) \quad (19)$$

$$(\theta_{1i} | a, b) \sim \Gamma[a, b] \quad (20)$$

$$(a, b) \sim p(a, b) \quad (21)$$

For any choice of priors in equations (17–21), the posterior inference cannot be performed analytically and thus we use a Gibbs sampler. The full conditionals for $\{\theta_{1i}\}$ are given by

$$(\theta_{1i} | a, b, \theta_1^{(-i)}, \theta_2, \beta, D) \sim \Gamma \left[a + r_i, b + S_i^{\theta_2} T_i(\beta) \right] \quad (22)$$

where $\theta_1^{(-i)} = \{\theta_{1j}; j \neq i\}$. Assuming $p(a, b) = p(a)p(b)$, and $b \sim \Gamma[c, d]$, the full conditional for b is obtained as

$$(b | \theta_1, \theta_2, \beta, a, D) \sim \Gamma \left[c + aK, d + \sum_{i=1}^K \theta_{1i} \right] \quad (23)$$

where $\theta_1 = (\theta_{1i}; i = 1, \dots, K)$. Similarly, the full conditional of a is

$$p(a | \theta_1, \theta_2, \beta, b, D) \propto \left[\frac{b^a}{\Gamma(a)} \prod_{i=1}^K (\theta_{1i})^{a-1} \right] p(a) \quad (24)$$

which is log concave as long as $p(a)$ is a log-concave density.

For θ_2 , we have

$$p(\theta_2 | \theta_1, \beta, a, b, D) \propto \left\{ \prod_{i=1}^K S_i^{r_i \theta_2} \right\} \exp \left[- \sum_{i=1}^K \theta_{1i} S_i^{\theta_2} T_i(\beta) \right] p(\theta_2) \quad (25)$$

which is log concave for a log-concave prior. Finally, for β , we get

$$p(\beta | \theta_1, \theta_2, a, b, D) \propto \beta^{\sum_{i=1}^K r_i} \left\{ \prod_{i=1}^K \left(\prod_{j=1}^{r_i} (x_{ij})^\beta \right) \right\} \times \exp \left[- \sum_{i=1}^K \theta_{1i} S_i^{\theta_2} T_i(\beta) \right] p(\beta) \quad (26)$$

which is also log concave for a log-concave prior $p(\beta)$. Thus, we can use the ARS method to sample from the conditionals for θ_2 and β .

Having obtained the posterior samples, we focus on the predictive reliability at the use stress

$$R(x_u | D) = \int R(x_u | \theta_{1u}, \theta_2, \beta) \times p(\theta_{1u}, \theta_2, \beta | D) d\theta_{1u} d\theta_2 d\beta \quad (27)$$

Given the posterior samples we can approximate equation (27) as

$$R(x_u|D) \simeq \frac{1}{J} \sum_{j=1}^J R(x_u|\theta_{1u}^{(j)}, \theta_2^{(j)}, \beta^{(j)}) \quad (28)$$

where $\theta_{1u}^{(j)}$ is generated from the gamma density $p(\theta_{1u}|\alpha^{(j)}, b^{(j)})$, conditional on the posterior realizations $(a^{(j)}, b^{(j)})$.

A Markov (Dynamic) Model

An alternate modeling strategy for capturing changes in the time transformation function is to consider a first-order Markov structure on $\alpha_i = \ln \theta_{1i}$ in equation (16). Since S_i and S_{i-1} are ordered neighboring stress environments, $S_i < S_{i-1}$, for $i = 1, \dots, K$, we assume that

$$(\alpha_i|\alpha_{i-1}, \phi, \tau) \sim N(\phi\alpha_{i-1}, \tau^{-1}) \quad (29)$$

where ϕ and τ are unknown quantities. Initially, we assume that $\alpha_0 \sim N(\mu_0, \tau^{-1})$. This is a generalization of the model in Mazzuchi and Soyer [1]. Note that equation (29) can be represented as a first-order autoregressive AR(1) dynamic model as

$$\alpha_i = \phi\alpha_{i-1} + \epsilon_i, \epsilon_i \sim N(0, \tau^{-1}) \quad (30)$$

The use stress can be considered as the $(K+1)$ st level in the above setup, that is, $\alpha_u = \alpha_{K+1}$. In cases where the stresses are not equally spaced, the precision τ can be discounted by a function $f(S_i, S_{i-1})$ as $\tau_i = \tau f(S_i, S_{i-1})$. A suitable choice for f is $f(S_i, S_{i-1}) = S_i/S_{i-1}$, where $S_i < S_{i-1}$.

Specification of the Markov model can be completed by assuming independent priors as $\tau \sim \Gamma(a, b)$ and $p(\phi)$ in the above. Although $-1 < \phi < 1$ would imply a stationarity of equation (30), it is not strictly necessary for our analysis, so we choose a normal prior for ϕ for convenience; $\phi \sim N(\mu_\phi, \tau_\phi^{-1})$. Finally, we assume independent priors on β and θ_2 . Thus, for $0 < i < K$, the full conditional of α_i is given by

$$\begin{aligned} p(\alpha_i|\alpha^{(-i)}, \theta_2, \phi, \beta, \tau, D) \\ \propto \exp \left[\alpha_i r_i - e^{\alpha_i} S_i^{\theta_2} T_i(\beta) \right. \\ \left. - \frac{\tau}{2} ((\alpha_i - \phi\alpha_{i-1})^2 + (\alpha_{i+1} - \phi\alpha_i)^2) \right] \quad (31) \end{aligned}$$

which is log concave. For $i = K$, that is, for the smallest stress environment the full conditional is obtained as

$$\begin{aligned} p(\alpha_K|\alpha^{(-K)}, \theta_2, \phi, \beta, \tau, D) \\ \propto \exp \left[\alpha_K r_K - e^{\alpha_K} S_K^{\theta_2} T_K(\beta) - \frac{\tau}{2} (\alpha_K - \phi\alpha_{K-1})^2 \right] \quad (32) \end{aligned}$$

which is also a log-concave density. For α_0 it can be shown that

$$\begin{aligned} (\alpha_0|\alpha, \theta_2, \phi, \beta, \tau, D) \\ \sim N \left[\frac{\mu_0 + \phi\alpha_1}{1 + \phi^2}, (\tau(1 + \phi^2))^{-1} \right] \quad (33) \end{aligned}$$

where $\alpha = \{\alpha_i; i = 1, \dots, K\}$, and for precision τ

$$(\tau|\alpha_0, \alpha, \theta_2, \phi, \beta, D) \sim \Gamma(a^*, b^*) \quad (34)$$

with parameters $a^* = a + (K+1)/2$ and $b^* = b + 1/2[(\alpha_0 - \mu_0)^2 + \sum_{i=1}^K (\alpha_i - \phi\alpha_{i-1})^2]$. The full conditional of ϕ is

$$\begin{aligned} p(\phi|\alpha_0, \alpha, \theta_2, \tau, \beta, D) \\ \propto \exp \left[-\frac{\tau}{2} \sum_{i=1}^K (\alpha_i - \phi\alpha_{i-1})^2 \right] p(\phi) \quad (35) \end{aligned}$$

If the prior on ϕ is normal, that is, $\phi \sim N(\mu_\phi, \tau_\phi^{-1})$, then the full conditional of ϕ will be a normal given by

$$\begin{aligned} (\phi|\alpha_0, \alpha, \theta_2, \tau, \beta, D) \\ \sim N \left[\frac{\tau_\phi \mu_\phi + \tau \sum_{i=1}^K (\alpha_i \alpha_{i-1})}{\tau_\phi + \tau \sum_{i=1}^K \alpha_{i-1}^2}, \left(\tau_\phi + \tau \sum_{i=1}^K \alpha_{i-1}^2 \right)^{-1} \right] \quad (36) \end{aligned}$$

For any other log-concave prior, the full conditional of ϕ is not available in closed form but is log concave.

The full conditionals of θ_2 and β are log concave and are given by

$$\begin{aligned} p(\theta_2|\alpha_0, \alpha, \phi, \tau, \beta, D) \\ \propto \left\{ \prod_{i=1}^K S_i^{r_i \theta_2} \right\} \exp \left[-\sum_{i=1}^K e^{\alpha_i} S_i^{\theta_2} T_i(\beta) \right] p(\theta_2) \quad (37) \end{aligned}$$

and

$$\begin{aligned}
& p(\beta|\alpha_0, \alpha, \theta_2, \tau, \phi, D) \\
& \propto \beta^{\sum_{i=1}^K r_i} \left\{ \prod_{i=1}^K \left(\prod_{j=1}^{r_i} (x_{ij})^\beta \right) \right\} \\
& \times \exp \left[- \sum_{i=1}^K e^{\alpha_i} S_i^{\theta_2} T_i(\beta) \right] p(\beta) \quad (38)
\end{aligned}$$

Once the posterior analysis is completed, we can obtain the predictive reliability at the use stress using equation (28), where $\theta_{1u} = \exp(\alpha_u)$ and $\alpha_u = \alpha_{K+1}$. Given the posterior samples, we can approximate the predictive reliability via equation (28) with $\theta_{1u}^{(j)} = \alpha_u^{(j)}$ and $\alpha^{(j)}$ is generated from the normal density $p(\alpha_u|\alpha_K^{(j)}, \phi^{(j)}, \tau^{(j)})$ for each posterior realization $(\alpha_K^{(j)}, \phi^{(j)}, \tau^{(j)})$.

Analysis of Accelerated Life Testing Data

In this section we illustrate the Bayesian analysis of the presented models using two accelerated life testing datasets, investigate the appropriateness of dynamic time transformation functions, and discuss model comparison. While doing so, we first discuss the deviance information criterion (*DIC*) for comparing our models.

Model Comparison

In comparing the performance of basic parametric model with hierarchical and Markov extensions presented above, computation of the Bayes factors (see Kass and Raftery [12] for a comprehensive review) is difficult. The marginal likelihoods for the competing models, which are needed to compute the Bayes factor, cannot be directly approximated from the Gibbs sampler.

An alternative approach is to use a model selection criterion such as the *DIC* of Spiegelhalter *et al.* [13]. For a generic parameter vector Θ , *DIC* is defined as

$$DIC = \bar{D} + p_D \quad (39)$$

where $D = -2 \log \mathcal{L}(\Theta)$, is two times the negative loglikelihood, $\bar{D} = E_{\Theta|\text{data}}[D]$ and $p_D = \bar{D} - D(\hat{\Theta})$, where $\hat{\Theta}$ is the posterior mean. The *DIC* has the general “fit+complexity” form used by many model selection criteria. In equation (39) $\bar{D}$ represents the “goodness of the fit” of the model where p_D represents a complexity penalty as reflected by the effective number of parameters of the model.

Example: Breakdowns of an Insulating Fluid

In what follows, we illustrate the use of our methodology to the ALT data published in Nelson [8]. The data is given in Table 1 and represents the times to

Table 1 Times to breakdown of an insulating fluid (in minutes) under various values of the stress

38 Kv	36 Kv	34 Kv	32 Kv	30 Kv	28 Kv	26 Kv
0.09	0.35	0.19	0.27	7.74	68.85	5.79
0.39	0.59	0.78	0.40	17.05	108.29	1579.52
0.47	0.96	0.96	0.69	20.46	110.59	2323.70
0.73	0.99	1.31	0.79	21.02	426.07	—
0.74	1.69	2.78	2.75	22.66	1067.6	—
1.13	1.97	3.16	3.91	43.40	—	—
1.40	2.07	4.15	9.88	47.30	—	—
2.38	2.59	4.67	13.95	139.07	—	—
—	2.71	4.85	15.93	141.12	—	—
—	2.90	6.50	27.80	175.88	—	—
—	3.67	7.35	53.24	194.90	—	—
—	3.99	8.01	82.85	—	—	—
—	5.35	8.27	89.29	—	—	—
—	13.77	12.06	100.58	—	—	—
—	25.50	31.75	215.10	—	—	—
—	—	32.52	—	—	—	—
—	—	33.91	—	—	—	—
—	—	36.71	—	—	—	—
—	—	72.89	—	—	—	—

breakdown of an insulating fluid subjected to various voltage levels. The accelerated stress levels are given by 26, 28, 30, 32, 34, 36, and 38 Kv and we are interested in making inference for the breakdown times for the insulating fluid at the use stress of 22 Kv. Following the analysis of Nelson [8] and Mazzuchi *et al.* [3], we assume a power law model for the time transformation function for the data.

We consider three models to analyze these data: the basic parametric model, the hierarchical Bayesian model, and the Markov model. Comparison among them will be made on the basis of the *DIC*. Due to our unfamiliarity with the problem, we will use diffuse priors in our analysis. In all models we choose a uniform prior in the range (0, 10) for the shape parameter β and a uniform prior in the range (0, 100) for θ_2 . In the parametric model of Section 2, the prior for θ_1 is $\Gamma[0.01, 0.01]$. In the hierarchical Bayes model, we use independent gamma priors for a and b in equations (17–21) with the parameters 0.01 and 0.01. In the Markov model, the prior for ϕ is uniform over $(-1, 1)$ and for τ it is $\Gamma[0.01, 0.01]$.

A comparison of the four models using *DIC* is given in Table 2. We note that the *DIC*s for the three models are all very similar with the parametric model having the lowest *DIC*. It is interesting to note that the hierarchical and Markov models estimate p_D , the effective number of parameters, as 4 or 5. Thus, the data do not seem to support the hypothesis that the time transformation function varies with the stress level.

This conclusion is also supported by Figure 1, where we present the posterior distribution θ_1 under the parametric model (denoted by P) and the posterior distributions of $\theta_{11}, \dots, \theta_{17}$ under the hierarchical model. We note that the posterior distributions of θ_{1j} s do not change much from one stress level to the other.

In Figure 2, we present the marginal posterior distributions of β , θ_1 , and θ_2 under the basic model. The posterior distribution of β is concentrated around 0.8, implying a decreasing failure rate at any stress level. The distribution of parameter θ_2 is concentrated

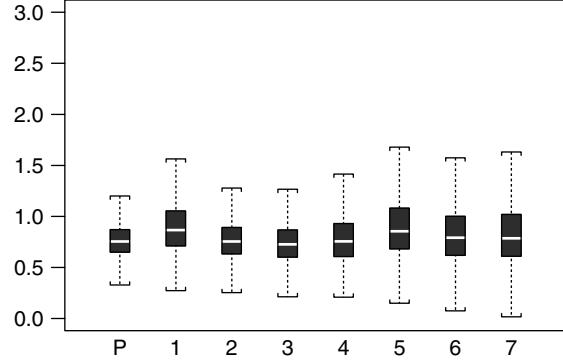


Figure 1 Comparison of posterior distributions of θ_1 and θ_{1j} s

around high positive values, implying that the failure rate increases with stress level as expected.

The predictive reliability function for the use stress $S_u = 22$ Kv is given in Figure 3 with 0.05 and 0.95 reliability bounds. We note that the items are quite reliable at the use-stress level as implied by high-reliability values for $X_u < 10\,000$, which is expected in ALTs.

Example: Rolling Bearing Data

The second dataset we consider is taken from Nelson [14, pp. 156] and it represents complete life data from a load ALT of rolling bearings. Four test loads were used and 10 bearings were tested at each of the test loads as illustrated in Table 3. The author suggested a Weibull model with a power law relationship. Life data is given in 10^6 revolutions. The data was analyzed using the three models and a power law time transformation function as in the section titled “Example: Breakdowns of an Insulating Fluid”.

A comparison of the three models using *DIC* is given in Table 4. Note that the *DIC*s under hierarchical and Markov models are considerably lower than that of the basic parametric model. The estimated effective number of parameters, p_D , is given as six under both of these models implying that there are four different θ_{1j} s. Thus, there is evidence that the time transformation function changes with stress level in this data. We can also see this by looking at the posterior distributions of θ_{1j} s in the hierarchical model. These are shown in Figure 4 and compared with the distribution of θ_1 in the parametric

Table 2 Model comparison using *DIC*

Model	<i>DIC</i>	p_D
Parametric model	607.85	3.33
Hierarchical Bayes model	609.77	4.57
Markov model	609.66	4.66

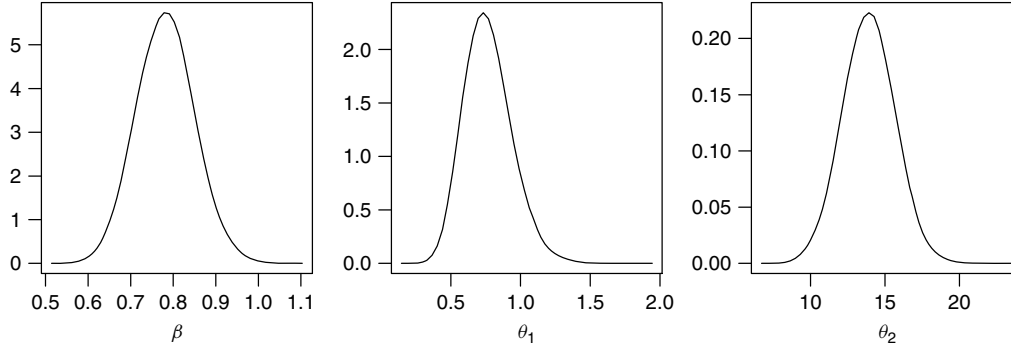


Figure 2 Posterior distributions of β , θ_1 , and θ_2 in the parametric model

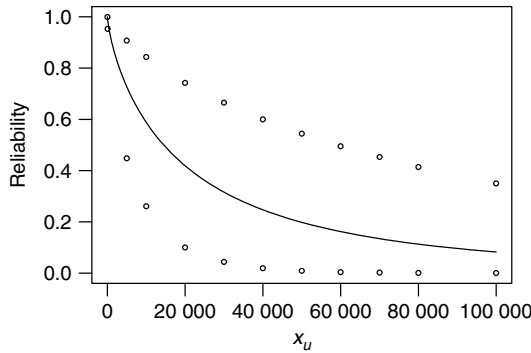


Figure 3 Predictive reliability function at $S_u = 22$ Kv

Table 3 Times to failure of roll bearings (in 10^6 revolutions) under various loads

0.87 Load	0.99 Load	1.09 Load	1.18 Load
1.67	0.80	0.012	0.073
2.20	1.00	0.18	0.098
2.51	1.37	0.20	0.117
3.00	2.25	0.24	0.135
3.90	2.95	0.26	0.175
4.70	3.70	0.32	0.262
7.53	6.07	0.32	0.27
14.70	6.65	0.42	0.35
22.76	7.05	0.44	0.386
37.40	7.37	0.88	0.456

model (denoted by P in Figure 4). The differences in the distributions are easily notable in this case. Similar results are obtained for the Markov model. Posterior distributions of the hyperparameters a and b are presented in Figure 5. We note that the posterior distributions for both of these parameters are very peaked even though the priors are flat.

Table 4 Model comparison using DIC

Model	DIC	p_D
Parametric model	114.82	3.02
Hierarchical Bayes model	109.08	5.59
Markov model	109.43	6.03

The posterior distributions of β and θ_2 for the hierarchical Bayes model are shown in Figure 6. The posterior distribution of β is concentrated around values greater than 1 implying increasing failure rate for the items at a given load. The posterior distribution of θ_2 is peaked around 20, showing a sharp increase in failure rate with increasing load.

In view of the above, the claim that the time transformation function may change with the stress environment is supported by some ALT data. On the basis of our analysis, it seems that both the hierarchical and dynamic models are able to capture such changes.

Concluding Remarks

In this article we have presented Bayesian analyses of two classes of parametric models for ALTs where the time transformation function changes with the stress environment. Our results show that such models may be more appropriate than the static models in some ALTs. In our setup, we have considered a power law model for the time transformation function and assumed a Weibull failure model. The approach can be easily extended to other time transformation functions and failure models. Alternatively, a semiparametric generalization of the models can be considered by relaxing the parametric assumptions on

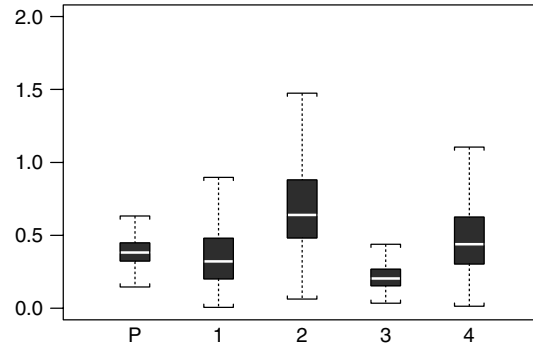


Figure 4 Posterior distributions of θ_1 and θ_{1j} s under the hierarchical model

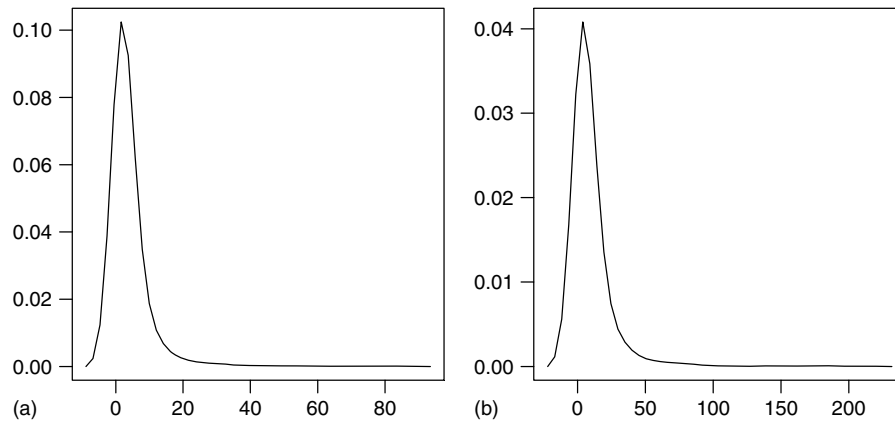


Figure 5 Posterior distributions of a and b under the hierarchical model

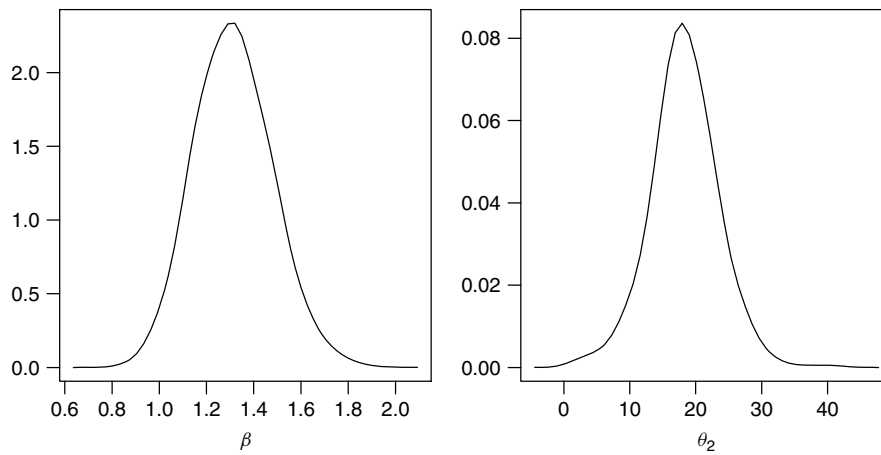


Figure 6 Posterior distributions of β and θ_2 under the hierarchical model

the distributions of the parameters of the time transformation function. Although semiparametric models have been used in Bayesian **survival analysis** (see for example, Kuo and Mallick [15]), consideration of these models are new in ALT literature. The semiparametric ALT models use mixtures of Dirichlet processes and inference for these involves the use of MCMC methods. Such work is presently under consideration.

References

- [1] Mazzuchi, T.A. & Soyer, R. (1992). A dynamic general linear model for inference from accelerated life testing, *Naval Research Logistics Quarterly* **39**, 757–773.
- [2] van Dorp, R., Mazzuchi, T.A., Fornell, G. & Pollock, L.R. (1996). A Bayes approach to step-stress accelerated life testing, *IEEE Transactions on Reliability* **R-45**, 491–498.
- [3] Mazzuchi, T.A., Soyer, R. & Vopatek, A.L. (1997). Linear Bayesian inference for accelerated Weibull model, *Lifetime Data Analysis* **3**, 1–13.
- [4] van Dorp, R. & Mazzuchi, T.A. (2004a). A general Bayes Weibull inference model for accelerated life testing, *Reliability Engineering and System Safety* **90**, 140–147.
- [5] van Dorp, R. & Mazzuchi, T.A. (2004b). A General Bayes exponential inference model for accelerated life testing, *Journal of Statistical Planning and Inference* **119**, 55–74.
- [6] Blackwell, L.M. & Singpurwalla, N.D. (1988). Inference from accelerated life tests using filtering in colored noise, *Journal of the Royal Statistical Society. Series B* **50**, 281–292.
- [7] Lindley, D.V. & Smith, A.F.M. (1972). Bayes estimates for the linear model (with discussion), *Journal of the Royal Statistical Society. Series B* **34**, 1–41.
- [8] Nelson, W.B. (1972). Graphical analysis of accelerated life test data with the inverse power law model, *IEEE Transactions on Reliability* **R-21**, 2–11.
- [9] Mann, N.R., Schafer, R.E. & Singpurwalla, N.D. (1974). *Methods for Statistical Analysis of Reliability and Life Data*, John Wiley & Sons, New York.
- [10] Gelfand, A.E. & Smith, A.F.M. (1990). Sampling-based approaches to calculating marginal densities, *Journal of the American Statistical Association* **85**, 398–409.
- [11] Gilks, W.R. & Wild, P. (1992). Adaptive rejection sampling for Gibbs sampling, *Applied Statistics* **41**, 337–348.
- [12] Kass, R.E. & Raftery, A.E. (1995). Bayes factors, *Journal of the American Statistical Association* **90**, 773–777.
- [13] Spiegelhalter, D.J., Best, N.G., Carlin, B.P. & van der Linde, A. (2002). Bayesian measures of model complexity and fit, *Journal of the Royal Statistical Society. Series B* **64**, 1–34.
- [14] Nelson, W.B. (1990). *Accelerated Life Testing*, John Wiley & Sons, New York.
- [15] Kuo, L. & Mallick, B.K. (1997). Bayesian semiparametric inference for the accelerated failure time model, *Canadian Journal of Statistics* **25**, 457–472.

REFIK SOYER, ALAATTIN ERKANLI AND JASON
R. MERRICK

Brainstorming

Introduction

Brainstorming is the term coined in the 1950s by advertising executive Alex Osborn [1] to describe a group process for generating ideas on a selected topic. It can be used whenever a diverse set of original ideas, approaches, theories, or points of view is desired. For example, brainstorming has been used in designing products and services [2], identifying topics for quality and reliability improvement efforts [3], collecting multiple theories about the causes of problems [4], generating potential solutions to problems or innovative changes to work flow [5], and engaging employees in improvement efforts [6].

Background, Principles, and Models

The psychologist Guilford [7] identified five traits associated with highly creative output in idea generation:

- *Fluency*: the rate of flow of ideas; roughly measured by the number of ideas generated in a given period of time.
- *Flexibility*: the ability to free-associate; roughly measured by the number of categories or themes represented in a stream of ideas.
- *Originality*: the ability to identify uncommon items; subjectively judged as the difference between the ideas and the usual approaches to the topic.
- *Awareness*: the ability to imagine beyond the immediate reality; subjectively judged as the ratio of conceptual to concrete content in the conversation within a group.
- *Drive*: Individuals' willingness in a group setting to take social risks and be unconventional in pursuit of a goal; subjectively judged by analysis of each individual's participation, and the absence of judgmental communication (verbal and non-verbal).

In each case, more is better. While subsequent studies [8] have added to Guilford's list, they also consistently reaffirm the importance of these five core traits.

Osborn [1] applied Guilford's research to the practical problem of generating ideas in an advertising agency. In order to create an environment that was conducive to creativity, Osborn suggested four rules:

- Criticism is ruled out.
- Freewheeling is welcomed; the wilder the ideas, the better.
- Quantity is wanted.
- Combination and improvement are sought.

Subsequent authors have variously restated or added to these (see [2] for example), but Osborn's basic rules remain the core of the brainstorming technique, and are widely used for group facilitation. Less well known today is Osborn's complete model for brainstorming (see Table 1). The power of brainstorming lies in the thoughtful application of the full approach.

Process

Craft a succinct statement of the issue

The issue statement should not be too broad, as this typically leads to only vague ideas. At the same time, the statement should not be so specific that it restricts thinking. For example, "How can we make our products and services more reliable?" is probably too broad and vague. Focusing on a specific product or service line, or on a few prominent failure modes, will likely produce more useful ideas. On the other

Table 1 Osborn's complete process for brainstorming

Fact finding

Problem definition: picking out and pointing up the problem

Preparation: gathering and analyzing the pertinent data

Idea finding

Idea production: generating a large number of thoughts and approaches to the problem

Idea development: selecting from among the resultant ideas and reprocessing then into even better ideas (by means of modification, combination, extraction of key concepts, and so on)

Solution finding

Evaluation: verifying the tentative solutions by means of tests and data collection

Adoption: deciding on and implementing the final solution

hand, “How do the policies of the purchasing department contribute to lack of reliability?” immediately shuts down other avenues of thinking that might lead to improved reliability.

Prepare background information and invite participants

Brainstorming will yield better results if all participants in the idea generation session know something about the issue. This might involve gathering data (*see Statistics in Data and Information Quality*) on current performance, knowledge of how others approach similar issues (*see Benchmarking in Project Definition*), or information about what has already been tried. Some experts further encourage formal observation in the field [2]. For example, in preparation for a brainstorming session to discuss possible causes of delayed shipments, a planning group might summarize data by product lines or geography, as well as identify themes uncovered in analysis of customer complaints and observations of the process.

Present this background information thoughtfully. The intent is to generally inform everyone and stimulate new insights. Avoid overwhelming or implying a specific line of thinking.

Idea generation sessions most often involve relatively small groups of 5–12 individuals, but can be conducted with groups numbering in the hundreds (with appropriate attention to equal participation and idea capture). Select participants from among the various stakeholders associated with the issue, along with a few individuals with uncommon backgrounds, skills, or points of view. For example, a session to identify new product ideas in a telecommunications company might include design engineers, managers, marketing personnel, customers (current and potential) – along with an anthropologist who can provide insight about social communication, or a child who simply often questions “why?”

Send the preparation information several days in advance and encourage participants to come to the idea generation session with a few ideas already in mind. This often results in participants noticing useful insights from chance observations prior to the session [9]. For example, a participant invited to the brainstorming session on delayed shipments described above might notice a geographic and seasonal pattern in the data and, while watching the television weather report, be stimulated to think about

integrating weather forecasts into the shipment routing system.

Prepare the physical space for the idea generation session

Idea generation works best in oversized rooms with flexible furnishings, where one can create an interactive and appropriately playful environment [2, 10].

Plan appropriate technology for recording ideas. This might include individual pads of paper, an easel pad for the group, audio or video recording equipment, or computers. The key factors to consider in selecting idea-recording technology are that it not inhibit active participation or slow down the rapid flow of ideas.

Begin the session with a brief review of the background information, an introduction to brainstorming methodology, and a warm-up exercise

The goal is to create a comfortable, nonthreatening, and information-rich atmosphere. Briefly go over the traits of creative thought and rules of brainstorming in order to clearly distinguish the session from other types of meetings. A 5–10 min warm-up exercise (for example, “How many different uses can we find for ordinary wire coat hangers?”) helps illustrate the desired atmosphere of fluency, flexibility, originality, awareness, and drive that one wants the group to apply to the issue that is the subject of the session.

Refocus the group on the issue statement and ask for ideas

Individual facilitators and organizations have a variety of preferences about how to orchestrate this group input. For example, some solicit group members in turn, while others allow them to voice ideas as they emerge. Some facilitators ask individuals to first write down several ideas in silence, and then share these with the group. Resources on brainstorming offer other variants [9, 11].

Be prepared to help the group through any lulls in the flow by either introducing an idea that opens up a new category, or injecting provocations. Osborn [1] supplied a list of verbs, like “magnify” or “substitute”, for use in quickly creating provocations. For example, one could ask, “How could we magnify certain failure modes if we desired, and what additional causes of the overall problem does that suggest that we haven’t yet thought about?” Others suggest removing taken for granted resources [12] or

creating an outrageous scenario that would demand fresh thinking [13]. For example, “What if we had to do this with no staff?” or “What if an earthquake destroyed our factory?” (see [5, 9, 11] and **Creativity Tools** for additional tools). Such provocations often reenergize the group and result in renewed fluency and originality.

Idea generation on the subject topic is typically limited to no more than 30–45 min. Additional topics can be handled in a similar period following a break.

Review, clarify, and initially process the ideas

The group can now eliminate redundant ideas, and obviously light-hearted ones. In the fast pace of the session, ideas might have been incompletely captured. These can now be embellished and clarified. If it seems useful, the group could categorize ideas. However, do not allow vague, high-level categories to take the place of the concrete and more specific ideas that were generated.

Conduct a formal harvesting of the best ideas

If the group has responsibility for acting on the ideas and time is available, immediately transition into selecting ideas for further work. For example, a quality improvement team brainstorming about possible causes of defects might select several theories and plan for **data collection**.

If action depends on others, or on the allocation of significant resources, the idea generation session should close with a brief description of how the ideas will be harvested by decision makers at a later time. For instance, ideas might be referred to a new product development committee, multidepartmental quality improvement council, or specific managers. Kotter and Cohen [14] speak of the importance of a guiding coalition of key decision makers who have access to the resources needed to bring about change. The odds of implementation are increased when these individuals are personally involved in selecting the ideas (see **Change Management, Stakeholder Assessment in**).

In either case, the harvesting group should first agree on the selection criteria. Examples include, cost, potential impact, immediately actionable *versus* longer term, or potential acceptance by employees. Any simple decision-making tool can then be used to select the ideas that best fit the criteria (see, for example [15]).

Plan for the further development, testing, and eventual implementation of ideas

The process of brainstorming is not complete until at least some of the ideas are acted upon (see Table 1). Organizations use a variety of project management (see **Project Management, Stage-Gate Approach to**) and improvement methodologies (see **Kaizen; Statistical Quality Control; Six Sigma**) to accomplish this. The key point is that failure to act on the ideas is likely to create an organizational culture that undermines future attempts at brainstorming.

Discussion

Research and Evaluation of Brainstorming

Experimental research on brainstorming tends to conclude that group brainstorming is no more effective, and in some studies is less effective, than idea generation by individuals [16].

However, Sutton and Hargadon [17] point out that these studies often fundamentally misrepresent how brainstorming is practically applied in organizations. The experimental study subjects, most often university students, typically have no prior or expected future relationship. They often have no particular interest or knowledge about the topic. Groups are often led by students, who are briefly presented the rules of brainstorming, and then asked to oversee 10–15 min of idea generation. Quantity of ideas is the most commonly measured result, with only a minority of studies assessing quality of ideas.

In contrast, observational research in organizations suggests that individuals in these settings bring a variety of subject matter expertise. Organizational groups most often have histories and expectations of future interactions, and care a great deal about the topics they are working on. Although quantity of ideas (fluency) is important within the process of idea generation, in the end, the effort only needs to yield a single successful idea to have a positive outcome for the organization. Sutton and Hargadon [17] provide case study evidence for the effectiveness of brainstorming, as practiced in organizational settings, and conclude that the experimental research, although rigorous and interesting, is largely irrelevant.

Strengths of Brainstorming

Group brainstorming is an efficient, effective, and flexible way to combine the knowledge and insights

of group members. It promotes **team building** by stressing equality among group members; deferral of judgment; and appreciation for the varying backgrounds, knowledge sets, and experiences of members.

Weaknesses of Brainstorming

An implicit assumption in brainstorming is that group members naturally have fresh insights and creative thoughts on the given topic and these simply need to be invited out. This might not always be the case. Group members who are immersed in “the way we have always done it” can become perceptually blocked from seeing in new ways [2, 6, 9]. When this occurs, a brainstorming session can generate a long list of items, none of which are significantly different from ideas that have been stated before. Additional tools for observation and creative thinking should be used in conjunction with the rules of brainstorming to create fresh, more creative insights (see, for example [9, 11–13] and **Creativity Tools**).

References

- [1] Osborn, A. (1953). *Applied Imagination*, Charles Scribner, New York.
- [2] Kelly, T. & Littman, J. (2001). *The Art of Innovation*, Doubleday, New York.
- [3] Plsek, P.E., Onnias, A. & Early, J.F. (1989). *Quality Improvement Tools*, Juran Institute, Inc., Wilton.
- [4] Pande, P.S., Neuman, R.P. & Cavanagh, R.R. (2000). *The Six Sigma Way: How GE, Motorola, and Other Top Companies are Honing Their Performance*, McGraw-Hill, New York.
- [5] Plsek, P.E. (1999). Innovative thinking for the improvement of medical systems, *Annals of Internal Medicine* **131**, 438–444.
- [6] Key, M.K. (ed) (1999). *Managing Change in Healthcare: Innovative Solutions for People-Based Organizations*, McGraw-Hill: New York.
- [7] Guilford, J.P. (1950). Creativity, *American Psychologist* **5**, 444–445.
- [8] Sternberg, R.J. (ed) (1999). *Handbook of Creativity*, Cambridge University Press, Cambridge.
- [9] Plsek, P.E. (1997). *Creativity, Innovation, and Quality*, Quality Press, Milwaukee.
- [10] Schrage, M. (2000). *Serious Play: How the World's Best Companies Stimulate Innovation*, Harvard Business School Press, Cambridge.
- [11] Higgins, J.M. (1994). *101 Creative Problem Solving Techniques*, New Management Publishing, Winter Park.
- [12] de Bono, E. (1992). *Serious Creativity*, Harper-Collins, New York.
- [13] Michalko, M. (1991). *Thinkertoys: A Handbook for Business Creativity for the 90s*, Ten Speed Press, Berkeley.
- [14] Kotter, J.P. & Cohen, D.S. (2002). *The Heart of Change: Real-Life Stories of How People Change Their Organizations*, Harvard Business School Press, Boston.
- [15] Scholtes, P.R., Joiner, B.L. & Streibel, B.J. (2003). *The Team Handbook*, 3rd Edition, Oriel, Milwaukee.
- [16] Mullen, B., Johnson, C. & Salas, E. (1991). Productivity loss in brainstorming groups: a meta-analytic integration, *Basic and Applied Social Psychology* **12**(1), 3–23.
- [17] Sutton, R.I. & Hargadon, A. (1996). Brainstorming groups in context: effectiveness in a product design firm, *Administrative Science Quarterly* **41**, 685–718.

PAUL E. PLSEK